

ENGFIS FIS

2022/23

Álgebra Linear e Geometria Analítica para Ciências

Salvatore Cosentino

Departamento de Matemática - Universidade do Minho

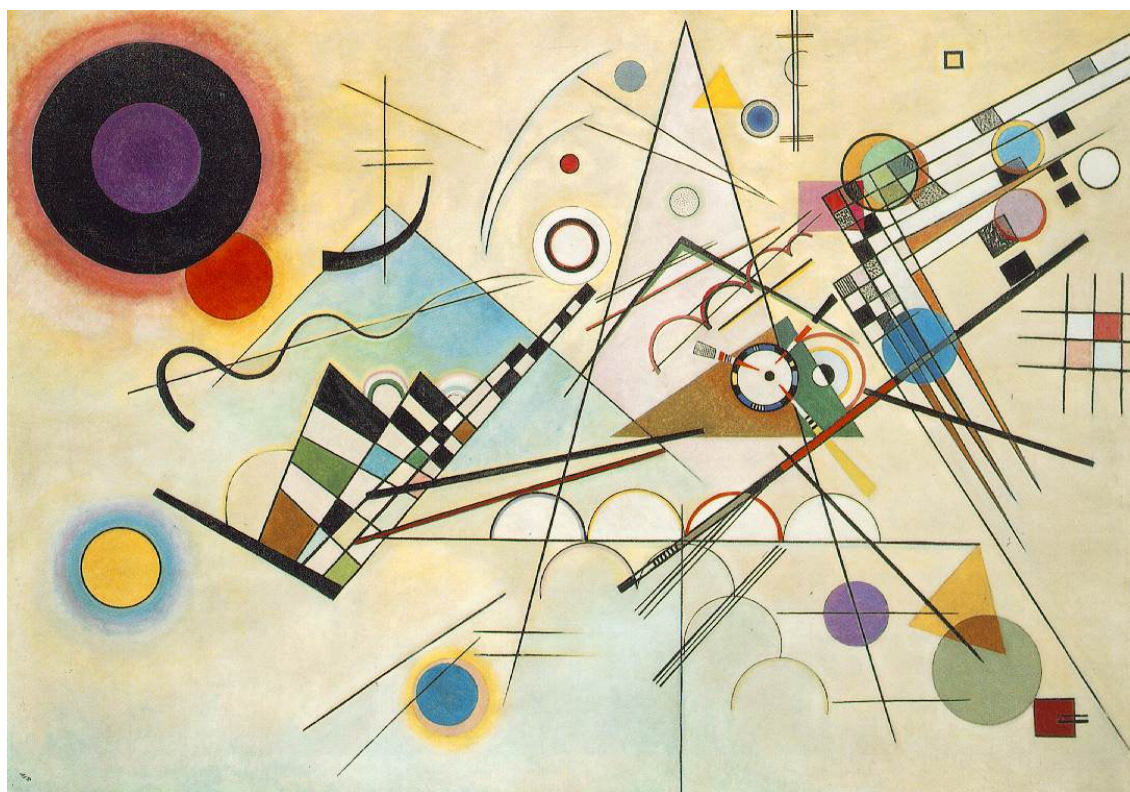
Campus de Gualtar - 4710 Braga - PORTUGAL

gab CG - Edifício 6 - 3.48, tel 253 604086

e-mail scosentino@math.uminho.pt

url <http://w3.math.uminho.pt/~scosentino>

12 de Janeiro de 2023



Resumo

This is not a book! These are notes written for personal use while preparing lectures on “Álgebra Linear e Geometria Analítica para Ciências” for students of FIS and ENGFIS during the a.y.’s 2021/22 and 2022/23. They are rather informal and certainly contain mistakes (indeed, they are constantly actualized). I tried to be as synthetic as I could, without missing the observations that I consider important.

Chapters correspond, at least in my intention during the first draft, to weeks, i.e. four hours lectures. Most probably I will not lecture all I wrote, and did not write all I plan to lecture. So, I included sketched or even empty paragraphs, about material that I think should/could be lectured within the same course, given enough time.

References contain some introductory manuals that I like, some classics, books where I have learnt things in the past century, recent books which I find interesting. Almost all material can be found in [Ap69], [La97] and [Ax15]. More advanced material may be found in [Ha58], [Po82] or [La87]. Those who love mathematics may look at what really algebra is about in the historical [Wa91] and then at the classic [MB99] or at the Bourbakist [Go96].

Everything about the course may be found in my web pages

<http://w3.math.uminho.pt/~scosentino/salteaching.html>

The notation is as follows:

e.g. means EXEMPLI GRATIA, that is, “for example”, and is used to introduce worked and important and (I hope!) interesting examples.

ex: means “exercise”, to be solved at home or in the classroom.

ref: means “references”, places where you can find and study what follows inside each section.

Black paragraphs form the main text.

Blue paragraphs deal with curiosities and ideas relevant in physics, engineering or other sciences. They may be the main reason why all this maths is worth studying for you. Some of them will probably only be understood and appreciated much later in your career.

Red paragraphs (mostly written in english) are more advanced or non trivial facts and results which may be skipped in a first (and also second) reading.

□ indicates the end of a proof.

Pictures were made with *Grapher*, *SketchBook* or *Paintbrush* on my MacBook, or taken from [Wikipedia](#), or produced with *Python* and *Matlab*.



This work is licensed under a
[Creative Commons Attribution-NonCommercial-ShareAlike 2.5 Portugal License](#).

Conteúdo

1	Vetores	6
1.1	O corpo dos números reais	6
1.2	Vetores	14
1.3	O espaço $\mathbb{R}^n$	17
1.4	Translações e homotetias	18
2	Produto escalar, norma e distância	21
2.1	Produto escalar euclidiano	21
2.2	Norma euclidiana	22
2.3	Geometria euclidiana elementar	24
3	Retas e planos	28
3.1	Equações paramétricas e cartesianas	28
3.2	Retas	29
3.3	Planos	33
3.4	Segmentos e convexos	37
4	Subespaços, bases e dimensão	40
4.1	Subespaços e geradores	40
4.2	Conjuntos livres	42
4.3	Bases e dimensão	43
4.4	Bases e ortogonalidade	45
5	Produto vetorial, área e volume	47
5.1	Independência no plano e determinante.	47
5.2	Produto vetorial	48
5.3	Produto misto	51
6	Números complexos	56
6.1	O corpo dos números complexos	56
6.2	Representação polar e raízes	61
6.3	Polinômios e fatorização	65
7	Espaços lineares	68
7.1	Espaços lineares	68
7.2	Espaços de funções	70
7.3	Independência, bases e dimensão	73
7.4	Produtos e quocientes	77
8	Transformações lineares	81
8.1	Formas lineares e espaço dual	81
8.2	Núcleo e hiperplanos	83
8.3	Transformações lineares	87
8.4	Núcleo e imagem	89
8.5	Álgebra dos operadores	91
8.6	Transformações lineares invertíveis	95
9	Transformações lineares e matrizes	98
9.1	Matrizes	98
9.2	Matrizes e transformações lineares	102
9.3	Composição e produto de matrizes	104
9.4	Transposta e dualidade	105

10 Operadores e matrizes quadradas	109
10.1 Álgebra das matrizes quadradas	109
10.2 Exemplos geométricos	111
10.3 Matrizes invertíveis	114
10.4 Mudança de bases e matrizes semelhantes	117
10.5 Traço	120
11 Sistemas lineares	122
11.1 Sistemas lineares no plano e no espaço	122
11.2 Sistemas lineares	124
11.3 Eliminação de Gauss-Jordan	127
12 Volume e determinantes	134
12.1 Formas alternadas e volumes	134
12.2 Determinante	137
12.3 Fórmula de Laplace	142
12.4 Determinante de um operador	144
12.5 Regra de Cramer	146
13 Valores e vetores próprios	148
13.1 Subespaços invariantes	148
13.2 Valores e vetores próprios	150
13.3 Operadores diagonalizáveis	155
13.4 Polinómio caraterístico	156
13.5 Diagonalização de matrizes	159
14 Estrutura dos operadores	164
14.1 Espaços próprios generalizados	164
14.2 Decomposição de Jordan-Chevalley	166
14.3 Blocos de Jordan	169
14.4 Complexificação	173
14.5 Forma normal de Jordan real	174

Notações

Conjuntos. $a \in A$ quer dizer que a é um elemento do conjunto A . $A \subset B$ quer dizer que o conjunto A é um subconjunto do conjunto B . $A \cap B$ é a interseção dos conjuntos A e B , e $A \cup B$ é a reunião dos conjuntos A e B . $A \times B$ é o produto cartesiano dos conjuntos A e B , o conjunto dos pares ordenados (a, b) com $a \in A$ e $b \in B$.

Funções. Uma *função* $f : X \rightarrow Y$, com *domínio* o conjunto X e *conjunto de chegada* o conjunto Y , é um subconjunto $R \subset X \times Y$ tal que para cada $x \in X$ existe um único $y := f(x) \in Y$, dito *imagem* de x , tal que $(x, y) \in R$. Quando domínio e contradomínio são claros, uma função pode ser denotada apenas por $x \mapsto f(x)$, ou seja, identificada com a “regra” que determina $y = f(x)$ a partir de x . A *imagem* do subconjunto $A \subset X$ é o conjunto $f(A) := \{f(a) \text{ com } a \in A\} \subset Y$. Em particular, a *imagem/contradomínio* da função $f : X \rightarrow Y$ é o conjunto $f(X) := \{f(x) \text{ com } x \in X\} \subset Y$ dos valores da função. O *gráfico* da função $f : X \rightarrow Y$ é o subconjunto

$$\text{Graph}(f) := \{(x, y) \in X \times Y \text{ t.q. } y = f(x)\} \subset X \times Y$$

do produto cartesiano do domínio e o conjunto de chegada. A função *identidade* $\mathbf{1}_X : X \rightarrow X$ é definida por $\mathbf{1}_X(x) = x$, e o seu gráfico é a *diagonal* $\{(x, x) \text{ com } x \in X\} \subset X \times X$.

A *restrição* da função $f : X \rightarrow Y$ ao subconjunto $A \subset X$ é a função $f|_A : A \rightarrow Y$ definida por $f|_A(a) := f(a)$ se $a \in A$.

A *composição* das funções $f : X \rightarrow Y$ e $g : f(X) \subset Y \rightarrow Z$ é a função $g \circ f : X \rightarrow Z$ definida por $(g \circ f)(x) := g(f(x))$, ou seja,

$$x \mapsto y = f(x) \mapsto z = g(y) = g(f(x))$$

Uma função $f : X \rightarrow Y$ é *injetiva* se $x \neq x'$ implica $f(x) \neq f(x')$, e portanto a imagem $f(X)$ é uma “cópia” de X . De fato, uma função injetiva admite uma *inversa esquerda*, uma função $g : f(X) \rightarrow X$ tal que $g(f(x)) = x$ para todo $x \in X$. Uma função $f : X \rightarrow Y$ é *sobrejetiva* se todo $y \in Y$ é imagem $y = f(x)$ de algum $x \in X$, ou seja, se $Y = f(X)$. Uma função $f : X \rightarrow Y$ é *bijetiva* se é injetiva e sobrejetiva, e portanto admite uma função *inversa* $f^{-1} : Y \rightarrow X$, que verifica $f^{-1}(f(x)) = x$ e $f(f^{-1}(y)) = y$ para todos os $x \in X$ e $y \in Y$.

1 Vetores

ref: [Ap69] Vol 1, 12.1-4 ; [La97] Ch. I, 1-2

A linguagem da filosofia. “... Signor Sarsi, la cosa non istà così. La filosofia è scritta in questo grandissimo libro che continuamente ci sta aperto innanzi agli occhi (io dico l’universo), ma non si può intendere se prima non s’impara a intender la lingua, e conoscer i caratteri, ne’ quali è scritto. Egli è scritto in lingua matematica, e i caratteri son triangoli, cerchi, ed altre figure geometriche, senza i quali mezzi è impossibile a intenderne umanamente parola; senza questi è un aggirarsi vanamente per un oscuro laberinto.”¹

1.1 O corpo dos números reais

A reta real $\mathbb{R}$ é um corpo ordenado e completo.

A reta real. Fixada uma origem (ou seja, um ponto 0), um unidade de medida (ou seja, a distância entre 0 e 1) e uma orientação (ou seja, uma direção “positiva”), é possível representar cada ponto de uma reta com um número real $x \in \mathbb{R}$. Vice-versa, ao número $x \in \mathbb{R}$ corresponde o ponto da reta colocado a uma distância $\sqrt{x^2}$ da origem, na direção positiva se $x > 0$ e negativa se $x < 0$.

Esta é, pelo menos, a nossa ideia “intuitiva”. A reta real pode ser definida, axiomáticamente, como um “corpo ordenado e completo”. Se compreendem esta frase, podem passar à seguinte subseção. Para conveniência do leitor, lembramos brevemente estas noções, omitindo provas e construções. Umas ótimas referências são a introdução de [Ap69] e o capítulo III de [Li95].

Counting. We count finite collections of similar objects (as fingers in our hand, years, molecules in a mole of gas, baryons in the Universe) using the numbers

$$1, 2, 3, 4, 5, \dots, 33, \dots, 6 \times 10^{23}, \dots, 10^{80}, \dots$$

We may “sum” 33 goats and 66 goats, to get a flock of $33 + 66 = 99$ goats. Also, we may need a surface of $23 \times 23 = 529$ square meters to build our pyramid with side of 23 meters. Conversely, we may sell 2 of our 99 goats and stay with the remaining flock of $99 - 2 = 97$ goats. Or we may store the visible mass $\sim 4 \times 10^{41}$ kg of the Milky Way into $\sim (4 \times 10^{41}) / (2 \times 10^{30}) = 2 \times 10^{11}$ stars of the same size of our Sun (estimated to be $\sim 2 \times 10^{30}$ kg).

Peano axioms for the natural numbers. We use the notation $\mathbb{N} := \{1, 2, 3, 4, 5, \dots\}$ for the set of *natural numbers*. It is convenient to define $\mathbb{N}$ by a (minimal) set of “axioms”, and this is what Giuseppe Peano² did back in 1889:

- N1** any natural $n \in \mathbb{N}$ has a “successor” $n^+ \in \mathbb{N}$ (which, a posteriori, we think as $n + 1$), different from n , and no two different naturals have the same successor;
- N2** there is an element, called “one” and denoted by $1 \in \mathbb{N}$, which is not the successor of any natural;
- N3** (*induction principle*) a subset $A \subset \mathbb{N}$ which contains 1 and such that $n \in A$ implies $n^+ \in A$ is the whole $\mathbb{N}$.

The third axiom is the key to prove that certain statements about numbers are valid for all naturals (since we humans have no time to check for all of them!). It is also the property that makes possible recursive definitions, as we’ll see soon.

Once accepted the axioms, we set $2 := 1^+$, $3 := 2^+$, $4 := 3^+$, ... and so on (but of course any other list of symbols, as $\clubsuit, \spadesuit, \heartsuit$... would do).

¹Galileo Galilei, *Il Saggiatore*, 1623.

²G. Peano, *Arithmetices principia, nova methodo exposita*, 1889.

We define *sums* inductively, starting from $n + 1 := n^+$, and setting $n + (m^+) := (n + m)^+$. The sum of two numbers represents a cardinality of an union. For example, $3 + 4 = 7$ means

$$\boxed{\bullet \bullet \bullet} + \boxed{\bullet \bullet \bullet \bullet} = \boxed{\bullet \bullet \bullet \bullet \bullet \bullet \bullet}$$

We define *products* inductively, starting from $n \cdot 1 = n$, and setting $n \cdot (m^+) := n \cdot m + n$. Thus, $d \cdot a$ is the sum of d times a , i.e.

$$d \cdot a = \underbrace{a + a + \cdots + a}_{d \text{ times}}$$

and actually represents an “area”. For example, $4 \cdot 3 = 12$ means

$$\boxed{\bullet \bullet \bullet} \times \boxed{\bullet \bullet} = \boxed{\begin{array}{ccc} \bullet & \bullet & \bullet \\ \bullet & \bullet & \bullet \end{array}}$$

If $a + b = c$, we say that b is the *difference* between c and a , and write $b = c - a$. Thus, for example, $7 = 13 - 6$.

If $q \cdot r = p$, we say that r is the *ratio* between p and q , and write $r = \frac{p}{q}$ or p/q . Thus, for example, $3 = 21/7$.

e.g. Triangular numbers. The sum of the first n naturals is given by the formula

$$1 + 2 + 3 + \cdots + n = \frac{n(n+1)}{2},$$

which you may conjecture summing the last with the first numbers (hence $n + 1$), then the last and the first of what remain (again $n - 1 + 2 = n + 1$), and so on, up to a total of $n/2$ such pairs, or, following the Greeks, observing the following picture (red bullets form the “gnomon”):



If you are not satisfied with that, you check the formula for $n = 1$ (this gives $1 = 1 \cdot 2/2$), assume it holds for n , sum the next term, which is $n+1$, and verifies that $n(n+1)/2 + (n+1) = (n+1)(n+2)/2$.

Observe that for large n this grows as

$$1 + 2 + 3 + \cdots + n \sim \frac{n^2}{2}$$

ex: Show that the sum of the first n odd numbers is

$$1 + 3 + 5 + 7 + \cdots + (2n - 1) = n \cdot n$$

(i.e. n^2 , but we have not introduced this notation yet!), as the following picture suggests (again, red bullets form the “gnomon”):



and guess a formula for the sum of the first n even numbers

$$2 + 4 + 6 + \cdots + 2n$$

ex: Guess the first significant term in the sum

$$1 + 4 + 9 + 16 + 25 + \cdots + n^2$$

of the first n square numbers (look at the picture above, and put one square over another ...).

Well-ordering principle. We may define an order in $\mathbb{N}$ saying that $n < m$ (“ n is *smaller* than m ”) if there exists $x \in \mathbb{N}$ such that $n + x = m$. We say that $n \leq m$ (“ n is *not greater* than m ”) if $n < m$ or $n = m$. This relation is stable under sums and products: if $n \leq m$ then also

$$n + x \leq m + x \quad \text{and} \quad n \cdot x \leq m \cdot x$$

for all $x \in \mathbb{N}$. It is clear that 1 is the “smallest” of all the numbers, i.e. $1 \leq x$ for all $x \in \mathbb{N}$. The induction principle N3 is equivalent to the statement that any subset of the naturals has a smallest element:

WO (*well-ordering principle*) every subset $A \subset \mathbb{N}$ has a first element (or minimum), i.e. an element $a \in A$ such that $a \leq x$ for all $x \in A$.

Integers. It turns out (but this took quite a large time to mankind!) that even elementary problems are solved with much ease if we enlarge our numbers allowing *negative* numbers, like -237 , hence a *zero* number, that we denote 0. The set thus obtained is the set of *integer numbers*

$$\mathbb{Z} := \{ \dots, -3, -2, -1, 0, 1, 2, 3, \dots \}.$$

(from the german *zahlen* = numbers) The two operations, $+$ and $\times$ (but we’d rather use “dots” for multiplication, like in $7 \cdot 3 = 21$, or even nothing when there is no possible confusion, like in $ab = a \cdot b$) are then characterised by the following properties, which define a structure that mathematicians call *commutative ring*.

R1 (*associativity* of both $+$ and $\times$) $(x + y) + z = x + (y + z) \quad (x \cdot y) \cdot z = x \cdot (y \cdot z)$

R2 (*commutativity* of both $+$ and $\times$) $x + y = y + x \quad x \cdot y = y \cdot x$

R3 (existence of *neutral elements* for both $+$ and $\times$) there exist two element, 0 and 1, called *neutral element* for the sum and for the multiplication, respectively, such that $x + 0 = x$ and $x \cdot 1 = x$

R4 (existence of the *inverse* for $+$) for any x there exists x' , called *additive inverse* or *opposite* of x , such that $x + x' = 0$

R5 (*distributive law*) $x \cdot (y + z) = x \cdot y + x \cdot z$

It is plain that all primary school arithmetical rules may be derived from these properties/axioms (but you should try to prove them by yourself!). To start with, you may derive the useful rule

$$a + x = a + y \quad \Rightarrow \quad x = y \tag{1.1}$$

summing to both sides of the first equality the additive inverse of a , using R1 and R3. In particular, this implies that the additive inverse in R4 is unique, for if $x + x' = 0 = x + x''$, then by the above rule (1.1) we get $x' = x''$.

This also implies that *subtraction* is possible: given a and b , there exists a unique x such that

$$a + x = b$$

(just take $x = b + a'$, where a' is the opposite of a , and check that it solves the problem, uniqueness being a consequence of (1.1)). Such difference is conveniently denoted by $x = b - a$. In particular, the opposite of a is $0 - a$, and should be better denoted simply by $-a$. Finally, one easily check the following rules concerning the minus sign:

$$b - a = b + (-a) \quad -(-a) = a \quad a(b - c) = ab - ac$$

It is useful to have a short notation for repeated sums

$$\sum_{n=1}^N x_n := x_1 + x_2 + \dots + x_N$$

and products

$$\prod_{n=1}^N x_n := x_1 \cdot x_2 \cdot \dots \cdot x_N$$

This is possible thanks to the associativity of both sums and products. The product of the first n naturals is ubiquitous when counting cardinalities, and deserves a name: it is called *n factorial*, and denoted by

$$n! := 1 \cdot 2 \cdot 3 \cdot \dots \cdot (n-1) \cdot n.$$

It is convenient to have a short notation for repeated products of a fixed number. For example, the product $x \cdot x$ is said “ x squared”, and denoted by x^2 (for, if $x > 0$, it is the area of a square with side x). Similarly, $x \cdot x \cdot x$ is said “cube of x ”, and denoted by x^3 (if $x > 0$, it is the volume of a cube with side x). For integer $n = 1, 2, 3, \dots$, the n -th power of the (rational) number x is defined by

$$x^n := \underbrace{x \cdot x \cdot \dots \cdot x}_{n \text{ times}}$$

(to be pedant, recursively according to $x^1 := x$ and $x^{n+1} := x^n \cdot x$ for $n \geq 1$). It is useful to set $x^0 := 1$, so that the law of exponents $x^n \cdot x^m = x^{n+m}$ holds for all exponents $m, n \geq 0$.

Division, factorization and primes. We say that the number a *divides* (or *is a divisor of*) the number b , and we write $a \mid b$, if there exists a $d \in \mathbb{N}$ such that $ad = b$. If a does not divide b , we write $a \nmid b$. Given any $p \leq q$, either $p \mid q$, so that $q = dp$ for some $d \in \mathbb{N}$, or there exist a unique $d \in \mathbb{N}$ and a unique “rest” $0 < r < q$ such that

$$q = dp + r.$$

We say that a natural number p is *prime* if it is not divided by any other natural but 1 and itself. Thus, 2, 3, 5, 7, 11, 13, ... are primes. It is a fundamental fact of arithmetic (which you could find in the classic book by Hardy and Wright ³) that any natural n can be uniquely factorised (up to order!) into prime factors, i.e. written as $n = p_1^{n_1} p_2^{n_2} \dots p_k^{n_k}$ for some primes p_i and exponents $n_i \in \mathbb{N}$. Thus, primes are the building blocks with which all naturals are constructed.

Here is Proposition 20 of Book IX of the *Elements* by Euclid:

Οἱ πρῶτοι ἀριθμοὶ πλείους εἰσὶ παντὸς τοῦ προτεθέντος πλήθους πρώτων ἀριθμῶν ⁴

or, in modern language,

Teorema 1.1 (Euclid). *The set of prime numbers is not finite.*

Indeed, following Euclid, assume that $p_1, p_2, \dots, p_n$ are all the primes. We could take their product and sum one, i.e. form the number $x = p_1 p_2 \dots p_n + 1$, and observe that x is not divisible by any of the p_k , since the rest of the division is always 1. Since x is larger than any of the p_k , it must have a prime divisor larger than all of them.

Clock arithmetics. Less obvious is that there exist other commutative rings. For any integer $n \geq 2$, we may equip the quotient $\mathbb{Z}/n\mathbb{Z} := \{k + n\mathbb{Z}, \text{ with } k \in \mathbb{Z}\} \approx \{0, 1, \dots, n-1\}$ with the obvious ring structure inherited from $\mathbb{Z}$. Thus,

$$(a + n\mathbb{Z}) + (b + n\mathbb{Z}) = a + b + n\mathbb{Z} \quad \text{and} \quad (a + n\mathbb{Z}) \cdot (b + n\mathbb{Z}) = a \cdot b + n\mathbb{Z}.$$

Such rings may have divisors of zero. Indeed, if n is not prime, but a product $n = pq$, then $(p + n\mathbb{Z}) \cdot (q + n\mathbb{Z}) = 0 + n\mathbb{Z}$.

³G.H. Hardy and E.M. Wright, *An Introduction to the Theory of Numbers*, fourth edition, Oxford University Press, 1959.

⁴“Prime numbers are more than any assigned multitude of prime numbers” [Euclid, *Elements*, Book IX, Proposition 20].

Ratios and proportions. Integers are not enough to measure the relative sizes of different quantities. As explained by Euclid in Book V of his *Elements*, we may need to consider *ratios* $x : y$ between couples of magnitudes of the same type. Equalities between two ratios are then called *proportions*. For example, two “commensurable” magnitudes (i.e. magnitude that can be measured by a same quantity/unit) x and y are in the ratio $7 : 3$ if there is a smaller quantity z such that $x = 7z$ and $y = 3z$ or, equivalently, if $3 \cdot x = 7 \cdot y$. Thus, if the recipe of a cake for 4 persons uses 6 eggs, and you need the same cake for 7 guests, you must use a number x of eggs which solves the proportion $6 : 4 = x : 7$ (but Greeks didn’t like ratios between quantities of different kind!).

Rationals and the four operations. We may form quotients p/q of integer numbers with “denominator” q not being 0 (but to be pedantic we should speak about “ordered pairs” of integers (p, q) ...), and define their sum and product as

$$\frac{a}{b} + \frac{c}{d} := \frac{ad + bc}{bd} \quad \text{and} \quad \frac{a}{b} \cdot \frac{c}{d} := \frac{ca}{bd}$$

Of course we interpret $a/1$ as a . In addition to the properties of a commutative ring, axioms R1-R5, the set of fractions also satisfies the following axiom

F6 (existence of the *inverse* for $\times$) $\forall x \neq 0$ there exists x' such that $x \cdot x' = 1$

Indeed, the inverse of a non-zero fraction p/q is simply q/p . A set equipped with two binary operations satisfying properties R1-5 and F6 is called a *field*. The set of fractions is therefore called *rational field*, and denoted by $\mathbb{Q}$.

A first important consequence of F6, together with the R1-5, is the rule

$$\lambda x = \lambda y \quad \text{and} \quad \lambda \neq 0 \quad \Rightarrow \quad x = y. \quad (1.2)$$

as follows multiplying both sides of the first equality by the multiplicative inverse of λ and using associativity R1 and F6. In particular, this implies that multiplicative inverses are unique. Indeed, if x' and x'' are two multiplicative inverses of the same $x \neq 0$, then $xx' = 1 = xx''$, and therefore, by the above rule, $x' = x''$.

This also implies that *division* is possible: given $a \neq 0$ and b , there exists a unique x such that

$$ax = b$$

(just take $x = ba'$, where a' is the multiplicative inverse of a , and check that it solves the problem, uniqueness being a consequence of (1.2)). Such x is called *quotient* between b and a , and denoted by b/a . In particular, the unique multiplicative inverse of $a \neq 0$ is then denoted as a^{-1} or $1/a$. Finally, one checks the following rules concerning division:

$$b/a = b \cdot a^{-1} \quad (a^{-1})^{-1} = a \quad (\text{if } a \neq 0)$$

One also check that a field has no divisors of zero: if $ab = 0$ then either $a = 0$ or $b = 0$.

Given $x \neq 0$, we also define negative powers according to $x^{-n} := 1/x^n$, for $n = 1, 2, 3, \dots$. Then, for all $n, m \in \mathbb{Z}$, and all $x \neq 0$, we have the law of exponents $x^n \cdot x^m = x^{n+m}$.

Linear equations. In a field, like the rationals (or, as you will see, the reals), we are able to solve a first degree equation like

$$ax + b = 0$$

(as usual, the notation above means that we are given the numbers a and b , and we want to find possible values for the “unknown” x). Indeed, we simply put b on the right hand side, multiplying by -1 , and then divide by a (the case $a = 0$ being trivial: it is no equation at all!). The solutions, which is obviously unique, is

$$x = -b/a.$$

Binomial formulas. The square of a binomial is

$$(a \pm b)^2 = a^2 \pm 2ab + b^2$$

Also,

$$(a + b)(a - b) = a^2 - b^2$$

Less obvious is that one can give a formula for the n -th power of a binomial. This has been found by Newton, and is

$$(a + b)^n = \sum_{k=0}^n \binom{n}{k} a^k b^{n-k}$$

where the *binomial coefficient* is defined as

$$\binom{n}{k} := \frac{n!}{k!(n-k)!}$$

Finite fields. There exist finite fields, i.e. fields formed by only a finite number of elements. Examples are $\mathbb{Z}_p := \mathbb{Z}/p\mathbb{Z}$ when p is a prime number. The smallest is $\mathbb{Z}_2$, the field formed by just two elements 0 and 1, equipped with the arithmetic rules: $0+0=0$, $0+1=1$, $1+1=0$, $0 \cdot 0=0$, $0 \cdot 1=1$ and $1 \cdot 1=1$.

Order. The field of rationals is an ordered field, i.e. may be “ordered”. This means that we may define a subset $\mathbb{Q}^+ := \{p/q \text{ with } p, q \in \mathbb{N}\}$ of *positive* rationals satisfying the “axioms of order”

O1 $0 \notin \mathbb{Q}^+$,

O2 if $a, b \in \mathbb{Q}^+$, then also $a + b \in \mathbb{Q}^+$ and $a \cdot b \in \mathbb{Q}^+$,

O3 if $x \neq 0$, then either $x \in \mathbb{Q}^+$ or $-x \in \mathbb{Q}^+$ (and not both!).

We then define $\mathbb{Q}^- := \mathbb{Q} \setminus (\mathbb{Q}^+ \cup \{0\})$, the set of *negative* rationals. We say that $a < b$ if there exists a $c \in \mathbb{Q}^+$ such that $a + c = b$. We say that $a > b$ if $b < a$. In particular, all $a \in \mathbb{Q}^+$, as for example 1, are $a > 0$, and all $b \in \mathbb{Q}^-$ are $b < 0$. We also say that $a \leq b$ if $a < b$ or $a = b$, and then that $a \geq b$ if $a \leq b$.

A first consequence is that for all couple of numbers a and b we have a trichotomy: either $a < b$, or else $a = b$ or else $a > b$. Another consequence is the following transitivity:

$$a < b \quad \text{and} \quad b < c \quad \Rightarrow \quad a < c$$

Clearly,

$$a < b \quad \Rightarrow \quad a + c < b + c$$

and also

$$a < b \quad \text{and} \quad c < d \quad \Rightarrow \quad a + c < b + d$$

Moreover,

$$a < b \quad \Rightarrow \quad \begin{cases} ad < bd & \text{if } d > 0 \\ ad > bd & \text{if } d < 0 \end{cases}$$

In particular,

$$a < b \quad \Rightarrow \quad -b < -a$$

Also, if $ab > 0$, then a and b are either both positive or both negative. Finally,

$$a \neq 0 \quad \Rightarrow \quad a \cdot a > 0$$

i.e. squares of non-zero numbers are positive. In particular, $1 > 0$.

ex: Prove, by induction, the *Bernoulli inequality*: for any $n = 1, 2, 3, \dots$ and any $x > -1$

$$(1 + x)^n \geq 1 + nx$$

Pythagora's theorem. Take a right triangle, set to 1 the length of the hypotenuse, and call α and β the lengths of the other sides. The altitude from the vertex opposed to the hypotenuse divides the latter into two pieces of lengths α^2 and β^2 , because they are sides of right triangles similar to the first one, having hypotenuses the two sides of length α and β , respectively. Therefore,

$$\alpha^2 + \beta^2 = 1.$$

For example, the diagonal ℓ of a square with unit side satisfies $1 + 1 = \ell^2$.

Babylonians-Heron method to compute square roots. Consider the problem to find the side ℓ of a square given its area $a > 0$, that is, the number which we modern call $\ell = \sqrt{a}$. A method, described by Heron of Alexandria⁵, but most probably already known to the Babylonians^{6 7}, consists in constructing recursively rectangles with fixed area a and sides which are nearer and nearer. If x_1 and y_1 are the base and the height of the first rectangle (chosen arbitrarily!), and therefore $x_1 y_1 = a$, then the second rectangle has for base the arithmetic mean $x_2 = (x_1 + y_1)/2$ and consequently height $y_2 = a/x_2$, the third rectangle has for base the arithmetic mean $x_3 = (x_2 + y_2)/2$, ... and so on. The recursive equation for the basis is

$$x_{n+1} = \frac{1}{2} \left(x_n + \frac{a}{x_n} \right).$$

Observe that if the area a and the initial conjecture x_1 are rationals, then all the x_n are rationals too.

The algorithm converges, and quite fast. Consider, for example, $a = 2$, so that we are looking for $\sqrt{2}$. We could, as the Babylonians, start from an initial guess $x_1 = 3/2$ for $\sqrt{2}$ (since $1^2 < 2 < 2^2$), and find

$$x_2 = \frac{17}{12} \simeq 1.41666666666 \quad x_3 = \frac{577}{408} \simeq 1.41421568627 \quad x_4 = \frac{665857}{470832} \simeq 1.41421356237$$

As you see, the sequence stabilizes quite fast.

As a first attempt to explain this miracle, we could start looking at the recursive equations for the bases and the heights of the rectangles:

$$x_{n+1} = \frac{x_n + y_n}{2} \quad 1/y_{n+1} = \frac{1/x_n + 1/y_n}{2}$$

(so, the next height is the “harmonic mean” of the base and height). We see that the x_n 's and the y_n 's form decreasing and increasing sequences, respectively (disregarding the first guess, of course), namely

$$y_2 \leq y_3 \leq \dots \leq y_n \leq \dots \leq x_n \leq \dots \leq x_3 \leq x_2,$$

The real root is somewhere between, namely $y_n \leq \sqrt{a} \leq x_n$. Hence, we have an explicit control of the error: the difference between x_n (or y_n) and the real value of $\sqrt{a}$ is not greater than $|x_n - y_n|$. A computation shows that the lengths of those intervals, the differences $\varepsilon_n = x_n - y_n$ satisfy the recursion

$$\varepsilon_{n+1} < \frac{1}{2} \cdot \varepsilon_n$$

So, and initial “error” $\varepsilon_1 \leq 1$ (an easy achievement, since we easily recognise squares of integers) reduces to at least $\varepsilon_n \leq 2^{-n}$ after n iterations. The true error is actually much smaller. Indeed, in our example we may compute

$$\varepsilon_2 = \frac{17}{12} - 2\frac{12}{17} = \frac{1}{204} \simeq 0.005 \quad \text{and} \quad \varepsilon_3 = \frac{577}{408} - 2\frac{408}{577} = \frac{1}{235416} \simeq 0.000004$$

So that the first improved guess x_2 has already one correct decimal, and the second, x_3 has already four correct decimals!

⁵Heron of Alexandria, *Metrica*, Book I.

⁶Carl B. Boyer, *A history of mathematics*, John Wiley & Sons, 1968.

⁷O. Neugebauer, *The exact sciences in antiquity*, Dover, 1969.

Irrationals. What Babylonians didn't suspect is that if you start with a rational guess for $\sqrt{2}$, you get an infinite sequence of rational approximations, but the process never stops. This is due to the following great discovery of Greek mathematicians.

Teorema 1.2 (Pythagoras). *There is no rational number whose square is equal to 2.*

In modern language this means that the square root of 2 is not rational. Indeed, assume that such a rational p/q exists, and assume it is reduced. Squaring we get $(p/q)^2 = 2$, that is, $p^2 = 2q^2$. Therefore, p^2 is divisible by 2, hence by 2^2 (because the factorisation of a square must contain even exponents). But this implies the existence of an integer r such that $2^2r = 2q^2$, hence also q^2 is divisible by 2, contrary to our hypothesis that the fraction was reduced.

It is clear that the same proof work with other square roots.

Supremum axiom. Pythagoras theorem suggests the need to enlarge the field $\mathbb{Q}$ of rational numbers and get a larger field. This is done by admitting a new axiom, in addition to the axioms of field and order. This is a rather technical point, but it amounts to saying that the reals “have no holes”, and may be thought as a continuous line of points.

First, we need some terminology. A *upper bound* (*limite superior*) of a set A is any number M such that $a \leq M$ for any $a \in A$. If a upper bound of A belongs to A (and therefore is the unique one belonging to A !), then it is called a *maximum* of A , and denoted $\max A$. Similarly, a *lower bound* (*limite inferior*) of a set A is any number m such that $m \leq a$ for any $a \in A$. A lower bound which belongs to A itself is called *minimum* of A , and denoted $\min A$.

Clearly, a set may have no upper and/or lower bound. A set of numbers A is *bounded from above* if it admits an upper bound, and *bounded from below* if it admits a lower bound. It is called *bounded* if it admits both upper and lower bounds (i.e. if there exists a number K such that $|a| \leq K$ for any $a \in A$). It may also happens that a set is bounded above and/or below without having maximum and/or minimum.

We define the *supremum* of A , notation $\sup A$, as the smallest of all the upper bounds of A . This means that $M = \sup A$ if $a \leq M$ for all $a \in A$, and if no $b < M$ is an upper bound for A . Similarly, we define the *infimum* of A , notation $\inf A$, as the largest of all lower bounds. Both supremum and infimum, if they exist, are clearly unique.

This is the final axiom, to be added to the field and order axioms, which entirely defines the real line:

S1 (*the supremum axiom*) Any not-empty subset $A \subset \mathbb{R}$ of the real line which is bounded from above has a supremum.

Of course, also any not-empty subset $B \subset \mathbb{R}$ which is bounded from below has a infimum (just reverse the signs of the numbers forming the set).

Finally, the *real line* $\mathbb{R}$ is the unique (up to isomorphism!) ordered and complete field, i.e. is characterised by the axioms R1-R5, F6, O1-O3 and S1. For a proof of uniqueness you may consult [Ap69]. Existence of such a field uses constructions by Dedekind (cuts) or Cantor (sequences).

Existence of the square root of two. So, for example, consider the set of decreasing rationals $\dots \leq x_n \leq \dots \leq x_3 \leq x_2$ obtained by the Heron method as basis of rectangles of area equal to 2. Since they all satisfy $x_n^2 > 2$, they admits an infimum, say a , which clearly satisfies $a^2 \geq 2$. Similarly, the heights $y_2 \leq y_3 \leq \dots \leq y_n \leq \dots$ satisfy $y_n^2 < 2$, and therefore their supremum b satisfies $b^2 \leq 2$. But the difference $|x_n - y_n|$ is arbitrarily small, since it is bounded by $1/2^n$. There follows that $a = b$ and therefore $a^2 = 2$.

Archimedean property. An important consequence of the supremum axiom is that the set of natural numbers $\mathbb{N} \subset \mathbb{R}$ is unbounded from above (if it were bounded it would have a supremum $s = \sup \mathbb{N}$, but then there would exist some natural $n > s - 1$, and we could find another natural $n^+ = n + 1 > s$, contradicting the assumption that s is a upper bound for $\mathbb{N}$). There follows that any real $x \in \mathbb{R}$ is strictly less than some natural n (and therefore of all its successors). Now, take any positive real number $\varepsilon > 0$. We claim that for any $x \in \mathbb{R}$ we can find an integer $n \in \mathbb{N}$ so large that

$$n \cdot \varepsilon > x,$$

for otherwise x/ε would be an upper bound for $\mathbb{N}$. This property of numbers, that “multiples of a given positive quantity (no matter how small) may be as large as we want”, is called *Archimedean property*.

Completeness. Even more important, the supremum axiom says that the real line does not lack points, i.e. is “complete” in a very technical sense: any monotone and bounded sequence of real numbers is convergent. This in turns implies the Bolzano theorem on the zeroes of continuous functions, hence most deep results of calculus. See [Ap69] or [Li95].

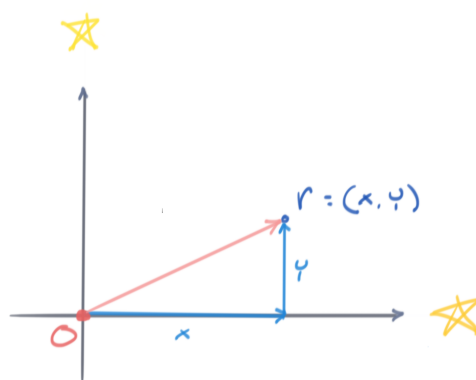
1.2 Vetores

Listas de números reais descrevem/modelam o plano euclidiano, o espaço da física de Galileo, o espaço de fases da mecânica de Newton, ...

22 set 2022

Listas. Sejam A um conjunto e n um número natural. Uma *lista* de comprimento n , ou *n-úpla*, nos elementos de A é uma coleção ordenada $(a_1, a_2, \dots, a_n)$ de elementos $a_k \in A$ (que pode ser pensada, omitindo vírgulas e parêntesis, como uma “palavra” de comprimento n nas “letras” do “alfabeto” A). O elemento a_k é chamado *k-ésima coordenada* da lista $(a_1, a_2, \dots, a_n)$. Duas listas $(a_1, a_2, \dots, a_n)$ e $(b_1, b_2, \dots, b_n)$ são iguais sse $a_k = b_k$ para cada $k = 1, 2, \dots, n$. Ou seja, numa lista a ordem conta (assim, “rima” e “amor” são duas listas diferentes nas letras do alfabeto português) e são possíveis repetições de letras (“sosa” é uma palavra possível). O conjunto de todas as listas de comprimento n nas letras do conjunto A é chamado *produto cartesiano de n-vezes A*, e denotado por A^n .

O plano cartesiano. O *plano cartesiano*⁸ $\mathbb{R}^2 := \mathbb{R} \times \mathbb{R}$ é o conjunto dos *pontos* $\mathbf{r} = (x, y)$, com “coordenadas (cartesianas)” $x, y \in \mathbb{R}$ (chamadas *abscissa* e *ordenada*, respetivamente). A *origem* é o ponto $\mathbf{0} := (0, 0)$. Os eixos correspondem a duas direções, por exemplo determinadas por duas estrelas fixas, X e Y . O ponto de coordenadas (x, y) é o ponto atingido ao deslocar-se x “passos” na direção da estrela X e depois y “passos” na direção da estrela Y , começando pela origem. O ponto $\mathbf{r} = (x, y)$ pode também ser pensado como o *vetor* (o segmento orientado) entre a origem e o ponto $\mathbf{r}$.



É possível definir duas operações naturais no plano cartesiano. A *soma* dos vetores $\mathbf{r} = (x, y)$ e $\mathbf{r}' = (x', y')$ é o vetor

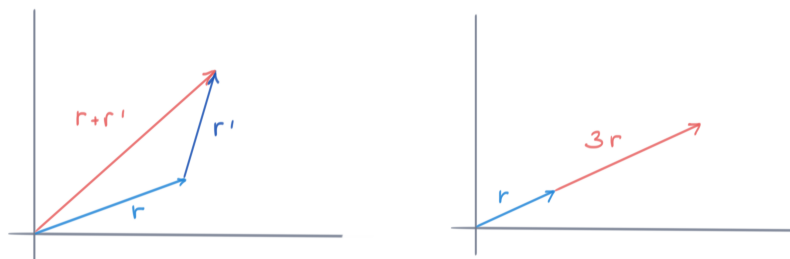
$$\mathbf{r} + \mathbf{r}' := (x + x', y + y'),$$

que representa uma diagonal do paralelogramo de lados $\mathbf{r}$ e $\mathbf{r}'$. Por exemplo, $(1, 2) + (3, 4) = (4, 6)$. O *produto* do número/escalar $\lambda \in \mathbb{R}$ pelo vetor $\mathbf{r} = (x, y)$ é o vetor

$$\lambda \mathbf{r} := (\lambda x, \lambda y)$$

que representa uma dilatação/contração (e uma inversão se $\lambda < 0$) de razão λ do vetor $\mathbf{r}$. Por exemplo, $3(1, 2) = (3, 6)$, e $-\frac{1}{2}(10, 12) = (-5, -6)$.

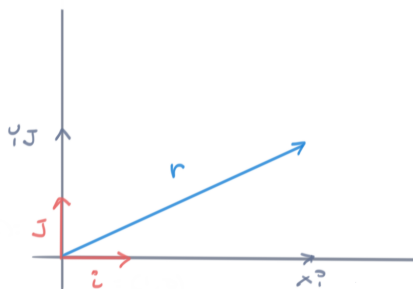
⁸René Descartes, *La Géométrie* [em *Discourse de la Méthode*, 1637].



Cada vetor pode ser representado de maneira única como soma

$$\mathbf{r} = (x, y) = x\mathbf{i} + y\mathbf{j},$$

onde $\mathbf{i} := (1, 0)$ e $\mathbf{j} := (0, 1)$ denotam os vetores da “base canónica” (a seguir daremos uma definição de “base”).



Lugares geométricos (pontos, retas, circunferências, parábolas, ...) podem ser descritos/definidos por equações algébricas, ditas “equações cartesianas”.

ex: Descreva as coordenadas cartesianas dos pontos da reta que passa por $(1, 2)$ e $(-1, 3)$.

ex: Descreva as coordenadas cartesianas do triângulo de vértices $(0, 0)$, $(1, 0)$ e $(0, 2)$.

ex: Esboce os lugares geométricos definidos pelas equações

$$xy = 1 \quad y = 2x - 7 \quad (x + 1)^2 + (y - 3)^2 = 9 \quad x - 2y^2 = 3$$

$$\begin{cases} x + y = 1 \\ x - 1 = 1 \end{cases} \quad \begin{cases} x + y = 3 \\ -2x - 2y = -6 \end{cases} \quad \begin{cases} x + y = 1 \\ 3x + 3y = 1 \end{cases}$$

ex: Esboce os lugares geométricos definidos pelas seguintes desigualdades

$$x - y \leq 1 \quad \begin{cases} 0 \leq x \leq 1 \\ 0 \leq y \leq 1 \end{cases} \quad \begin{cases} x + y \leq 1 \\ x - y \leq 1 \end{cases}$$

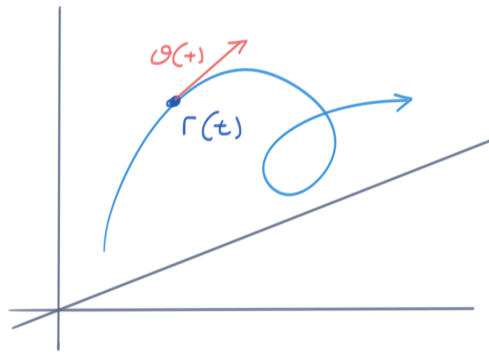
ex: Determine umas desigualdades (cartesianas) que definem os pontos do paralelogramo de lados $(2, 1)$ e $(3, 5)$.

O espaço, o espaço-tempo e o espaço de fases da física newtoniana. O espaço onde acontece a física newtoniana é o *espaço 3-dimensional* $\mathbb{R}^3 := \mathbb{R} \times \mathbb{R} \times \mathbb{R}$. A posição de uma partícula é um ponto/vetor

$$\mathbf{r} = (x, y, z) := x\mathbf{i} + y\mathbf{j} + z\mathbf{k} \in \mathbb{R}^3$$

onde $\mathbf{i} := (1, 0, 0)$, $\mathbf{j} := (0, 1, 0)$ e $\mathbf{k} := (0, 0, 1)$ denotam os vetores da base canônica. Ou seja, uma receita para deslocar um ponto material da *origem* $\mathbf{0} := (0, 0, 0)$ até a posição $\mathbf{r} = (x, y, z)$ consiste em fazer x “passos” na direção do vetor $\mathbf{i}$, depois y “passos” na direção do vetor $\mathbf{j}$ e enfim z “passos” na direção do vetor $\mathbf{k}$ (mas a ordem é indiferente!).

A *lei horária/trajetória*, de uma partícula é uma função $t \mapsto \mathbf{r}(t)$ que associa a cada tempo t num intervalo $I \subset \mathbb{R}$ a posição $\mathbf{r}(t) = (x(t), y(t), z(t)) \in \mathbb{R}^3$ da partícula no instante t . A *velocidade* da partícula no instante t é o vetor $\mathbf{v}(t) := \dot{\mathbf{r}}(t) = (\dot{x}(t), \dot{y}(t), \dot{z}(t))$ formado pelas derivadas das coordenadas.



A *aceleração* da partícula no instante t é o vetor $\mathbf{a}(t) := \dot{\mathbf{v}}(t) = \ddot{\mathbf{r}}(t) = (\ddot{x}(t), \ddot{y}(t), \ddot{z}(t))$ formado pelas segundas derivadas das coordenadas. A aceleração de uma partícula de massa $m > 0$ num referencial inercial é determinada pela equação (segunda lei) de Newton⁹

$$m \mathbf{a}(t) = \mathbf{F}(\mathbf{r}(t), \dot{\mathbf{r}}(t))$$

onde $\mathbf{F} : \mathbb{R}^3 \times \mathbb{R}^3 \rightarrow \mathbb{R}^3$ é um campo de forças.

Por exemplo, a lei horária do movimento retilíneo uniforme é $\mathbf{r}(t) = \mathbf{a} + \mathbf{v}t$ onde $\mathbf{a}$ é a posição inicial e $\mathbf{v}$ a velocidade, um vetor constante.

O *espaço-tempo* da física newtoniana é o produto cartesiano $\mathbb{R} \times \mathbb{R}^3 \approx \mathbb{R}^4$, o espaço dos *eventos* $(t, x, y, z) \in \mathbb{R}^4$, onde $\mathbf{r} = (x, y, z) \in \mathbb{R}^3$ representa uma posição num referencial inercial, e $t \in \mathbb{R}$ é o *tempo absoluto*.

O *estado* de uma partícula, a informação necessária e suficiente para resolver a equação de Newton e portanto determinar a sua trajetória futura (e passada), é um ponto $(\mathbf{r}, \mathbf{p}) \in \mathbb{R}^3 \times \mathbb{R}^3 \approx \mathbb{R}^6$ do *espaço dos estados/de fases*, onde $\mathbf{r}$ é a posição e $\mathbf{p} := m\mathbf{v}$ é o *momento (linear)*.

ex: A lei horária da queda livre, próximo da superfície da terra, é $\mathbf{r}(t) = \mathbf{a} + \mathbf{v}t + (0, 0, -g/2)t^2$, onde $\mathbf{a}$ é a posição inicial, $\mathbf{v}$ a velocidade inicial e $g \simeq 9.8 \text{ m/s}^2$ a aceleração gravitacional. Calcule a velocidade $\dot{\mathbf{r}}(t)$ e a aceleração $\ddot{\mathbf{r}}(t)$.

Vetores aplicados. Um *vetor aplicado/geométrico* (uma força, uma velocidade, ...) é um segmento orientado $\vec{AB}$ entre um ponto de aplicação $A \in \mathbb{R}^3$ e um ponto final $B \in \mathbb{R}^3$. Dois vetores aplicados $\vec{AB}$ e $\vec{CD}$ são *paralelos* se $B - A = \lambda(D - C)$ com $\lambda \in \mathbb{R}$, e são *equivalentes* (e portanto definem o mesmo “vetor” $\mathbf{x} = B - A$) se $B - A = D - C$.

ex: Mostre que cada vetor aplicado é equivalente a um vetor aplicado na origem.

ex: Diga se são paralelos ou equivalentes $\vec{AB}$ e $\vec{CD}$ quando

$$\begin{array}{llll} A = (1, 2) & B = (-1, 1) & C = (2, 3) & D = (4, , 4) \\ A = (0, 1, \pi) & B = (-2, 3, 0) & C = (1, 0, -\pi) & D = (2, 3, 0) \end{array}$$

⁹Isaac Newton, *Philosophiæ Naturalis Principia Mathematica*, 1687.

ex: Determine $D \in \mathbb{R}^n$ de maneira tal que $\vec{AB}$ e $\vec{CD}$ sejam equivalentes quando

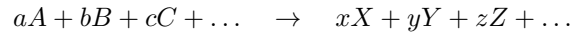
$$A = (1, 2) \quad B = (-1, 1) \quad C = (2, 3)$$

$$A = (0, 1, \pi) \quad B = (-2, 3, 0) \quad C = (0, 0, 0)$$

Composição de forças. Se duas forças $\mathbf{F}$ e $\mathbf{G}$ atuam sobre uma partícula colocada num certo ponto do espaço, então a “resultante” é uma força $\mathbf{F} + \mathbf{G}$.

ex: Determine a “dimensão” do espaço de fases de um sistema composto por 8 planetas (como, por exemplo, Mercúrio, Vênus, Terra, Marte, Júpiter, Saturno, Urano, Netuno) e de um sistema composto por 6×10^{23} moléculas.

Reações químicas. O estado de uma reação química



entre os n reagentes $A, B, C, \dots$ e os m produtos $X, Y, Z, \dots$ é descrito usando as concentrações $[A], [B], [C], \dots, [X], [Y], [Z], \dots$, e portanto $n + m$ números.

1.3 O espaço $\mathbb{R}^n$

Os exemplos físicos na seção anterior motivam a definição abstrata de espaço vetorial real de dimensão arbitrária n .

O espaço vetorial $\mathbb{R}^n$. O *espaço vetorial real de dimensão n* é o conjunto

$$\mathbb{R}^n := \underbrace{\mathbb{R} \times \mathbb{R} \times \dots \times \mathbb{R}}_{n \text{ vezes}}$$

das n -uplas $\mathbf{x} = (x_1, x_2, \dots, x_n)$ de números reais, ditas *vetores* ou *pontos*, munido das operações *adição/soma*, que envia dois vetores $\mathbf{x}$ e $\mathbf{y}$ no vetor

$$\mathbf{x} + \mathbf{y} := (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n)$$

e *multiplicação/produto por um escalar*, que envia um vetor $\mathbf{x}$ e um *escalar* $\lambda \in \mathbb{R}$ no vetor

$$\lambda \mathbf{x} := (\lambda x_1, \lambda x_2, \dots, \lambda x_n)$$

O vetor $\lambda \mathbf{x}$ é dito *proporcional* ao vetor $\mathbf{x}$. O *vetor nulo/origem* é o vetor

$$\mathbf{0} := (0, 0, \dots, 0)$$

e satisfaz

$$\mathbf{x} + \mathbf{0} = \mathbf{x}$$

para todo $\mathbf{x} \in \mathbb{R}^n$. O *simétrico* do vetor $\mathbf{x} = (x_1, x_2, \dots, x_n)$ é o vetor

$$-\mathbf{x} := (-1)\mathbf{x} = (-x_1, -x_2, \dots, -x_n)$$

que satisfaz

$$\mathbf{x} + (-\mathbf{x}) = \mathbf{0}$$

Isto justifica a notação $\mathbf{x} - \mathbf{y} := \mathbf{x} + (-\mathbf{y})$. A soma de vetores é comutativa, ou seja,

$$\mathbf{x} + \mathbf{y} = \mathbf{y} + \mathbf{x}$$

e associativa, ou seja,

$$\mathbf{x} + (\mathbf{y} + \mathbf{z}) = (\mathbf{x} + \mathbf{y}) + \mathbf{z}$$

(e portanto as parêntesis são inúteis). Em particular, a ordem dos vetores numa soma (finita) $\mathbf{x} + \mathbf{y} + \cdots + \mathbf{z}$ é irrelevante. É também evidente que

$$\lambda(\mu\mathbf{x}) = (\lambda\mu)\mathbf{x}$$

e que valem as propriedades distributivas

$$(\lambda + \mu)\mathbf{x} = \lambda\mathbf{x} + \mu\mathbf{x} \quad \text{e} \quad \lambda(\mathbf{x} + \mathbf{y}) = \lambda\mathbf{x} + \lambda\mathbf{y}$$

Finalmente, a multiplicação pelo escalar 1 transforma um vetor em si próprio, ou seja, $1\mathbf{x} = \mathbf{x}$, e a multiplicação pelo escalar “zero” produz o vetor nulo, ou seja, $0\mathbf{x} = \mathbf{0}$ (assim que, com abuso de linguagem, o escalar 0 pode também denotar o vetor nulo $\mathbf{0}$, desde que seja claro no contexto).

A *combinação linear* (ou *sobreposição*) dos vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k \in \mathbb{R}^n$ com *coeficientes* $\lambda_1, \lambda_2, \dots, \lambda_k \in \mathbb{R}$ é o vetor

$$\sum_{i=1}^k \lambda_i \mathbf{v}_i := \lambda_1 \mathbf{v}_1 + \lambda_2 \mathbf{v}_2 + \cdots + \lambda_k \mathbf{v}_k.$$

A *base canónica* de $\mathbb{R}^n$ é o conjunto ordenado dos vetores

$$\mathbf{e}_1 = (1, 0, \dots, 0) \quad \mathbf{e}_2 = (0, 1, 0, \dots, 0) \quad \dots \quad \mathbf{e}_n = (0, \dots, 0, 1)$$

assim que cada vetor $\mathbf{x} = (x_1, x_2, \dots, x_n) \in \mathbb{R}^n$ é uma combinação linear única

$$\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \cdots + x_n \mathbf{e}_n$$

dos vetores da base canónica. O número x_k é chamado *k-ésima coordenada*, ou *componente*, do vetor $\mathbf{x}$ (relativamente à base canónica).

No *plano* $\mathbb{R}^2$ os pontos costumam ser denotados por $\mathbf{r} = (x, y)$, e no *espaço* (3-dimensional) $\mathbb{R}^3$ por $\mathbf{r} = (x, y, z)$.

ex: Calcule

$$(1, 2, 3) + (2, 3, 4) \qquad 6 \cdot (-1, -6, 0) \qquad (1, -1) - (3, 2)$$

ex: Calcule e esboce os pontos $A + B$, $A - B$, $2A - 3B$ e $-A + \frac{1}{2}B$ quando

$$A = (1, 2) \text{ e } B = (-1, 1) \quad \text{ou} \quad A = (0, 1, 7) \text{ e } B = (-2, 3, 0)$$

ex: [Ap69] 12.4.

1.4 Translações e homotetias

Translações e homotetias. A soma de vetores e a multiplicação por um escalar, as duas operações algébricas definidas no espaço vetorial $\mathbb{R}^n$, descrevem transformações geométricas que são translações e homotetias.

Uma *translação* do espaço $\mathbb{R}^n$ é uma transformação $T_{\mathbf{a}} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ definida por

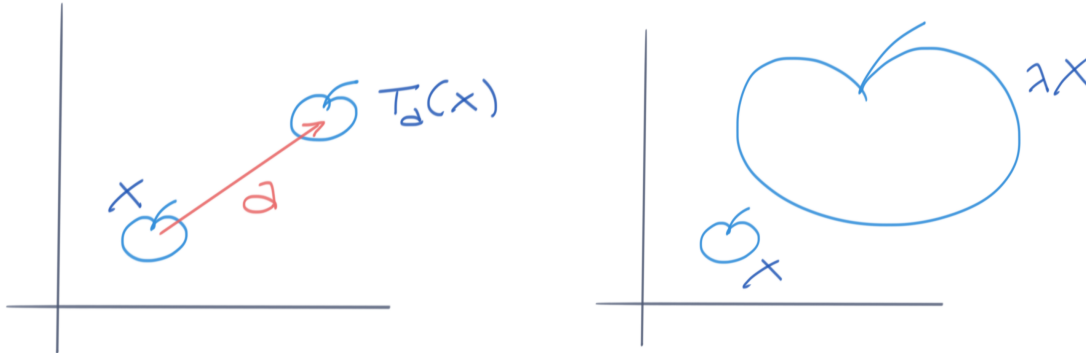
$$\mathbf{x} \mapsto T_{\mathbf{a}}(\mathbf{x}) := \mathbf{x} + \mathbf{a},$$

com $\mathbf{a} \in \mathbb{R}^n$. A imagem de um subconjunto $X \subset \mathbb{R}^n$ é também denotada $\mathbf{a} + X := T_{\mathbf{a}}(X)$.

A *homotetia* de razão $\lambda \neq 0$ é a transformação $M_{\lambda} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ definida por

$$\mathbf{x} \mapsto M_{\lambda}(\mathbf{x}) := \lambda\mathbf{x}.$$

Representa uma dilatação quando $\lambda > 1$, e uma contração quando $\lambda < 1$. A imagem de um subconjunto $X \subset \mathbb{R}^n$ é também denotada $\lambda X := M_{\lambda}(X)$.



A composição de uma homotetia e uma translação é a transformação

$$\mathbf{x} \mapsto T_{\mathbf{a}} \circ M_{\lambda}(\mathbf{x}) = \lambda \mathbf{x} + \mathbf{a}.$$

Em particular, é possível definir uma *homotetia* de centro $\mathbf{c} \in \mathbb{R}^n$ e razão $\lambda \neq 0$ como sendo a transformação $H_{\mathbf{c},\lambda} : \mathbb{R}^n \rightarrow \mathbb{R}^n$

$$\mathbf{x} \mapsto H_{\mathbf{c},\lambda}(\mathbf{x}) := \mathbf{c} + \lambda(\mathbf{x} - \mathbf{c}).$$

Observe que o centro é um ponto fixo de uma homotetia, i.e. $H_{\mathbf{c},\lambda}(\mathbf{c}) = \mathbf{c}$.

ex: Mostre que $T_{\mathbf{a}} \circ T_{\mathbf{b}} = T_{\mathbf{a}+\mathbf{b}}$ e deduza que $T_{\mathbf{a}} \circ T_{-\mathbf{a}} = T_{-\mathbf{a}} \circ T_{\mathbf{a}} = \mathbf{1}$.

ex: Mostre que $\forall \mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ existe uma translação $T_{\mathbf{a}}$ tal que $T_{\mathbf{a}}(\mathbf{x}) = \mathbf{y}$.

ex: Descreva o efeito das transformações M_{λ} , com $\lambda < 1$ e $\lambda > 1$, no plano $\mathbb{R}^2$.

ex: Determine e compare as transformações compostas $T_{\mathbf{a}} \circ M_{\lambda}$ e $M_{\lambda} \circ T_{\mathbf{a}}$. São iguais?

Graus Celsius, Fahrenheit e Kelvin. A temperatura pode ser medida em graus Celsius (C), Fahrenheit (F) e Kelvin (K), e

$$F = 1.8 \cdot C + 32 \quad K = (F + 459.67)/1.8$$

ex: Determine a relação entre graus Kelvin e Celsius.

ex: Determine a relação entre um grau Kelvin e um grau Fahrenheit.

Invariância galileiana/sistemas inerciais. Seja $\mathbf{r} \in \mathbb{R}^3$ a posição de uma partícula num referencial inercial $\mathcal{R}$. Num referencial $\mathcal{R}'$ em movimento retilíneo uniforme com velocidade (constante) $\mathbf{V} \in \mathbb{R}^3$ e origem $\mathbf{R}$ no instante $t = 0$, a posição da partícula é dada pela “transformação de Galileio” [LL78]

$$\mathbf{r}' = \mathbf{r} - (\mathbf{R} + \mathbf{V}t). \quad (1.3)$$

ex: Verifique que a diferença $\mathbf{r}_1 - \mathbf{r}_2$ entre as posições de duas partículas é invariante para as transformações (1.3), ou seja, $\mathbf{r}'_1 - \mathbf{r}'_2 = \mathbf{r}_1 - \mathbf{r}_2$. Verifique que a aceleração $\mathbf{a} := \ddot{\mathbf{r}}$ da trajetória $t \mapsto \mathbf{r}(t)$ de uma partícula também é invariante para as transformações (1.3), ou seja, a aceleração não depende do sistema inercial no qual é calculada. Deduza que se a força entre duas partículas apenas depende da diferença entre as posições (interação gravitacional, interação elétrica, ...), então a lei de Newton $m\ddot{\mathbf{r}} = \mathbf{F}$ é invariante.

ex: Mostre que o momento linear $\mathbf{p} := m\dot{\mathbf{r}}$ transforma segundo a lei

$$\mathbf{p}' = \mathbf{p} + m\mathbf{V}.$$

Centro de massas. O *centro de massas* do sistema de partículas de massas $m_1, m_2, \dots, m_N$ colocadas nos pontos $\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N \in \mathbb{R}^3$ é

$$\mathbf{R} := \frac{1}{M} \sum_{k=1}^N m_k \mathbf{r}_k = \frac{m_1 \mathbf{r}_1 + m_2 \mathbf{r}_2 + \dots + m_N \mathbf{r}_N}{M}$$

onde $M := m_1 + m_2 + \dots + m_N$ é a massa total do sistema.

ex: Três massas unitárias são colocadas nos vértices de um triângulo no plano. Caracterize o centro de massa.

2 Produto escalar, norma e distância

ref: [Ap69] Vol 1, 12.5-11 ; [La97] Ch. I, 3-4

2.1 Produto escalar euclidiano

29 set 2022

A geometria euclidiana do plano, e em geral dos espaços $\mathbb{R}^n$, pode ser reconstruída a partir da noção algébrica de “produto escalar”.

Módulo e distância na reta real. O *módulo*, ou *valor absoluto*, do número real $x \in \mathbb{R}$ é

$$|x| := \max\{x, -x\} = \begin{cases} x & \text{se } x \geq 0 \\ -x & \text{se } x < 0 \end{cases}.$$

A *distância* entre os pontos x e y da reta real $\mathbb{R}$ é

$$d(x, y) := |x - y|$$

O módulo e a distância satisfazem as desigualdades do triângulo

$$|x + y| \leq |x| + |y| \quad \text{e} \quad d(x, y) \leq d(x, z) + d(y, z).$$

ex: Mostre que as desigualdades podem ser estritas.

O plano euclidiano. De acordo com o teorema de Pitágoras, o *comprimento* do vetor $\mathbf{r} = (x, y) \in \mathbb{R}^2$, ou seja, a distância entre o ponto representado pelo vetor $\mathbf{r}$ e a origem, é dado pela expressão

$$d(\mathbf{r}, \mathbf{0}) = \sqrt{x^2 + y^2}.$$

Mais em geral, pela homogeneidade do plano euclidiano, somos levados a definir a *distância* entre os pontos $\mathbf{r} = (x, y)$ e $\mathbf{r}' = (x', y')$ como sendo o comprimento do vetor $\mathbf{r}' - \mathbf{r}$, ou seja,

$$d(\mathbf{r}', \mathbf{r}) = \sqrt{(x - x')^2 + (y - y')^2}.$$

Acontece que toda a geometria euclidiana do plano (distâncias, ângulos, paralelismo e perpendicularidade, ...) pode ser deduzida a partir da noção algébrica de *produto escalar/interno*

$$\mathbf{r} \cdot \mathbf{r}' := xx' + yy'.$$

Por exemplo, os vetores $\mathbf{r} = (x, y)$ e $\mathbf{r}' = (x', y')$ são *perpendiculares/ortogonais* quando $\mathbf{r} \cdot \mathbf{r}' = 0$, ou seja, quando $xx' + yy' = 0$. O comprimento do vetor $\mathbf{r}$ fica então igual a $\sqrt{\mathbf{r} \cdot \mathbf{r}}$.

ex: Determine um vetor ortogonal ao vetor $(2, 3)$.

ex: Calcule o comprimento do vetor $(3, 4)$.

ex: Calcule a distância entre os pontos $(2, 1)$ e $(1, 2)$.

Produto escalar euclidiano. O *produto escalar/interno euclidiano* entre os vetores $\mathbf{x}$ e $\mathbf{y}$ de $\mathbb{R}^n$ é o escalar

$$\boxed{\mathbf{x} \cdot \mathbf{y} := x_1y_1 + x_2y_2 + \cdots + x_ny_n} \quad (2.1)$$

Usando o símbolo de somatório, o produto escalar fica $\mathbf{x} \cdot \mathbf{y} := \sum_{i=1}^n x_iy_i$. Também conveniente é utilizar o *símbolo de Kroneker* δ_{ij} , definido por $\delta_{ii} = 1$ e $\delta_{ij} = 0$ se $i \neq j$, assim que o produto escalar é $\mathbf{x} \cdot \mathbf{y} = \sum_{i,j} \delta_{ij}x_iy_j$. Outra notação tradicional para o produto escalar é $\langle \mathbf{x}, \mathbf{y} \rangle$.

O produto escalar euclidiano satisfaz as seguintes propriedades fundamentais, embora elementares, das quais podem ser deduzidas a quase totalidade das consequências interessantes deste capítulo. O produto escalar é *simétrico*, ou seja,

$$\boxed{\mathbf{x} \cdot \mathbf{y} = \mathbf{y} \cdot \mathbf{x}} \quad (2.2)$$

para todos os vetores $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$. O produto escalar é *linear* na primeira variável, ou seja é homogêneo e aditivo,

$$\boxed{(\lambda \mathbf{x}) \cdot \mathbf{y} = \lambda(\mathbf{x} \cdot \mathbf{y}) \quad \text{e} \quad (\mathbf{x} + \mathbf{y}) \cdot \mathbf{z} = \mathbf{x} \cdot \mathbf{z} + \mathbf{y} \cdot \mathbf{z}} \quad (2.3)$$

assim que

$$(\lambda \mathbf{x} + \mu \mathbf{y}) \cdot \mathbf{z} = \lambda(\mathbf{x} \cdot \mathbf{z}) + \mu(\mathbf{y} \cdot \mathbf{z})$$

para todos os vetores $\mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbb{R}^n$ e os escalares $\lambda, \mu \in \mathbb{R}$. Pela simetria (2.2), o mesmo acontece quando consideramos combinações lineares na segunda variável, ou seja, o produto escalar é também linear na segunda variável (os matemáticos dizem então que o produto escalar euclidiano é *bilinear*). Finalmente, um cálculo mostra que o produto escalar de um vetor $\mathbf{x}$ com si próprio é uma soma de quadrados das coordenadas, pois $\mathbf{x} \cdot \mathbf{x} = x_1^2 + x_2^2 + \cdots + x_n^2$. Consequentemente, o produto escalar é *positivo*, ou seja,

$$\boxed{\mathbf{x} \cdot \mathbf{x} \geq 0 \quad \text{e} \quad \mathbf{x} \cdot \mathbf{x} = 0 \text{ sse } \mathbf{x} = \mathbf{0}} \quad (2.4)$$

A positividade (2.4) claramente implica que o único vetor $\mathbf{x}$ tal que $\mathbf{x} \cdot \mathbf{y} = 0$ para todos os $\mathbf{y} \in \mathbb{R}^n$ é o vetor nulo $\mathbf{0}$.

Os vetores $\mathbf{x}$ e $\mathbf{y}$ são ditos *ortogonais/perpendiculares* quando $\mathbf{x} \cdot \mathbf{y} = 0$. Uma notação conveniente pode ser $\mathbf{x} \perp \mathbf{y}$.

ex: Mostre que o produto interno é simétrico, bilinear e positivo.

ex: Verifique que os vetores da base canônica são ortogonais dois a dois, ou seja, $\mathbf{e}_i \cdot \mathbf{e}_j = 0$ se $i \neq j$.

ex: Se $\mathbf{v}$ é ortogonal a todos os vetores $\mathbf{x} \in \mathbb{R}^n$, então $\mathbf{v} = \mathbf{0}$.

ex: Calcule o produto interno entre $\mathbf{x} = (1, 2)$ e $\mathbf{y} = (-1, 1)$, e entre $\mathbf{x} = (0, 1, \pi)$ e $\mathbf{y} = (-2, 3, 0)$.

ex: Determine se são ortogonais $\mathbf{x} = (1, 2)$ e $\mathbf{y} = (-1, 1)$, ou $\mathbf{x} = (0, 1, \pi)$ e $\mathbf{y} = (-2, 3, 0)$.

ex: Se $\mathbf{x} \cdot \mathbf{y} = \mathbf{x} \cdot \mathbf{z}$ então $\mathbf{y} = \mathbf{z}$?

ex: [Ap69] 12.8.

2.2 Norma euclidiana

Norma euclidiana. Pela positividade (2.4), o produto escalar de um vetor $\mathbf{x} \in \mathbb{R}^n$ com si próprio é um número não negativo, $\mathbf{x} \cdot \mathbf{x} \geq 0$. A *norma euclidiana* do vetor $\mathbf{x}$ é a raiz quadrada não negativa deste número, ou seja,

$$\boxed{\|\mathbf{x}\| := \sqrt{\mathbf{x} \cdot \mathbf{x}}}$$

Naturalmente, é mais fácil “calcular” o quadrado da norma, $\|\mathbf{x}\|^2 = \mathbf{x} \cdot \mathbf{x}$, que é apenas uma soma de produtos, do que a própria norma, que envolve uma raiz quadrada. Como já observado, no plano euclidiano a norma representa o comprimento do vetor, de acordo com o teorema de Pitágoras.

É claro, pela (2.4), que a norma é *positiva*, ou seja

$$\boxed{\|\mathbf{x}\| \geq 0 \quad \text{e} \quad \|\mathbf{x}\| = 0 \text{ sse } \mathbf{x} = \mathbf{0}} \quad (2.5)$$

ou seja, o único vetor com norma 0 é o vetor nulo $\mathbf{0}$. O cálculo $(\lambda \mathbf{x}) \cdot (\lambda \mathbf{x}) = \lambda^2 (\mathbf{x} \cdot \mathbf{x})$ mostra que a norma é *positivamente homogênea*, ou seja, satisfaz

$$\|\lambda \mathbf{x}\| = |\lambda| \|\mathbf{x}\| \quad (2.6)$$

Um cálculo elementar, que usa apenas as propriedades do produto escalar, mostra que

$$\|\mathbf{x} \pm \mathbf{y}\|^2 = \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 \pm 2\mathbf{x} \cdot \mathbf{y} \quad (2.7)$$

Uma primeira consequência é o *teorema de Pitágoras*: se $\mathbf{x}$ e $\mathbf{y}$ são ortogonais então

$$\|\mathbf{x} \pm \mathbf{y}\|^2 = \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2.$$

Ao somar as duas equações (2.7) (ou seja, a expressão com $+$ e a expressão com $-$), deduzimos que a norma euclidiana satisfaz a *identidade do paralelogramo*

$$\|\mathbf{x} + \mathbf{y}\|^2 + \|\mathbf{x} - \mathbf{y}\|^2 = 2\|\mathbf{x}\|^2 + 2\|\mathbf{y}\|^2 \quad (2.8)$$

Por outro lado, ao calcular a diferença das duas equações (2.7), obtemos a *identidade de polarização*

$$4\mathbf{x} \cdot \mathbf{y} = \|\mathbf{x} + \mathbf{y}\|^2 - \|\mathbf{x} - \mathbf{y}\|^2. \quad (2.9)$$

Esta identidade é particularmente importante, pois mostra que o produto escalar pode ser reconstruído a partir da norma euclidiana que define.

Um vetor é dito *unitário* se a sua norma é igual a um. Todo vetor não nulo $\mathbf{x}$ é proporcional a um vetor unitário, por exemplo $\mathbf{u} = \mathbf{x}/\|\mathbf{x}\|$. Este processo é chamado “normalização”. O vetor unitário $\mathbf{u}$ proporcional a um vetor não nulo $\mathbf{x}$ não é único, pois sempre podemos multiplicar $\mathbf{u}$ por ± 1 . O conjunto dos vetores unitários do espaço euclidiano $\mathbb{R}^n$ é chamado *esfera unitária* de dimensão $n - 1$, e denotado por $\mathbb{S}^{n-1}$.

ex: Mostre que a norma é positivamente homogênea e positiva.

ex: Verifique que os vetores $\mathbf{e}_1 = (1, 0, \dots)$, $\mathbf{e}_2 = (0, 1, 0, \dots)$, $\dots$ da base canônica são unitários, ou seja $\|\mathbf{e}_i\| = 1$.

ex: Verifique que se $\mathbf{v} \neq \mathbf{0}$ então $\mathbf{u} = \pm \mathbf{v}/\|\mathbf{v}\|$ é um vetor unitário paralelo a $\mathbf{v}$.

ex: Mostre que $\|\mathbf{x} + \mathbf{y}\| = \|\mathbf{x} - \mathbf{y}\|$ sse $\mathbf{x} \cdot \mathbf{y} = 0$.

ex: Verifique e interprete geometricamente a identidade do paralelogramo (2.8).

ex: Verifique que o produto escalar euclidiano pode ser deduzido da norma usando a identidade de polarização

$$\mathbf{x} \cdot \mathbf{y} = \frac{1}{4} (\|\mathbf{x} + \mathbf{y}\|^2 - \|\mathbf{x} - \mathbf{y}\|^2),$$

ou também

$$\mathbf{x} \cdot \mathbf{y} = \frac{1}{2} (\|\mathbf{x} + \mathbf{y}\|^2 - \|\mathbf{x}\|^2 - \|\mathbf{y}\|^2).$$

ex: Calcule a norma dos vetores $\mathbf{x} = (1, -1)$, $\mathbf{y} = (-1, 1, -1)$ e $\mathbf{z} = (1, 2, 3, 4)$.

ex: Sejam $\mathbf{x}$ e $\mathbf{y}$ dois vetores não paralelos (em particular, não nulos), e $\lambda > 0$ um escalar positivo. Então

$$\frac{\|\lambda \mathbf{x}\|}{\|\mathbf{x}\|} = \frac{\|\lambda \mathbf{y}\|}{\|\mathbf{y}\|} = \frac{\|\lambda(\mathbf{x} + \mathbf{y})\|}{\|\mathbf{x} + \mathbf{y}\|} = \lambda.$$

Fixado um ponto, por exemplo a origem O , desenhe os triângulos OAB e $OA'B'$, onde $A = O + \mathbf{x}$, $A' = O + \lambda \mathbf{x}$, $B = O + (\mathbf{x} + \mathbf{y})$ e $B' = O + \lambda(\mathbf{x} + \mathbf{y})$, e deduza/reconheça o *teorema de Tales* (ou *teorema da interseção*).

Normas e métricas não euclidianas. Existem outras noções naturais de norma, e relativa distância, que não são definidas a custa de um produto escalar. Por exemplo,

$$\|\mathbf{x}\|_\infty := \max_{1 \leq i \leq n} |x_i| \quad \text{e} \quad \|\mathbf{x}\|_1 := \sum_{i=1}^n |x_i|$$

É imediato verificar que $\|\cdot\|_\infty$ e $\|\cdot\|_1$ são positivas, positivamente homogêneas e subaditivas. Consequentemente, definem distâncias, $d_\infty(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|_\infty$ e $d_1(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|_1$, que são positivas, simétricas e satisfazem a desigualdade do triângulo.

2.3 Geometria euclidiana elementar

Projeções. Os vetores $\mathbf{x}$ e $\mathbf{y}$ de $\mathbb{R}^n$ são ditos *ortogonais/perpendiculares* quando $\mathbf{x} \cdot \mathbf{y} = 0$. Uma notação é $\mathbf{x} \perp \mathbf{y}$. Esta relação é claramente simétrica (mas não é nem reflexiva nem transitiva).

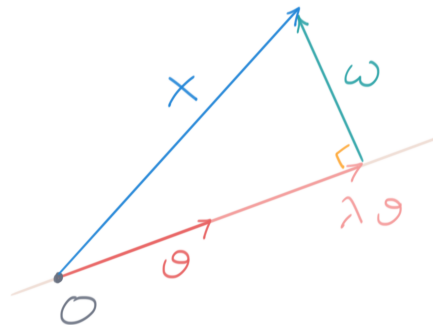
Como já observado, a positividade do produto escalar (2.4) implica que o único vetor $\mathbf{x}$ ortogonal a todos os vetores do espaço euclidiano, ou seja, tal que $\mathbf{x} \cdot \mathbf{y} = 0$ para todos os $\mathbf{y} \in \mathbb{R}^n$, é o vetor nulo $\mathbf{x} = \mathbf{0}$.

Seja $\mathbf{v} \in \mathbb{R}^n$ um vetor não nulo. Todo vetor $\mathbf{x} \in \mathbb{R}^n$ pode ser representado de maneira única como soma

$$\mathbf{x} = \lambda \mathbf{v} + \mathbf{w}$$

de um vetor $\lambda \mathbf{v}$ proporcional a $\mathbf{v}$ e um vetor $\mathbf{w}$ ortogonal a $\mathbf{v}$. De fato, a condição de ortogonalidade $(\mathbf{x} - \lambda \mathbf{v}) \cdot \mathbf{v} = 0$ obriga a escolher

$$\lambda = \frac{\mathbf{x} \cdot \mathbf{v}}{\|\mathbf{v}\|^2} \quad (2.10)$$



O vetor $\lambda \mathbf{v}$ é dito *projeção (ortogonal)* do vetor $\mathbf{x}$ sobre (a reta definida pelo) vetor $\mathbf{v}$, e o coeficiente λ é dito *componente* de $\mathbf{x}$ ao longo de $\mathbf{v}$. Em particular, a componente de $\mathbf{x}$ ao longo de um vetor unitário $\mathbf{u}$ é o produto escalar $\mathbf{x} \cdot \mathbf{u}$. Por exemplo, a componente do vetor $\mathbf{x}$ sobre o vetor $\mathbf{e}_i$ da base canônica é a coordenada $x_i = \mathbf{x} \cdot \mathbf{e}_i$.

Desigualdade de Schwarz e ângulos. Se $\mathbf{x}$ não é proporcional ao vetor não nulo $\mathbf{y}$ (e, em particular, não é nulo), então a norma da diferença $\mathbf{x} - \lambda \mathbf{y}$ entre o vetor $\mathbf{x}$ e a sua projeção $\lambda \mathbf{y}$ sobre $\mathbf{y}$ é estritamente positiva. Um cálculo mostra que

$$0 < \|\mathbf{x} - \lambda \mathbf{y}\|^2 = \|\mathbf{x}\|^2 - \frac{|\mathbf{x} \cdot \mathbf{y}|^2}{\|\mathbf{y}\|^2}$$

Consequentemente,

Teorema 2.1 (desigualdade de Schwarz). *O módulo do produto escalar entre dois vetores $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ é limitado por*

$$|\mathbf{x} \cdot \mathbf{y}| \leq \|\mathbf{x}\| \|\mathbf{y}\| \quad (2.11)$$

e a igualdade verifica-se sse os vetores $\mathbf{x}$ e $\mathbf{y}$ são proporcionais.

A desigualdade de Schwarz (também atribuída ao francês Cauchy e ao ucraniano Bunyakovski) implica que, se $\mathbf{x}$ e $\mathbf{y}$ são vetores não nulos, então

$$-1 \leq \frac{\mathbf{x} \cdot \mathbf{y}}{\|\mathbf{x}\| \|\mathbf{y}\|} \leq 1.$$

Isto permite definir o *ângulo* (ou melhor, o coseno do ângulo) entre os vetores não nulos $\mathbf{x}$ e $\mathbf{y}$ de $\mathbb{R}^n$ como sendo o único $\theta \in [0, \pi)$ tal que

$$\cos \theta = \frac{\mathbf{x} \cdot \mathbf{y}}{\|\mathbf{x}\| \|\mathbf{y}\|} \quad (2.12)$$

Esta definição é compatível com a noção de ortogonalidade: dois vetores não nulos são ortogonais quando $\cos \theta = 0$, logo quando $\theta = \pi/2$.

Métrica euclidiana. Consequência importante da desigualdade de Schwarz é que a norma euclidiana é *subaditiva*, ou seja,

$$\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\| \quad (2.13)$$

para todos $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$. De fato,

$$\begin{aligned} \|\mathbf{x} + \mathbf{y}\|^2 &= \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 + 2\mathbf{x} \cdot \mathbf{y} \\ &\leq \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 + 2\|\mathbf{x}\| \|\mathbf{y}\| \quad \text{pela (2.11)} \\ &= (\|\mathbf{x}\| + \|\mathbf{y}\|)^2 \end{aligned}$$

e a desigualdade (2.13) segue ao calcular a raiz quadrada.

A *distância/métrica euclidiana* entre os vetores/pontos $\mathbf{x}$ e $\mathbf{y}$ de $\mathbb{R}^n$ é definida por

$$d(\mathbf{x}, \mathbf{y}) := \|\mathbf{x} - \mathbf{y}\| \quad (2.14)$$

É imediato verificar que d é uma “distância”, ou seja, que é simétrica,

$$d(\mathbf{x}, \mathbf{y}) = d(\mathbf{y}, \mathbf{x})$$

não negativa e nula apenas quando os vetores são iguais,

$$d(\mathbf{x}, \mathbf{y}) \geq 0 \quad \text{e} \quad d(\mathbf{x}, \mathbf{y}) = 0 \quad \text{sse} \quad \mathbf{x} = \mathbf{y}$$

e satisfaz a *desigualdade do triângulo*

$$d(\mathbf{x}, \mathbf{y}) \leq d(\mathbf{x}, \mathbf{z}) + d(\mathbf{z}, \mathbf{y}).$$

Observe que a fórmula explícita pela distância é

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

uma generalização natural do teorema de Pitágoras.

ex: Verifique que os vetores $\mathbf{e}_1 = (1, 0, \dots)$, $\mathbf{e}_2 = (0, 1, 0, \dots)$, $\dots$ da base canônica são ortogonais dois a dois, ou seja, $\mathbf{e}_i \cdot \mathbf{e}_j = 0$ se $i \neq j$.

ex: Calcule a componente de $\mathbf{x} = (1, 2)$ ao longo de $\mathbf{v} = (-1, 1)$, e a projeção de $\mathbf{x} = (0, 1, \pi)$ sobre $\mathbf{v} = (-2, 3, 0)$.

ex: Sejam $\mathbf{x}$ e $\mathbf{y}$ dois vetores não nulos, e considere os vetores unitários $\mathbf{u} = \mathbf{x}/\|\mathbf{x}\|$ e $\mathbf{v} = \mathbf{y}/\|\mathbf{y}\|$. Calcule $\|\mathbf{u} \pm \mathbf{v}\|^2$ e mostre que $-1 \leq \mathbf{u} \cdot \mathbf{v} \leq 1$. Deduza a desigualdade de Schwarz 2.1.

ex: Mostre que se θ é o ângulo entre $\mathbf{x}$ e $\mathbf{y}$ então

$$\|\mathbf{x} \pm \mathbf{y}\|^2 = \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 \pm 2\|\mathbf{x}\|\|\mathbf{y}\|\cos\theta$$

ex: Calcule o cosseno do ângulo entre $\mathbf{x} = (1, 2)$ e $\mathbf{y} = (-1, 1)$. Calcule o cosseno do ângulo entre $\mathbf{x} = (0, 1, \pi)$ e $\mathbf{y} = (-2, 3, 0)$.

ex: Calcule o cosseno dos ângulos do triângulo de vértices $A = (1, 1)$, $B = (-1, 3)$ e $C = (0, 2)$. Calcule o cosseno dos ângulos do triângulo de vértices $A = (1, 2, 5)$, $B = (2, 1, 2)$ e $C = (0, 3, 0)$.

ex: Determine um vetor ortogonal ao vetor $(1, -1)$, e um vetor ortogonal ao vetor $(1, 3, 6)$.

ex: Determine a família dos vetores (ou seja, todos os vetores) de $\mathbb{R}^2$ ortogonais ao vetor (a, b) . Determine a família dos vetores de $\mathbb{R}^3$ ortogonais ao vetor (a, b, c) .

ex: Prove que $d(\lambda\mathbf{x}, \lambda\mathbf{y}) = |\lambda| d(\mathbf{x}, \mathbf{y})$.

ex: Calcule a distância entre $\mathbf{x} = (1, 2)$ e $\mathbf{y} = (-1, 1)$, e a distância entre $\mathbf{x} = (0, 1, \pi)$ e $\mathbf{y} = (-2, 3, 0)$.

ex: [Ap69] 12.11.

Trabalho e energia cinética. O *trabalho (mecânico)* infinitesimal que realiza um campo de forças $\mathbf{F}$ ao deslocar uma partícula (ao longo do segmento) do ponto $\mathbf{r}$ ao ponto $\mathbf{r} + d\mathbf{r}$ (se o deslocamento $d\mathbf{r}$ é “pequeno” então o campo pode ser considerado constante) é igual ao produto escalar

$$dT := \mathbf{F} \cdot d\mathbf{r}.$$

Em particular, é nulo se o deslocamento é ortogonal à força. Usando a equação de Newton $\mathbf{F} = \dot{\mathbf{p}}$, com $\mathbf{p} = m\mathbf{v}$, e considerando m constante, temos que

$$dT = \mathbf{F} \cdot d\mathbf{r} = m \frac{d\mathbf{v}}{dt} \cdot \mathbf{v} dt = m \mathbf{v} \cdot d\mathbf{v} = m \frac{1}{2} d(\mathbf{v} \cdot \mathbf{v}).$$

Portanto $dT = dK$, ou seja, o trabalho infinitesimal é igual a variação infinitesimal da *energia cinética*

$$K := \frac{1}{2} m \|\mathbf{v}\|^2.$$

Ao integrar, temos o *teorema trabalho-energia*: o trabalho $T := \int_{t_0}^{t_1} \mathbf{F} \cdot d\mathbf{r}$ realizado sobre uma partícula por um campo de forças $\mathbf{F}$ é igual a variação $\Delta K := K(t_1) - K(t_0)$ da energia cinética.

Bolas e esferas. A *bola aberta* e a *bola fechada* (ou *círculo*, se $n = 2$) de centro $\mathbf{x} \in \mathbb{R}^n$ e raio $r > 0$ são

$$B_r(\mathbf{x}) := \{\mathbf{y} \in \mathbb{R}^n \text{ t.q. } \|\mathbf{y} - \mathbf{x}\| < r\} \quad \text{e} \quad \overline{B_r(\mathbf{x})} := \{\mathbf{y} \in \mathbb{R}^n \text{ t.q. } \|\mathbf{y} - \mathbf{x}\| \leq r\}$$

respetivamente. A *esfera* (ou *circunferência*, se $n = 2$) de centro $\mathbf{x} \in \mathbb{R}^n$ e raio $r > 0$ é

$$S_r(\mathbf{x}) := \{\mathbf{y} \in \mathbb{R}^n \text{ t.q. } \|\mathbf{y} - \mathbf{x}\| = r\}$$

Em particular, a *esfera unitária* de dimensão $n - 1$ é o conjunto

$$\mathbb{S}^{n-1} := S_1(\mathbf{0}) = \{\mathbf{x} \in \mathbb{R}^n \text{ t.q. } \|\mathbf{x}\| = 1\}$$

dos vetores unitários de $\mathbb{R}^n$.

ex: Determine uma equação cartesiana da circunferência de centro $(2, -1)$ e raio 7.

ex: Diga quando (ou seja, para quais valores dos raios r e r' dependendo dos centros $\mathbf{x}$ e $\mathbf{x}'$) a interseção $B_r(\mathbf{x}) \cap B_{r'}(\mathbf{x}')$ é $\neq \emptyset$.

ex: Verifique que $B_r(\mathbf{x}) = T_{\mathbf{x}}(M_{\mathbf{r}}(B_1(\mathbf{0})))$ e $S_r(\mathbf{x}) = T_{\mathbf{x}}(M_{\mathbf{r}}(\mathbb{S}^{n-1}))$.

Centroide. O *centroide* do sistema de pontos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N \in \mathbb{R}^n$ é o ponto

$$\mathbf{C} := \frac{1}{N} (\mathbf{x}_1 + \mathbf{x}_2 + \dots + \mathbf{x}_N)$$

(o centro de massa de um sistema de partículas de massas unitárias colocadas nas posições $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N$). Observe que em dimensão $n = 1$ o centróide da coleção de números $x_1, x_2, \dots, x_N$ é a *média aritmética* $\bar{x} := (x_1 + x_2 + \dots + x_N)/N$.

ex: Mostre que o centróide é o ponto $\mathbf{y}$ que minimiza a função

$$\mathcal{E}(\mathbf{y}) = \sum_{k=1}^N \|\mathbf{x}_k - \mathbf{y}\|^2.$$

ex: Calcule o centróide do sistema composto pelos pontos $(0, 1)$, $(2, 2)$ e $(3, 0)$ do plano.

ex: Mostre que o centróide de 3 pontos do plano, A , B e C , é a interseção dos segmentos que unem os vértices do triângulo ABC aos pontos médios dos lados opostos.

3 Retas e planos

ref: [Ap69] Vol 1, 13.1-8 ; [La97] Ch. I, 5-6

3.1 Equações paramétricas e cartesianas

6 out 2022

Equações paramétricas e cartesianas. Curvas, superfícies e outros subconjuntos de $\mathbb{R}^n$ podem ser definidos de forma *paramétrica*, ou seja, como imagens

$$A = f(S) = \{f(s) \text{ com } s \in S\}$$

de funções $f : S \rightarrow \mathbb{R}^n$ definidas em espaços de “parâmetros” $S = [0, 1], \mathbb{R}, \mathbb{R}^2, \dots$ (por exemplo, uma “trajetória” $t \mapsto \mathbf{r}(t) \in \mathbb{R}^3$, onde t é o “tempo”), ou de forma *cartesiana*, ou seja, como “lugares geométricos”

$$B = \{\mathbf{x} \in \mathbb{R}^n \text{ t.q. } f_1(\mathbf{x}) = 0, f_2(\mathbf{x}) = 0, \dots\}$$

dos pontos de $\mathbb{R}^n$ onde as funções $f_1 : \mathbb{R}^n \rightarrow \mathbb{R}, f_2 : \mathbb{R}^n \rightarrow \mathbb{R}, \dots$ se anulam.

ex: Descreva e esboce os seguintes subconjuntos de $\mathbb{R}^2$ ou $\mathbb{R}^3$.

$$\begin{aligned} &\{(x, y) \in \mathbb{R}^2 \text{ t.q. } x^2 + y^2 = 2\} \quad \{(x, y) \in \mathbb{R}^2 \text{ t.q. } xy = 0\} \\ &\{(1, 1) + t(0, 3) \text{ com } t \in \mathbb{R}\} \quad \{(0, 4, 0) + t(2, 3, 4) \text{ com } t \in \mathbb{R}\} \\ &\{(\cos t, \sin t) \text{ com } t \in [0, 2\pi]\} \quad \{(\cos t, \sin t, s) \text{ com } t \in [0, 2\pi], s \in \mathbb{R}\} \\ &\{(x, y) \in \mathbb{R}^3 \text{ t.q. } x^2 + y^2 + z^2 < \pi\} \quad \{(t, |t|) \text{ com } t \in [-1, 1]\} \quad \{(t, t^2) \text{ com } t \in \mathbb{R}\} \\ &\{(x, y) \in \mathbb{R}^2 \text{ t.q. } \langle (1, 2), (x, y) \rangle = 0\} \quad \{(x, y) \in \mathbb{R}^2 \text{ t.q. } 2x - 3y = 0\} \\ &\{(x, y) \in \mathbb{R}^2 \text{ t.q. } \langle (3, -1), (x, y) \rangle = 2\} \quad \{(x, y) \in \mathbb{R}^2 \text{ t.q. } 2x - 3y = 1\} \\ &\{(x, y) \in \mathbb{R}^2 \text{ t.q. } -3x + y = 0 \text{ e } x - 7y = 0\} \quad \{(x, y, z) \in \mathbb{R}^3 \text{ t.q. } x + y + z = 0\} \\ &\{(x, y) \in \mathbb{R}^2 \text{ t.q. } x + y = 2 \text{ e } 2x - y = 1\} \quad \{(x, y, z) \in \mathbb{R}^3 \text{ t.q. } x + y + z = 1\} \\ &\{(x, y, z) \in \mathbb{R}^3 \text{ t.q. } 2x - 3y - z = 0 \text{ e } x + y + 11z = 3\} \end{aligned}$$

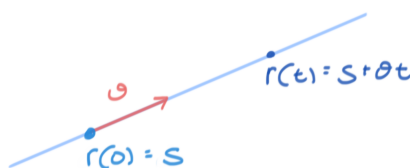
Partícula livre. A trajetória $t \mapsto \mathbf{r}(t) \in \mathbb{R}^3$ de uma partícula livre de massa m num referencial inercial é modelada pela equação de Newton

$$\frac{d}{dt}(m\mathbf{v}) = 0, \quad \text{ou seja, se } m \text{ é constante,} \quad m\mathbf{a} = 0,$$

onde $\mathbf{v}(t) := \dot{\mathbf{r}}(t)$ denota a velocidade da partícula no instante t , e $\mathbf{a}(t) := \ddot{\mathbf{r}}(t)$ denota a aceleração da partícula no instante t . Em particular, o “momento linear” $\mathbf{p} := m\mathbf{v}$, é uma constante do movimento, de acordo com o princípio de inércia de Galileio¹⁰ ou a primeira lei de Newton¹¹. As soluções da equação de Newton da partícula livre são as retas afins

$$\mathbf{r}(t) = \mathbf{s} + \mathbf{v}t$$

onde $\mathbf{s}, \mathbf{v} \in \mathbb{R}^3$ são a posição inicial e a velocidade (inicial), respetivamente.



¹⁰ “... il mobile durasse a muoversi tanto quanto durasse la lunghezza di quella superficie, né erta né china; se tale spazio fusse interminato, il moto in esso sarebbe parimenti senza termine, cioè perpetuo.”

[Galileo Galilei, *Dialogo sopra i due massimi sistemi del mondo*, 1623]

¹¹ “Corpus omne perseverare in statu suo quiescendi vel movendi uniformiter in directum, nisi quatenus a viribus impressis cogitur statum illum mutare.”

[Isaac Newton, *Philosophiæ Naturalis Principia Mathematica*, 1687]

ex: Determine a trajetória de uma partícula livre que passa, no instante $t_0 = 0$, pela posição $\mathbf{r}(0) = (3, 2, 1)$ com velocidade $\dot{\mathbf{r}}(0) = (1, 2, 3)$.

ex: Determine a trajetória de uma partícula livre que passa pela posição $\mathbf{r}(0) = (0, 0, 0)$ no instante $t_0 = 0$ e pela posição $\mathbf{r}(2) = (1, 1, 1)$ no instante $t_1 = 2$. Calcule a sua “velocidade escalar” (em inglês, *speed*), ou seja, a norma $v = \|\mathbf{v}\|$.

3.2 Retas

Retas. Um vetor não nulo $\mathbf{v} \in \mathbb{R}^n$ define/gera uma *reta*

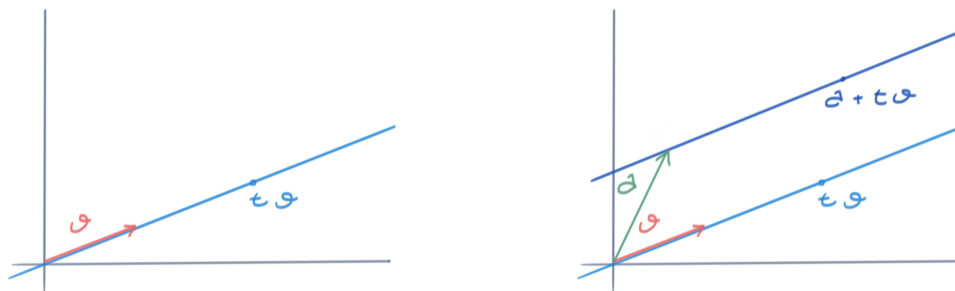
$$\mathbb{R}\mathbf{v} := \{t\mathbf{v} \text{ com } t \in \mathbb{R}\}$$

passando pela origem, o conjunto formado pelos vetores proporcionais a $\mathbf{v}$. Dois vetores não nulos e proporcionais geram a mesma reta passando pela origem. De fato, se $\mathbf{w} = \lambda\mathbf{v}$ com $\lambda \neq 0$, então $t\mathbf{v} = s\mathbf{w}$ quando $t = s\lambda$, e portanto $\mathbb{R}\mathbf{v} = \mathbb{R}\mathbf{w}$ (as retas diferem apenas pela parametrização, ou seja, a velocidade).

A *reta (afim)* paralela ao vetor não nulo $\mathbf{v} \in \mathbb{R}^n$ que passa pelo ponto $\mathbf{a} \in \mathbb{R}^n$ é

$$\mathbf{a} + \mathbb{R}\mathbf{v} := \{\mathbf{a} + t\mathbf{v} \text{ com } t \in \mathbb{R}\}$$

O vetor $\mathbf{v}$ é dito *vetor direccional* da reta, e pode ser pensado como a velocidade da trajetória/lei horária $\mathbf{r}(t) = \mathbf{a} + t\mathbf{v}$ de uma partícula em movimento retilíneo uniforme, interpretando o parâmetro t como um “tempo”. O ponto $\mathbf{a}$ pode então ser pensado como a “posição inicial” da partícula. É importante observar que dois pontos $\mathbf{r}(t)$ e $\mathbf{r}(s)$ da reta coincidem sse $t = s$, pois $\mathbf{r}(t) - \mathbf{r}(s) = (t - s)\mathbf{v}$, e o vetor $\mathbf{v}$ não é nulo. Isto permite falar de uma “ordem” entre os pontos da reta (a ordem temporal, uma das duas possíveis), e em particular dizer que quando $s < u < t$ então o ponto $\mathbf{r}(u)$ está “entre” os pontos $\mathbf{r}(s)$ e $\mathbf{r}(t)$.



Duas retas afins, $\mathbf{a} + \mathbb{R}\mathbf{v}$ e $\mathbf{b} + \mathbb{R}\mathbf{w}$, são ditas *paralelas* quando $\mathbf{v}$ e $\mathbf{w}$ são proporcionais. Dadas estas definições, puramente algébricas, de pontos e retas, os axiomas e os teoremas da geometria de Euclides são simples consequências.

É claro que duas retas afins $\mathbf{a} + \mathbb{R}\mathbf{v}$ e $\mathbf{b} + \mathbb{R}\mathbf{w}$ são iguais quando são paralelas e têm (pelo menos) um ponto comum (por exemplo, quando $\mathbf{b}$ pertence à primeira reta, ou $\mathbf{a}$ à segunda). De fato, se $\mathbf{w} = \lambda\mathbf{v}$ com $\lambda \neq 0$ e existem tempos t_0 e s_0 tais que $\mathbf{a} + t_0\mathbf{v} = \mathbf{b} + s_0\mathbf{w}$, então

$$\mathbf{b} + s\mathbf{w} = \mathbf{a} + t_0\mathbf{v} + (s - s_0)\mathbf{w} = \mathbf{a} + t_0\mathbf{v} + (s - s_0)\lambda\mathbf{v} = \mathbf{a} + (t_0 + \lambda(s - s_0))\mathbf{v}$$

e, vice-versa,

$$\mathbf{a} + t\mathbf{v} = \mathbf{b} + s_0\mathbf{w} + (t - t_0)\mathbf{v} = \mathbf{b} + s_0\mathbf{w} + (t - t_0)\lambda^{-1}\mathbf{w} = \mathbf{b} + (s_0 + \lambda^{-1}(t - t_0))\mathbf{w}$$

Uma consequência é que dada uma reta $\mathbf{a} + \mathbb{R}\mathbf{v}$ e um ponto $\mathbf{b}$, existe uma única reta paralela a primeira e passando por $\mathbf{b}$, que é a reta $\mathbf{b} + \mathbb{R}\mathbf{v}$ (uma afirmação equivalente ao “quinto postulado” de Euclides, logo falsa nas geometrias não-euclidianas ...).

Outra consequência é que por dois pontos distintos $\mathbf{a}$ e $\mathbf{b}$ passa uma e uma única reta, a reta

$$\mathbf{a} + \mathbb{R}(\mathbf{b} - \mathbf{a})$$

De fato, é claro que esta reta passa pelos dois pontos. Por outro lado, se $\mathbf{c} + \mathbb{R}\mathbf{w}$ é uma reta passando por $\mathbf{a}$ e $\mathbf{b}$, então existem tempos t_0 e t_1 , necessariamente distintos, tais que $\mathbf{a} = \mathbf{c} + t_0\mathbf{w}$ e $\mathbf{b} = \mathbf{c} + t_1\mathbf{w}$. Consequentemente, $\mathbf{b} - \mathbf{a} = (t_1 - t_0)\mathbf{w}$, ou seja, o vetor $\mathbf{w}$ é proporcional a $\mathbf{b} - \mathbf{a}$. Tendo um ponto em comum, as duas retas coincidem.

Uma forma mais simétrica da equação paramétrica da reta passando por dois pontos distintos $\mathbf{a}$ e $\mathbf{b}$ é obtida se observamos que $\mathbf{a} + t(\mathbf{b} - \mathbf{a}) = (1 - t)\mathbf{a} + t\mathbf{b}$ e chamamos $t_1 = (1 - t)$ e $t_2 = t$. O resultado é que esta reta pode ser descrita como o conjunto dos pontos da forma

$$t_1\mathbf{a} + t_2\mathbf{b}$$

com tempos $t_1, t_2 \in \mathbb{R}$ que satisfazem $t_1 + t_2 = 1$.

No espaço euclidiano $\mathbb{R}^n$, munido do produto escalar (2.1), é também possível definir ângulos e ortogonalidade. O ângulo entre as retas afins $\mathbf{a} + \mathbb{R}\mathbf{v}$ e $\mathbf{b} + \mathbb{R}\mathbf{w}$ é o ângulo entre os vetores direccionais $\mathbf{v}$ e $\mathbf{w}$, ou seja, o único $\theta \in [0, \pi)$ cujo coseno é $\cos \theta = \mathbf{v} \cdot \mathbf{w} / (\|\mathbf{v}\| \|\mathbf{w}\|)$. Assim, as retas são ditas *ortogonais* quando $\mathbf{w} \cdot \mathbf{v} = 0$.

Retas no plano. No plano $\mathbb{R}^2$, é possível eliminar o parâmetro t das equações paramétricas $(x(t), y(t)) = \mathbf{a} + t\mathbf{v}$ de uma reta, e deduzir uma equação cartesiana. Por exemplo, se $\mathbf{a} = (a_1, a_2)$ e $\mathbf{v} = (v_1, v_2)$, então

$$\begin{cases} x = a_1 + tv_1 \\ y = a_2 + tv_2 \end{cases}$$

Para eliminar o “tempo” t podemos multiplicar a primeira equação por v_2 , a segunda por v_1 ,

$$\begin{cases} v_2x = v_2a_1 + tv_1v_2 \\ v_1y = v_1a_2 + tv_2v_1 \end{cases}$$

e depois calcular a diferença das duas equações. O resultado é a equação cartesiana

$$v_2x - v_1y = v_2a_1 - v_1a_2$$

com coeficientes v_1 e v_2 não todos nulos. Finalmente,

$$\mathbf{a} + \mathbb{R}\mathbf{v} = \{(x, y) \in \mathbb{R}^2 \text{ t.q. } v_2(x - a_1) - v_1(y - a_2) = 0\}$$

Retas no plano euclidiano. No plano euclidiano $\mathbb{R}^2$, munido do produto escalar (2.1), um vetor não nulo $\mathbf{n} \in \mathbb{R}^2$ define uma *reta normal*

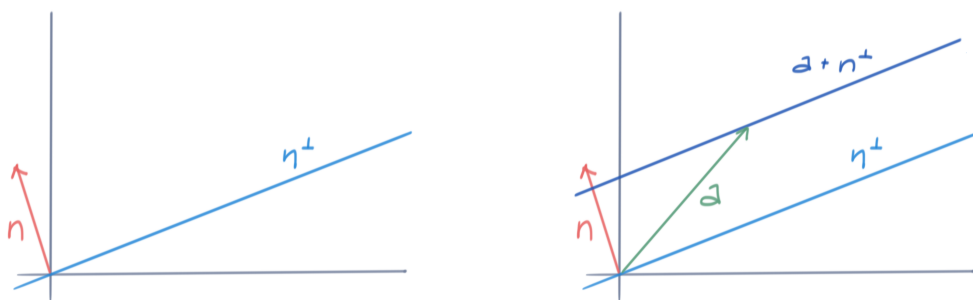
$$\mathbf{n}^\perp := \{\mathbf{x} \in \mathbb{R}^2 \text{ t.q. } \mathbf{x} \cdot \mathbf{n} = 0\}$$

formada pelos vetores ortogonais a $\mathbf{n}$. A reta perpendicular/normal ao vetor não nulo $\mathbf{n} \in \mathbb{R}^2$ que passa pelo ponto $\mathbf{a} \in \mathbb{R}^2$ é

$$\mathbf{a} + \mathbf{n}^\perp := \{\mathbf{x} \in \mathbb{R}^2 \text{ t.q. } (\mathbf{x} - \mathbf{a}) \cdot \mathbf{n} = 0\}$$

e o vetor $\mathbf{n}$ é dito vetor *normal* à reta. Por exemplo, se $\mathbf{a} = (a_1, a_2)$ e $\mathbf{n} = (n_1, n_2)$, então

$$\mathbf{a} + \mathbf{n}^\perp = \{(x, y) \in \mathbb{R}^2 \text{ t.q. } n_2(x - a_1) + n_1(y - a_2) = 0\}$$



ex: Mostre que se $L = \mathbf{a} + \mathbb{R}\mathbf{v}$ e $L' = \mathbf{a}' + \mathbb{R}\mathbf{v}'$ são duas retas paralelas, então existe um vetor $\mathbf{r}$ tal que $L' = \mathbf{r} + L$.

ex: Determine uma equação paramétrica da reta que passa por $(2, 3)$ e é paralela a $(-1, 2)$

ex: Determine uma equação paramétrica da reta que passa por $(5, 1, -2)$ e é paralela a $(3, -7, 2)$.

ex: Determine uma equação paramétrica da reta que passa pelos pontos $(3, 3)$ e $(-1, -1)$.

ex: Determine uma equação paramétrica da reta que passa pelos pontos $(0, 3, 4)$ e $(8, 3, 2)$.

ex: Determine uma equação paramétrica da reta $2x - 3y = 5$ do plano.

ex: Determine uma equação cartesiana da reta que passa por $(5, -1)$ e é paralela a $(-6, 2)$.

ex: Determine uma equação cartesiana da reta que passa por $(0, 0)$ e é perpendicular a $(-2, -3)$.

ex: Determine uma equação cartesiana da reta $(-2, 3) + \mathbb{R}(5, 1)$.

ex: Calcule o (coseno do) ângulo entre as retas $x - y = 0$ e $-x + y = -7$.

ex: Determine um vetor normal à reta que passa pelos pontos $(3, 0)$ e $(2, 1)$.

ex: Determine um vetor normal à reta $5x - 3y = 2$ do plano.

ex: Determine $P \in \mathbb{R}^2$ e $\mathbf{v} \in \mathbb{R}^2$ tais que

$$\{(x, y) \in \mathbb{R}^2 \text{ t.q. } x + 2y = -1\} = \{P + t\mathbf{v} \text{ com } t \in \mathbb{R}\}$$

ex: As retas

$$\{(x, y) \in \mathbb{R}^2 \text{ t.q. } 2x - 3y = 5\} \quad \text{e} \quad \{(x, y) \in \mathbb{R}^2 \text{ t.q. } 3x - 2y = 5\}$$

$$\{(x, y) \in \mathbb{R}^2 \text{ t.q. } x + 7y = 3\} \quad \text{e} \quad \{(x, y) \in \mathbb{R}^2 \text{ t.q. } -2x - 14y = 0\}$$

são paralelas? São perpendiculares?

ex: Determine as interseções entre as retas $x - 2y = 1$ e $-2x + 4y = 3$.

ex: Determine as interseções entre as retas $3x + 5y = 0$ e $x - y = -1$.

ex: Determine as interseções entre as retas $(3, 1) + \mathbb{R}(1, 3)$ e $(0, 1) + \mathbb{R}(-1, -2)$.

ex: Determine a família das retas paralelas ao vetor $\mathbf{v} = (a, b)$ do plano.

ex: Determine a família das retas que passam pelo ponto (a, b) do plano.

ex: Verifique que homotetias e translações enviam cada reta numa reta paralela.

ex: [Ap69] 13.5.

Ternos pitagóricos, método da corda de Diofanto. Um *terno Pitagórico* é uma solução inteira da equação

$$X^2 + Y^2 = Z^2$$

Tais $X, Y, Z \in \mathbb{Z}$ representam os comprimentos inteiros dos lados de um triângulo retângulo. Um exemplo bem conhecido é 3, 4, 5. Um terno Pitagórico é equivalente a uma solução racional (ou seja, com $x, y \in \mathbb{Q}$) de

$$x^2 + y^2 = 1$$

(basta dividir por $Z \neq 0$), e portanto um ponto racional $(x, y) \in \mathbb{Q}^2$ da circunferência unitária $\mathbb{S} := \{(x, y) \in \mathbb{R}^2 \text{ t.q. } x^2 + y^2 = 1\} \subset \mathbb{R}^2$.

ex: Determine a (outra) interseção entre uma reta que passa pelo ponto $(-1, 0)$ e a circunferência unitária $\mathbb{S} := \{(x, y) \in \mathbb{R}^2 \text{ t.q. } x^2 + y^2 = 1\}$.

ex: Mostre que quando o declive d da reta é racional, ou seja $d = U/V$ com $U, V \in \mathbb{Z}$, a interseção determina uma solução inteira de $X^2 + Y^2 = Z^2$. Verifique que esta solução corresponde à solução de Euclides

$$X = (U^2 - V^2)W, \quad Y = 2UVW, \quad Z = (U^2 + V^2)W,$$

com $U, V, W \in \mathbb{Z}$.

Distância entre um ponto e uma reta. No espaço euclidiano $\mathbb{R}^n$, munido do produto escalar (2.1), é possível definir distâncias entre subconjuntos. A distância entre $A \subset \mathbb{R}^n$ e $B \subset \mathbb{R}^n$ é

$$d(A, B) := \inf_{\mathbf{x} \in A, \mathbf{y} \in B} \|\mathbf{x} - \mathbf{y}\|$$

Por exemplo, é possível calcular a distância entre um ponto e uma reta. Seja $\mathbf{r}(t) = \mathbf{a} + t\mathbf{v}$, com $t \in \mathbb{R}$, a representação paramétrica dos pontos de uma reta em $\mathbb{R}^n$ passando pelo ponto $\mathbf{a}$ e paralela ao vetor não nulo $\mathbf{v}$. O quadrado da distância entre $\mathbf{r}(t)$ e um ponto $\mathbf{b}$ (por exemplo a origem) é um polinômio de segundo grau no tempo t ,

$$\|\mathbf{r}(t) - \mathbf{b}\|^2 = t^2 \|\mathbf{v}\|^2 + 2((\mathbf{a} - \mathbf{b}) \cdot \mathbf{v})t + \|\mathbf{a} - \mathbf{b}\|^2.$$

Esta função assume um mínimo quando o tempo é igual a $\bar{t} = -\mathbf{a} \cdot \mathbf{v} / \|\mathbf{v}\|^2$ (basta igualar a zero a derivada em ordem a t). Portanto o ponto da reta mais próximo do ponto $\mathbf{b}$ é

$$\mathbf{r}(\bar{t}) = \mathbf{a} + \frac{(\mathbf{b} - \mathbf{a}) \cdot \mathbf{v}}{\|\mathbf{v}\|^2} \mathbf{v}$$

e a sua distância é $\|\mathbf{b} - \mathbf{r}(\bar{t})\|$. Um cálculo mostra que este é o (único) ponto da reta onde $\mathbf{r}(\bar{t}) - \mathbf{b}$ é perpendicular ao vetor $\mathbf{v}$.

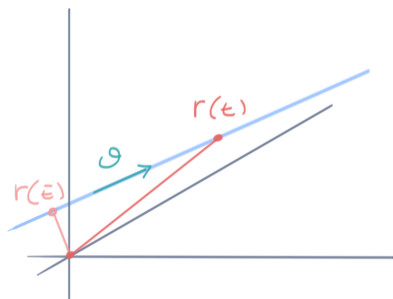
Outra forma de ver as coisas, sem utilizar ideias do cálculo, é considerar a representação de $\mathbf{b} - \mathbf{a}$ como soma

$$\mathbf{b} - \mathbf{a} = \lambda \mathbf{v} + \mathbf{w}$$

de um vetor $\lambda \mathbf{v}$ paralelo a $\mathbf{v}$ e de um vetor $\mathbf{w} = \mathbf{b} - (\mathbf{a} + \lambda \mathbf{v})$ ortogonal a $\mathbf{v}$. Pela (2.10), a componente de $\mathbf{b} - \mathbf{a}$ na direção de $\mathbf{V}$ é $\lambda = ((\mathbf{b} - \mathbf{a}) \cdot \mathbf{v}) / \|\mathbf{v}\|^2$. Então, se $\mathbf{r}(t)$ é qualquer outro ponto da reta, a diferença $\mathbf{r}(t) - (\mathbf{a} + \lambda \mathbf{v})$ é proporcional a $\mathbf{v}$, logo ortogonal a $\mathbf{w}$. Pelo teorema de Pitágoras

$$\|\mathbf{b} - \mathbf{r}(t)\|^2 = \|\mathbf{r}(t) - (\mathbf{a} + \lambda \mathbf{v})\|^2 + \|\mathbf{w}\|^2 \geq \|\mathbf{w}\|^2 = \|\mathbf{b} - (\mathbf{a} + \lambda \mathbf{v})\|^2$$

Consequentemente, o ponto $(\mathbf{a} + \lambda \mathbf{v})$, que corresponde a $\mathbf{r}(\bar{t})$, é o ponto da reta mais próximo de $\mathbf{b}$, e a sua distância é igual a $\|\mathbf{b} - (\mathbf{a} + \lambda \mathbf{v})\| = \|\mathbf{b} - \mathbf{r}(\bar{t})\|$.



Em particular, no plano euclidiano $\mathbb{R}^2$, a distância entre o ponto $\mathbf{x}$ e a reta $R = \mathbf{a} + \mathbf{n}^\perp$, normal ao vetor não nulo $\mathbf{n}$ e passando pelo ponto $\mathbf{a}$, é a norma da projeção de $\mathbf{x} - \mathbf{a}$ sobre o vetor normal $\mathbf{n}$, ou seja

$$d(\mathbf{r}, R) = \frac{|(\mathbf{x} - \mathbf{a}) \cdot \mathbf{n}|}{\|\mathbf{n}\|}$$

Se a reta é $R = \{(x, y) \in \mathbb{R}^2 \text{ t.q. } mx + ny + c = 0\}$, então

$$d((x, y), R) = \frac{|mx + ny + c|}{\sqrt{m^2 + n^2}}$$

ex: Mostre que $\|\mathbf{a} + t\mathbf{v}\| \geq \|\mathbf{a}\|$ para cada tempo $t \in \mathbb{R}$ sse $\mathbf{a} \cdot \mathbf{v} = 0$ (calcule o quadrado da norma de $\mathbf{a} + t\mathbf{v}$). Dê uma interpretação geométrica.

ex: Mostre que a norma de cada ponto $\mathbf{r} \in \mathbf{a} + \mathbf{n}^\perp$ da reta que passa por $\mathbf{a} \in \mathbb{R}^2$ e é perpendicular ao vetor não nulo $\mathbf{n} \in \mathbb{R}^2$ verifica

$$\|\mathbf{r}\| \geq d = \frac{|\mathbf{a} \cdot \mathbf{n}|}{\|\mathbf{n}\|}$$

e que $\|\mathbf{r}\| = d$ sse $\mathbf{r}$ é a projeção de $\mathbf{a}$ sobre $\mathbf{n}$, ou seja $\mathbf{r} = \frac{|\mathbf{a} \cdot \mathbf{n}|}{\|\mathbf{n}\|^2} \mathbf{n}$ (observe que a equação da reta é $\mathbf{r} \cdot \mathbf{n} = \mathbf{a} \cdot \mathbf{n}$ e use a desigualdade de Schwarz).

ex: Calcule a distância entre o ponto $(8, -3)$ e a reta $(1, 0) + \mathbb{R}(3, 3)$.

ex: Calcule a distância entre o ponto $(2, 4)$ e a reta $x - y = 0$.

ex: [Ap69] 13.5.

3.3 Planos

Planos. Dois vetores $\mathbf{v}, \mathbf{w} \in \mathbb{R}^n$ não proporcionais (em particular, não nulos pois um vetor nulo é proporcional a qualquer outro vetor) geram um *plano*

$$\mathbb{R}\mathbf{v} + \mathbb{R}\mathbf{w} := \{t\mathbf{v} + s\mathbf{w} \text{ com } t, s \in \mathbb{R}\} \subset \mathbb{R}^n \quad (3.1)$$

passando pela origem. Os vetores $\mathbf{v}$ e $\mathbf{w}$ são ditos *geradores* do plano, e os “tempos” t e s do ponto genérico $\mathbf{r}(t, s) = t\mathbf{v} + s\mathbf{w}$ podem ser pensados como “coordenadas”. Cada ponto do plano é determinado por um único par de coordenadas. De fato, se $t\mathbf{v} + s\mathbf{w} = t'\mathbf{v} + s'\mathbf{w}$, então $(t - t')\mathbf{v} = (s' - s)\mathbf{w}$, mas se $t \neq t'$ ou $s \neq s'$ isto significa que os vetores $\mathbf{v}$ e $\mathbf{w}$ são proporcionais. Em particular, a origem pode ser representada como o ponto $\mathbf{0} = t\mathbf{v} + s\mathbf{w}$ do plano apenas quando $t = s = 0$. isto também significa que a única interseção entre as retas $\mathbb{R}\mathbf{v}$ e $\mathbb{R}\mathbf{w}$ é a origem. Portanto podemos pensar este plano como a reunião disjunta das retas paralelas a $\mathbf{v}$ e passando pelos pontos da reta $\mathbb{R}\mathbf{w}$.

Os vetores $\mathbf{v}$ e $\mathbf{w}$ não são os únicos geradores do plano $\mathbb{R}\mathbf{v} + \mathbb{R}\mathbf{w}$, mas podem ser substituídos por qualquer outro par de vetores não proporcionais do plano. De fato, sejam

$$\mathbf{v}' = t_0\mathbf{v} + s_0\mathbf{w} \quad \text{e} \quad \mathbf{w}' = t_1\mathbf{v} + s_1\mathbf{w}$$

dois vetores não nulos do plano. Se $\mathbf{v}'$ e $\mathbf{w}'$ são proporcionais, ou seja $\mathbf{v}' = \lambda\mathbf{w}'$ com $\lambda \neq 0$ (a menos de trocar $\mathbf{v}'$ com $\mathbf{w}'$), então

$$t_0\mathbf{v} + s_0\mathbf{w} = \lambda(t_1\mathbf{v} + s_1\mathbf{w})$$

e portanto

$$(t_0 - \lambda t_1)\mathbf{v} + (s_0 - \lambda s_1)\mathbf{w} = \mathbf{0}$$

Isto implica que $t_0 - \lambda t_1 = s_0 - \lambda s_1 = 0$ (pois a única maneira de gerar a origem com $\mathbf{v}$ e $\mathbf{w}$ é escolher coordenadas nulas), e portanto que $s_0 t_1 - s_1 t_0 = 0$. Vice-versa, se $s_0 t_1 = s_1 t_0$ então

$$t_1 \mathbf{v}' = t_1 t_0 \mathbf{v} + t_1 s_0 \mathbf{w} = t_1 t_0 \mathbf{v} + s_1 t_0 \mathbf{w} = t_0 \mathbf{w}'$$

e

$$t_0 \mathbf{w}' = t_0 t_1 \mathbf{v} + t_0 s_1 \mathbf{w} = t_0 t_1 \mathbf{v} + s_0 t_1 \mathbf{w} = t_1 \mathbf{v}'$$

e portanto os vetores $\mathbf{v}'$ e $\mathbf{w}'$ são proporcionais (se $t_0 \neq 0$ ou $t_1 \neq 0$ pelas relações acima, se $t_0 = t_1 = 0$ porque são os dois proporcionais a $\mathbf{w}$). Consequentemente, $\mathbf{v}'$ e $\mathbf{w}'$ não são proporcionais sse $s_0 t_1 - s_1 t_0 \neq 0$. Esta é precisamente a condição que permite resolver o sistema

$$\begin{cases} \mathbf{v}' = t_0\mathbf{v} + s_0\mathbf{w} \\ \mathbf{w}' = t_1\mathbf{v} + s_1\mathbf{w} \end{cases}$$

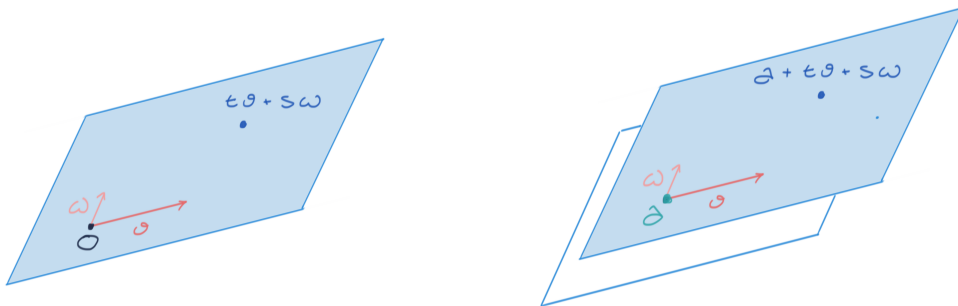
e obter umas representações de $\mathbf{v}$ e $\mathbf{w}$ como combinações lineares de $\mathbf{v}'$ e $\mathbf{w}'$, ou seja

$$\begin{cases} \mathbf{v} = \frac{1}{s_0 t_1 - s_1 t_0} (-s_1 \mathbf{v}' + s_0 \mathbf{w}') \\ \mathbf{w} = \frac{1}{s_0 t_1 - s_1 t_0} (t_1 \mathbf{v}' - t_0 \mathbf{w}') \end{cases}$$

(para obter $\mathbf{v}$ devem multiplicar a primeira identidade por s_1 , a segunda por s_0 , e calcular a diferença entre as duas, para obter $\mathbf{w}$ devem multiplicar a primeira identidade por t_1 , a segunda por t_0 , e calcular a diferença entre as duas) Consequentemente, todo ponto $t\mathbf{v} + s\mathbf{w}$ pode ser representado como $t'\mathbf{v}' + s'\mathbf{w}'$, com certas coordenadas t' e s' que dependem de t e s , e vice-versa. Finalmente, se $\mathbf{v}'$ e $\mathbf{w}'$ são dois vetores não proporcionais do plano gerado por $\mathbf{v}$ e $\mathbf{w}$, então $\mathbb{R}\mathbf{v} + \mathbb{R}\mathbf{w} = \mathbb{R}\mathbf{v}' + \mathbb{R}\mathbf{w}'$.

O plano (afim) gerado pelos vetores $\mathbf{v}$ e $\mathbf{w}$ que passa pelo ponto $\mathbf{a} \in \mathbb{R}^n$ é

$$\mathbf{a} + (\mathbb{R}\mathbf{v} + \mathbb{R}\mathbf{w}) := \{\mathbf{a} + t\mathbf{v} + s\mathbf{w} \text{ com } t, s \in \mathbb{R}\}$$



O plano que passa por três pontos distintos e não colineares $\mathbf{a}$, $\mathbf{b}$ e $\mathbf{c}$ é, por exemplo, o plano

$$\mathbf{a} + \mathbb{R}(\mathbf{b} - \mathbf{a}) + \mathbb{R}(\mathbf{c} - \mathbf{a})$$

Uma forma mais simétrica da equação paramétrica deste plano é

$$t_1\mathbf{a} + t_2\mathbf{b} + t_3\mathbf{c}$$

com tempos $t_1, t_2, t_3 \in \mathbb{R}$ que satisfazem $t_1 + t_2 + t_3 = 1$.

Planos no espaço de dimensão 3. No espaço $\mathbb{R}^3$, é possível eliminar os dois “tempos” t e s das equações paramétricas (3.1) de um plano e deduzir uma equação cartesiana. Por exemplo, o ponto genérico do plano $\mathbf{a} + \mathbb{R}\mathbf{v} + \mathbb{R}\mathbf{w}$ gerado pelos vetores não proporcionais $\mathbf{v} = (v_1, v_2, v_3)$ e $\mathbf{w} = (w_1, w_2, w_3)$ e passando por $\mathbf{a} = (a_1, a_2, a_3)$ é $\mathbf{r} = t\mathbf{v} + s\mathbf{w}$, ou seja, tem coordenadas

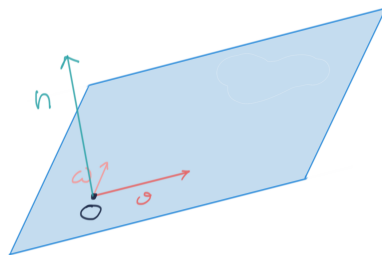
$$\begin{cases} x = a_1 + tv_1 + sw_1 \\ y = a_2 + tv_2 + sw_2 \\ z = a_3 + tv_3 + sw_3 \end{cases}$$

Para eliminar os dois parâmetros t e s contemporaneamente, podemos multiplicar a primeira equação por $\alpha = v_1w_2 - v_2w_1$, a segunda por $\beta = v_3w_1 - v_1w_3$ e a terceira por $\gamma = v_1w_2 - v_2w_1$. Ao somar as três equações resultantes acontece que ficamos apenas (para a explicação deste aparente milagre devem esperar a introdução do produto vetorial) com a equação cartesiana

$$\alpha x + \beta y + \gamma z = \delta$$

com coeficientes não todos nulos, e termo constante $\delta = \alpha a_1 + \beta a_2 + \gamma a_3$.

Planos no espaço euclidiano de dimensão 3. No espaço euclidiano $\mathbb{R}^3$, munido do produto escalar (2.1), um vetor não nulo $\mathbf{n} \in \mathbb{R}^3$ define um plano normal $\mathbf{n}^\perp := \{\mathbf{x} \in \mathbb{R}^3 \text{ t.q. } \mathbf{x} \cdot \mathbf{n} = 0\}$, formado pelos vetores ortogonais a $\mathbf{n}$.



O plano ortogonal/perpendicular/normal ao vetor não nulo $\mathbf{n} \in \mathbb{R}^3$ que passa pelo ponto $\mathbf{a} \in \mathbb{R}^3$ é

$$\mathbf{a} + \mathbf{n}^\perp := \{\mathbf{x} \in \mathbb{R}^3 \text{ t.q. } (\mathbf{x} - \mathbf{a}) \cdot \mathbf{n} = 0\}$$

O vetor $\mathbf{n}$, que é definido a menos de um múltiplo não nulos, é dito *vetor normal* ao plano, e define a reta normal ao plano. Em particular, é sempre possível escolher um vetor normal unitário.

Por exemplo, uma equação cartesiana do plano perpendicular ao vetor $\mathbf{n} = (m, n, p) \in \mathbb{R}^3$ que passa pelo ponto $\mathbf{a} = (a, b, c) \in \mathbb{R}^3$ é

$$m(x - a) + n(y - b) + p(z - c) = 0$$

O ângulo entre dois planos de $\mathbb{R}^3$ pode ser definido como sendo o ângulo entre dois vetores normais aos planos.

ex: Determine uma equação paramétrica do plano que passa pela origem e é gerado pelos vetores $(3, -7, 2)$ e $(-1, 0, -1)$.

ex: Determine uma equação paramétrica do plano que passa pelo ponto $(0, 3, 4)$ e é gerado pelos vetores $(0, 5, 0)$ e $(8, 3, 2)$.

ex: Determine uma equação paramétrica do plano $\{(x, y, z) \in \mathbb{R}^3 \text{ t.q. } x + y + z = 1\}$.

ex: Determine uma equação paramétrica do plano $\{(x, y, z) \in \mathbb{R}^3 \text{ t.q. } z = 0\}$.

ex: Determine uma equação cartesiana do plano que passa por $(-1, 1, 11)$ e é gerado pelos vetores $(1, 0, 0)$ e $(0, 1, 0)$.

ex: Determine uma equação cartesiana do plano que passa por $(0, 0, 0)$ e é gerado pelos vetores $(3, -7, 2)$ e $(-1, 0, -1)$.

ex: Determine uma equação cartesiana do plano que passa por $(3, 3, 3)$ e é paralelo ao plano $x + y + z = 0$.

ex: Determine uma equação cartesiana do plano que passa pelos pontos $(0, 3, 4)$, $(0, 5, 0)$ e $(8, 3, 2)$.

ex: Determine uma equação cartesiana do plano que passa por $(0, 0, 0)$ e é perpendicular ao vetor $(-2, -3, -4)$.

ex: Determine uma equação cartesiana do plano que passa por $(2, 1, 0)$ e é perpendicular ao vetor $(9, 3, 0)$.

ex: Calcule o (coseno do) ângulo entre os planos $x - y + z = 0$ e $-x + 3y + 5z = -7$.

ex: Determine um vetor normal ao plano que passa pelos pontos $(0, 0, 0)$, $(1, 0, 0)$ e $(0, 1, 0)$.

ex: Determine um vetor normal ao plano $x + y + z = 1$.

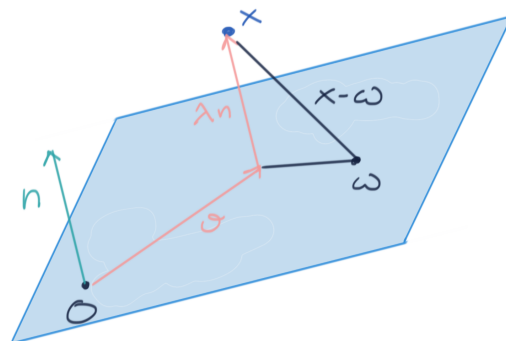
ex: Determine as intersecções entre os planos $x + 2y + 3z = -1$ e $-2x + 4y - z = 3$.

ex: [Ap69] 13.8 e 13.17.

Distância entre um ponto e um plano em $\mathbb{R}^3$. Consideramos, no espaço euclidiano $\mathbb{R}^3$, munido do produto escalar (2.1), um plano $\mathbf{n}^\perp$ passando pela origem e normal ao vetor não nulo $\mathbf{n}$. Todo vetor $\mathbf{x} \in \mathbb{R}^3$ pode ser representado como soma $\mathbf{x} = \lambda \mathbf{n} + \mathbf{v}$ da projeção de $\mathbf{x}$ na direção de $\mathbf{n}$ (com $\lambda = \mathbf{x} \cdot \mathbf{n} / \|\mathbf{n}\|^2$) e de um vetor $\mathbf{v}$ ortogonal a $\mathbf{n}$, logo um vetor do plano $\mathbf{n}^\perp$. Se $\mathbf{w}$ é qualquer outro ponto do plano $\mathbf{n}^\perp$, então o teorema de Pitágoras diz que

$$\|\mathbf{x} - \mathbf{w}\|^2 = \|\lambda \mathbf{n}\|^2 + \|\mathbf{w} - \mathbf{v}\|^2 \geq \|\lambda \mathbf{n}\|^2$$

Consequentemente, a distância entre $\mathbf{x}$ e qualquer vetor do plano é sempre superior ou igual a norma da projeção $\lambda \mathbf{n}$, e esta distância mínima é atingida precisamente quando $\mathbf{v} = \mathbf{w}$.



Em geral, a distância entre o ponto $\mathbf{x}$ e o plano $P = \mathbf{a} + \mathbf{n}^\perp$ é portanto igual a norma da projeção de $\mathbf{x} - \mathbf{a}$ sobre $\mathbf{n}$, ou seja,

$$d(\mathbf{x}, P) = \frac{|(\mathbf{x} - \mathbf{a}) \cdot \mathbf{n}|}{\|\mathbf{n}\|}$$

Se o plano é $P = \{(x, y, z) \in \mathbb{R}^3 \text{ t.q. } mx + ny + pz + c = 0\}$, então

$$d((x, y, z), P) = \frac{|mx + ny + pz + c|}{\sqrt{m^2 + n^2 + p^2}}$$

ex: Calcule a distância entre o ponto $(2, 4, 1)$ e o plano $x + y + z = 0$.

ex: Calcule a distância entre o ponto $(5, 7, 1)$ e o plano que passa por $(2, 0, 3)$ e é normal ao vetor $(1, 1, 0)$.

ex: Calcule a distância entre o ponto $(15, 11, 17)$ e o plano $x-y$.

ex: [Ap69] 13.17.

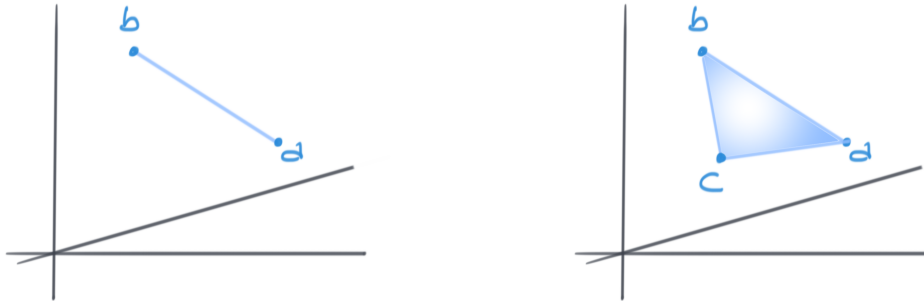
3.4 Segmentos e convexos

Conjuntos convexos. O segmento (*afim*) entre os pontos $\mathbf{a}$ e $\mathbf{b}$ de $\mathbb{R}^n$ é

$$\overline{\mathbf{ab}} = \{t\mathbf{a} + (1-t)\mathbf{b} \text{ com } t \in [0, 1]\}$$

(a órbita de uma partícula livre que viaja com velocidade $\mathbf{v} = \mathbf{b} - \mathbf{a}$ de uma posição inicial $\mathbf{a}$ durante um tempo $0 \leq t \leq 1$). Uma equação paramétrica mais simétrica é $t_1\mathbf{a} + t_2\mathbf{b}$ com $t_1 \geq 0$ e $t_2 \geq 0$ que satisfazem $t_1 + t_2 = 1$.

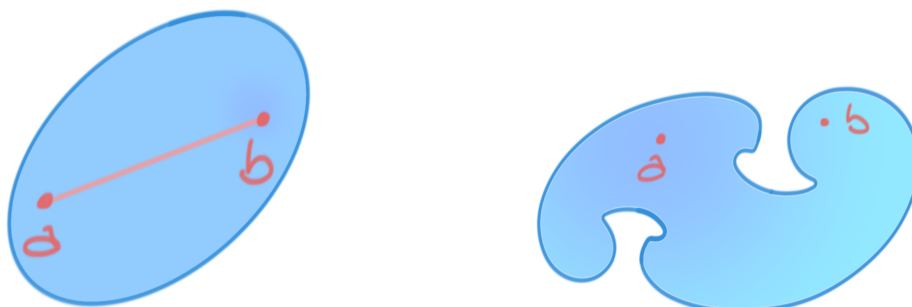
É natural definir o *triângulo* de vértices $\mathbf{a}, \mathbf{b}$ e $\mathbf{c}$ (não colineares) como sendo o conjunto dos pontos $t_1\mathbf{a} + t_2\mathbf{b} + t_3\mathbf{c}$ com coeficientes $t_1 \geq 0, t_2 \geq 0$ e $t_3 \geq 0$ que satisfazem $t_1 + t_2 + t_3 = 1$. Observe que esta condição também implica que $0 \leq t_i \leq 1$, e portanto o “peso” de cada vértice nestas combinações lineares pode ser pensado como uma “probabilidade”. É também claro que este triângulo é a reunião dos segmentos entre $\mathbf{c}$ e os pontos do segmento $\overline{\mathbf{ab}}$.



Pontos, segmentos e triângulos admitem a seguinte generalização natural. A *combinação convexa* dos pontos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m \in \mathbb{R}^n$ é o conjunto

$$\text{Conv}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m) = \{t_1\mathbf{x}_1 + t_2\mathbf{x}_2 + \dots + t_m\mathbf{x}_m \text{ com } t_i \geq 0 \text{ e } t_1 + t_2 + \dots + t_m = 1\}$$

Podemos pensar que os coeficientes t_k 's são “probabilidades”, de acordo com as quais os diferentes pontos $\mathbf{x}_k$'s são “pesados”. O nome tem a ver com a noção de “convexidade”. Um subconjunto $C \subset \mathbb{R}^n$ é *convexo* se contém o segmento entre cada par de seus pontos, ou seja, se $\mathbf{x}, \mathbf{y} \in C$ implica que também $t_1\mathbf{x} + t_2\mathbf{y} \in C$ se $t_1 \geq 0$ e $t_2 \geq 0$ com $t_1 + t_2 = 1$.



O próprio $\mathbb{R}^n$ é convexo. Pontos, segmentos e triângulos são também convexos. É claro, em geral, que combinações convexas são conjuntos convexos. De fato, sejam $\mathbf{a}$ e $\mathbf{b}$ dois pontos de $\text{Conv}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m)$, assim que $\mathbf{a} = \sum_i t_i \mathbf{x}_i$ e $\mathbf{b} = \sum_i s_i \mathbf{x}_i$ com coeficientes não negativos t_i e s_j tais que $\sum_i t_i = \sum_i s_i = 1$. Então, se $0 \leq t \leq 1$,

$$(1-t)\mathbf{a} + t\mathbf{b} = \sum_i ((1-t)t_i + ts_i) \mathbf{x}_i$$

Os coeficientes $(1-t)t_i + ts_i$ são também não negativos e somam

$$\sum_i (1-t)t_i + ts_i = (1-t) \sum_i t_i + t \sum_i s_i = (1-t) + t = 1$$

Isto prova que também $(1-t)\mathbf{a} + t\mathbf{b} \in \text{Conv}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m)$ se $0 \leq t \leq 1$.

Fecho convexo. É claro que a interseção de uma família arbitrária de convexos é um conjunto convexo. Consequentemente, por cada subconjunto $A \subset \mathbb{R}^n$, existe um convexo minimal que contém A (a interseção de todos os convexos que contêm A) chamado *envoltória/invólucro/fecho convexa/o* de A , e denotado por $\text{Conv}(A)$. Calcular o fecho convexo de um subconjunto genérico pode ser uma tarefa difícil. Um caracterização algébrica simples é possível quando o conjunto é formado por um número finito de pontos.

Teorema 3.1. *O menor convexo que contém os pontos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m \in \mathbb{R}^n$ é a combinação convexa $\text{Conv}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m)$*

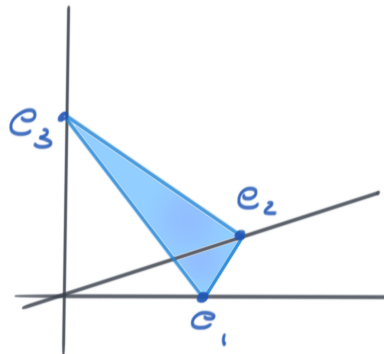
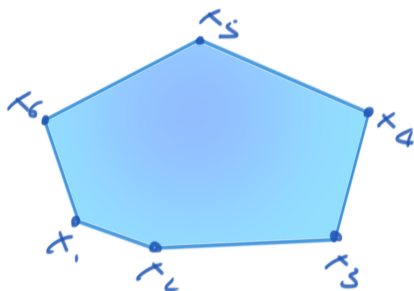
Demonstração. A prova é por indução sobre a cardinalidade m , e depende da observação que $\text{Conv}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m, \mathbf{x}_{m+1})$ pode ser pensado como a reunião dos segmentos entre $\mathbf{x}_{m+1}$ e os pontos de $\text{Conv}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m)$.

Se $m = 1$, o resultado é trivial. Assumimos o resultado válido para m pontos, e consideramos um conjunto convexo C que contém os $m+1$ pontos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{m+1}$. Seja $\mathbf{x} = \sum_{i=1}^{m+1} t_i \mathbf{x}_i$ um ponto de $\text{Conv}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{m+1})$, logo com $t_i \geq 0$ e $\sum_i t_i = 1$. Se $t_{m+1} = 1$, então $\mathbf{x} = \mathbf{x}_{m+1}$, e claramente $\mathbf{x} \in C$. Se, por outro lado, $t_{m+1} \neq 1$, podemos dividir por $(1-t_{m+1})$ e representar

$$\mathbf{x} = (1-t_{m+1}) \left(\sum_{i=1}^m \frac{t_i}{1-t_{m+1}} \mathbf{x}_i \right) + t_{m+1} \mathbf{x}_{m+1}$$

ou seja, como ponto do segmento entre $\mathbf{x}_{m+1}$ e $\mathbf{y} = \sum_{i=1}^m \frac{t_i}{1-t_{m+1}} \mathbf{x}_i$. Mas $\mathbf{y}$ é um ponto de $\text{Conv}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m)$, pois $\sum_{i=1}^m t_i = 1-t_{m+1}$. Pela hipótese indutiva, $\mathbf{y} \in C$ (pois C é um convexo que contém os $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m$). Pela convexidade de C , também $\mathbf{x} \in C$. $\square$

Se os $\mathbf{x}_k$'s são pontos do plano, então o fecho convexo é um polígono convexo. Um caso particular importante é o fecho convexo dos vetores $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ da base canônica de $\mathbb{R}^n$, chamado *simplexo unitário* e denotado por $\Delta^{n-1} = \text{Conv}(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n) \subset \mathbb{R}^n$, que é claramente um objeto de dimensão $n-1$.



ex: Mostre que segmentos e triângulos são convexos.

ex: O *paralelogramo* definido pelos vetores $\mathbf{a}$ e $\mathbf{b}$, o conjunto das combinações lineares $t_1\mathbf{a} + t_2\mathbf{b}$ com $0 \leq t_1 \leq 1$ e $0 \leq t_2 \leq 1$, é convexo.

ex: No espaço euclidiano $\mathbb{R}^n$ munido do produto escalar (2.1), as bolas abertas $B_r(\mathbf{x})$ ou fechadas $\overline{B_r(\mathbf{x})}$ são convexas.

ex: Dados um vetor não nulo $\mathbf{n}$ o espaço euclidiano $\mathbb{R}^n$ munido do produto escalar (2.1), e um escalar $b \in \mathbb{R}$, os *semi-espacos* $\{\mathbf{x} \in \mathbb{R}^n \text{ t.q. } \mathbf{n} \cdot \mathbf{x} \geq b\}$ ou $\{\mathbf{x} \in \mathbb{R}^n \text{ t.q. } \mathbf{n} \cdot \mathbf{x} > b\}$ são convexas.

ex: Translações e homotetias preservam os convexas, ou seja, se $C \subset \mathbb{R}^n$ é convexo, então também $C + \mathbf{a}$ e λC são convexas, para todo vetor $\mathbf{a} \in \mathbb{R}^n$ e todo escalar $\lambda \in \mathbb{R}$.

Medidas de probabilidades. Uma (*medida de*) *probabilidade* num “espaço dos acontecimentos” finito $\Omega = \{\omega_1, \omega_2, \dots, \omega_n\}$ é uma função $\mathbb{P} : 2^\Omega := \{\text{subconjuntos } A \subset \Omega\} \rightarrow [0, 1]$ aditiva, i.e. tal que

$$\mathbb{P}(A \cup B) = \mathbb{P}(A) + \mathbb{P}(B) \quad \text{se } A \cap B = \emptyset,$$

(a probabilidade do evento “ A ou B ” é igual à probabilidade do evento A mais a probabilidade do evento B se A e B são eventos mutuamente exclusivos) que verifica $\mathbb{P}(\emptyset) = 0$ (a probabilidade do “evento impossível” é nula) e $\mathbb{P}(\Omega) = 1$ (a probabilidade do “evento certo” é um).

Cada vetor $\mathbf{p} = (p_1, p_2, \dots, p_n) \in \mathbb{R}^n$ cujas coordenadas estão limitadas por $0 \leq p_i \leq 1$ e tal que $p_1 + p_2 + \dots + p_n = 1$ define uma probabilidade $\mathbb{P}$, por meio de

$$\mathbb{P}(A) = \sum_{\omega_i \in A} p_i$$

ou seja, $p_i = \mathbb{P}(\{\omega_i\})$. Portanto, o espaço das medidas de probabilidades em Ω é o *simplexo* (*unitário*)

$$\Delta^{n-1} := \{\mathbf{p} = (p_1, p_2, \dots, p_n) \in \mathbb{R}^n \text{ com } 0 \leq p_i \leq 1 \text{ e } p_1 + p_2 + \dots + p_n = 1\}. \quad (3.2)$$

4 Subespaços, bases e dimensão

ref: [Ap69] Vol 1, 12.12-17 ; [La97] Ch. III, 1-5

4.1 Subespaços e geradores

A generalização natural de retas e planos passando pela origem, e relativos geradores, é a noção de subespaço linear.

Combinações lineares. Uma *combinação linear* dos vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ de $\mathbb{R}^n$ com *coeficientes* $t_1, t_2, \dots, t_m \in \mathbb{R}$ (que são escalares, logo números reais neste caso) é um vetor

13 out 2022

$$\mathbf{v} = t_1 \mathbf{v}_1 + t_2 \mathbf{v}_2 + \dots + t_m \mathbf{v}_m$$

Um tal vetor $\mathbf{v}$, obtido como combinação linear dos $\mathbf{v}_k$'s, é também dito (*linearmente*) *dependente* dos $\mathbf{v}_k$'s.

ex: O vetor $\mathbf{v}$ é dependente do vetor $\mathbf{w}$ sse é proporcional a $\mathbf{w}$.

ex: Cada $\mathbf{v}_k$, com $1 \leq k \leq m$, é dependente dos vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$.

ex: Se $\mathbf{v}$ é dependente dos $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ mas não é dependente dos $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{m-1}$, então $\mathbf{v}_m$ é dependente dos $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{m-1}, \mathbf{v}$.

ex: Se $\mathbf{v}$ é dependente dos $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ e cada $\mathbf{v}_k$, com $1 \leq k \leq m$, é dependente dos $\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_p$, então $\mathbf{v}$ é dependente dos $\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_p$.

Subespaços e geradores. Um *subespaço vetorial/linear* de $\mathbb{R}^n$ é um subconjunto não vazio $\mathbf{V} \subset \mathbb{R}^n$ que é preservado pelas operações de soma de vetores e produto por um escalar, ou seja, tal que

$$\text{se } \mathbf{v}, \mathbf{w} \in \mathbf{V} \text{ e } \lambda \in \mathbb{R} \text{ então também } \mathbf{v} + \mathbf{w} \in \mathbf{V} \text{ e } \lambda \mathbf{v} \in \mathbf{V}.$$

Isto implica que se $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m \in \mathbf{V}$ e $t_1, t_2, \dots, t_m$ são escalares arbitrários então também

$$\sum_{k=1}^m t_k \mathbf{v}_k = t_1 \mathbf{v}_1 + t_2 \mathbf{v}_2 + \dots + t_m \mathbf{v}_m \in \mathbf{V}$$

Em outras palavras, um subespaço é um subconjunto não vazio de $\mathbb{R}^n$ que contém todas as possíveis combinações lineares dos seus vetores.

Em particular, todo subespaço de $\mathbb{R}^n$ contém o vetor nulo $\mathbf{0}$ (a combinação linear trivial, com todos os coeficientes nulos), e o subespaço minimal é o subespaço trivial $\{\mathbf{0}\}$. O subespaço maximal é o próprio $\mathbb{R}^n$. É claro também que a interseção $\cap_k \mathbf{V}_k$ de uma família de subespaços (finita ou não) é também um subespaço de $\mathbb{R}^n$.

Exemplos de subespaços próprios e não triviais de $\mathbb{R}^n$ são os subespaços do género

$$\mathbf{V}_m = \{ \mathbf{x} \in \mathbb{R}^n \text{ t.q. } x_k = 0 \forall k > m \}$$

com $1 \leq m < n$.

Retas e planos passando pela origem são exemplos de subespaços de $\mathbb{R}^n$. Uma maneira natural de construir subespaços, que generaliza a construção de retas e planos passando pela origem, é a seguinte. Seja $S \subset \mathbb{R}^n$ um subconjunto, finito ou não. O conjunto das combinações lineares finitas dos elementos de S

$$\text{Span}(S) := \{ t_1 \mathbf{s}_1 + t_2 \mathbf{s}_2 + \dots + t_m \mathbf{s}_m \text{ com } \mathbf{s}_k \in S, t_k \in \mathbb{R} \}$$

é um subespaço vetorial de $\mathbb{R}^n$, dito *subespaço gerado* por S , ou também *expansão linear* de S (em inglês, *linear span*). É o menor subespaço de $\mathbb{R}^n$ que contém S , ou seja a interseção de

todos os subespaços que contêm S (que é um conjunto não vazio, pois contém o próprio $\mathbb{R}^n$). Os vetores de S são então chamados *geradores* do subespaço $\text{Span}(S)$, pois qualquer vetor de $\text{Span}(S)$ é “gerado”, ou seja, “produzido”, por elementos de S ao fazer combinações lineares finitas.

Por exemplo, o subespaço gerado por um vetor não nulo $\mathbf{v}$ é a reta $\text{Span}(\mathbf{v}) = \mathbb{R}\mathbf{v}$. O subespaço gerado por dois vetores não nulos e não proporcionais (ou seja, não dependentes) $\mathbf{v}$ e $\mathbf{w}$ é o plano $\text{Span}(\mathbf{v}, \mathbf{w}) = \mathbb{R}\mathbf{v} + \mathbb{R}\mathbf{w}$. De fato, é tautológico que um subconjunto $V \subset \mathbb{R}^n$ é um subespaço sse contém todas as retas e os planos gerados pelos seus vetores.

Decidir se um dado vetor $\mathbf{v} \in \mathbb{R}^n$ pertence ao subespaço $\text{Span}(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m)$ gerado pelos m vetores $\mathbf{v}_k$'s significa decidir se existem constantes t_k 's tais que $\mathbf{v} = t_1\mathbf{v}_1 + t_2\mathbf{v}_2 + \dots + t_m\mathbf{v}_m$, ou seja, resolver um sistema de n equações (uma para cada coordenada) nas m incógnitas t_k 's. Se $n > m$ isto é muito pouco provável para vetores “genéricos” $\mathbf{v}$.

ex: Determine o subespaço de $\mathbb{R}^2$ gerado por $(3, -2)$ e $(-6, 4)$.

ex: Determine o subespaço de $\mathbb{R}^3$ gerado por $(1, 0, 0)$ e $(0, 1, 1)$.

ex: O vetor $(2, 3, 4)$ pertence ao subespaço gerado pelos vetores $(1, 2, 3)$ e $(1, 1, 1)$? E o vetor $(4, 3, 2)$?

ex: Diga se são subespaços vetoriais os seguintes subconjuntos de $\mathbb{R}^2$ ou $\mathbb{R}^3$

$$\{(t, 3t) \text{ com } t \in \mathbb{R}\} \quad \{(2t, 3t) \text{ com } t \in [0, 1]\} \quad \{(t-1, 2+3t) \text{ com } t \in \mathbb{R}\}$$

$$\{(t, 3s-t, s) \text{ com } (t, s) \in \mathbb{R}^2\} \quad \{(1-t, t+s, 5) \text{ com } (t, s) \in \mathbb{R}^2\}$$

$$\{(x, y) \in \mathbb{R}^2 \text{ t.q. } x-2y=0\} \quad \{(x, y) \in \mathbb{R}^2 \text{ t.q. } x-2y=1\}$$

$$\{(x, y, z) \in \mathbb{R}^3 \text{ t.q. } x-y+z=0 \text{ e } -2x+y-z=1\}$$

$$\{(x, y, z) \in \mathbb{R}^3 \text{ t.q. } x+y+z=0 \text{ e } x-y-3z=0\}$$

ex: A reunião de dois subespaços vetoriais é um subespaço vetorial?

Subespaços ortogonais. No espaço euclidiano $\mathbb{R}^n$, munido do produto escalar (2.1), subespaços podem ser definidos à custa de relações de ortogonalidade. Dado $\mathbf{n} \in \mathbb{R}^n$, o conjunto

$$\mathbf{n}^\perp := \{\mathbf{x} \in \mathbb{R}^n \text{ t.q. } \mathbf{n} \cdot \mathbf{x} = 0\}$$

é um subespaço linear de $\mathbb{R}^n$. Observe que $\mathbf{n}^\perp$ é o próprio $\mathbb{R}^n$ apenas quando $\mathbf{n}$ é o vetor nulo, pois se $\mathbf{n} \neq \mathbf{0}$ então $\mathbf{n}^\perp$ não contém a reta $\mathbb{R}\mathbf{n}$.

Dados $\mathbf{n}_1, \mathbf{n}_2, \dots, \mathbf{n}_m \in \mathbb{R}^n$, a interseção

$$\bigcap_{k=1}^m \mathbf{n}_k^\perp = \{\mathbf{x} \in \mathbb{R}^n \text{ t.q. } \mathbf{n}_i \cdot \mathbf{x} = 0 \quad \forall i = 1, 2, \dots, m\}$$

é também um subespaço linear de $\mathbb{R}^n$. Em geral, dado um subconjunto arbitrário $N \subset \mathbb{R}^n$ (não necessariamente um conjunto finito ou um subespaço), o conjunto

$$N^\perp := \{\mathbf{x} \in \mathbb{R}^n \text{ t.q. } \mathbf{n} \cdot \mathbf{x} = 0 \quad \forall \mathbf{n} \in N\}$$

é um subespaço linear de $\mathbb{R}^n$.

ex: Fixado $\mathbf{n} \in \mathbb{R}^n$, o conjunto $\{\mathbf{x} \in \mathbb{R}^n \text{ t.q. } \mathbf{n} \cdot \mathbf{x} = 1\}$ é um subespaço linear de $\mathbb{R}^n$?

ex: [Ap69] 12.15.

4.2 Conjuntos livres

Um conjunto minimal de geradores de um subespaço é chamado conjunto livre, ou linearmente independente.

Conjuntos livres/linearmente independentes. É conveniente definir a dependência linear, e portanto a sua negação, como propriedade de um conjunto de vetores.

O conjunto finito formado pelos vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ de $\mathbb{R}^n$ é dito (*linearmente*) *dependente* se existe uma combinação linear não trivial dos $\mathbf{v}_i$'s que representa o vetor nulo, ou seja, se existem coeficientes t_k 's, não todos nulos, tais que

$$t_1\mathbf{v}_1 + t_2\mathbf{v}_2 + \dots + t_m\mathbf{v}_m = \mathbf{0}$$

Isto significa que (pelo menos) um deles (um dos que têm coeficiente não nulo na expressão acima) é dependente dos outros, ou seja, pertence ao subespaço gerado pelos outros. Por exemplo, se $t_1 \neq 0$, então $\mathbf{v}_1 = -(t_2/t_1)\mathbf{v}_2 - \dots - (t_m/t_1)\mathbf{v}_m$.

É claro que todo conjunto que contém o vetor nulo é dependente, pois $\mathbf{0} + 0\mathbf{v}_2 + \dots + 0\mathbf{v}_m$ é uma combinação linear não trivial dos vetores $\mathbf{0}, \mathbf{v}_2, \dots, \mathbf{v}_m$ igual ao vetor nulo.

O conjunto finito formado pelos vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ é *livre*, ou (*linearmente*) *independente* (ou os vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ são *livres* ou (*linearmente*) *independentes*), se gera o vetor nulo duma única maneira, ou seja, se a única combinação linear nula

$$t_1\mathbf{v}_1 + t_2\mathbf{v}_2 + \dots + t_m\mathbf{v}_m = \mathbf{0}$$

é a combinação linear trivial, com coeficientes $t_1 = t_2 = \dots = t_m = 0$ (ou seja, se a única solução do sistema homogêneo acima, de m equações nas m incógnitas $t_1, t_2, \dots, t_m$, é a solução trivial).

Teorema 4.1. *O conjunto finito $S = \{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m\} \subset \mathbb{R}^n$ é livre sse gera cada vetor de $\text{Span}(S)$ duma única maneira, ou seja, se cada vetor $\mathbf{v} \in \text{Span}(S)$ admite uma única representação*

$$\mathbf{v} = t_1\mathbf{v}_1 + t_2\mathbf{v}_2 + \dots + t_m\mathbf{v}_m$$

como combinação linear dos $\mathbf{v}_k$'s.

Demonstração. Se $\mathbf{v} = s_1\mathbf{v}_1 + s_2\mathbf{v}_2 + \dots + s_m\mathbf{v}_m$ for outra representação, com pelo menos um $t_k \neq s_k$, então $(t_1 - s_1)\mathbf{v}_1 + (t_2 - s_2)\mathbf{v}_2 + \dots + (t_m - s_m)\mathbf{v}_m$ seria uma representação não trivial do vetor nulo. A outra implicação é evidente. $\square$

ex: Os vetores $\mathbf{v}$ e $\mathbf{w}$ são dependentes sse são paralelos.

ex: Um conjunto que contém o vetor nulo não é independente.

ex: Se $S \subset \mathbb{R}^n$ é independente, então qualquer subconjunto $T \subset S$ é independente. Equivalentemente, se $S \subset \mathbb{R}^n$ é dependente, então qualquer conjunto $T \supset S$ é também dependente.

ex: Se $\mathbf{s}_1, \mathbf{s}_2, \dots, \mathbf{s}_m$ são independentes mas $\mathbf{s}_1, \mathbf{s}_2, \dots, \mathbf{s}_m, \mathbf{v}$ são dependentes, então $\mathbf{v}$ é dependente dos $\mathbf{s}_1, \mathbf{s}_2, \dots, \mathbf{s}_m$.

ex: Mostre que os vetores (a, b) e (c, d) de $\mathbb{R}^2$ são independentes sse $ad - bc \neq 0$.

ex: Os vetores $(1, 2)$ e $(-1, 2)$ são independentes?

ex: Os vetores $(1, 1, 0)$, $(0, 1, 1)$ e $(1, 0, 1)$ são independentes?

ex: Os vetores $(1, 2)$, $(2, 3)$ e $(3, 4)$ são independentes?

ex: Os vetores $(\sqrt{2}, 1, 0)$, $(1, \sqrt{2}, 1)$ e $(0, 1, \sqrt{2})$ são independentes?

ex: Verifique se $(7, -1)$ é combinação linear de $(3, 8)$ e $(-1, 0)$.

ex: Verifique se $(1, 5, 3)$ é combinação linear de $(0, 1, 2)$ e $(-1, 0, 5)$.

ex: Determine os valores de t para os quais os vetores $(t, 1, 0)$, $(1, t, 1)$ e $(0, 1, t)$ são dependentes.

ex: [Ap69] 12.15.

4.3 Bases e dimensão

Todo subespaço é gerado por um número minimal de vetores, chamados base, cuja cardinalidade é chamada dimensão.

Bases e dimensão. É razoável esperar que um subespaço de $\mathbb{R}^n$, por exemplo o próprio $\mathbb{R}^n$, deve ser gerado por um número minimal de vetores, e deve conter um número maximal de vetores independentes. A observação importante é que o número de vetores independentes que cabem no subespaço gerado por um conjunto finito de vetores, independentes ou não, satisfaz um “princípio de conservação” ([Ap69] theorem 12.8 ou [La97] theorem 5.1).

Teorema 4.2. *No espaço gerado por m vetores $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m$ de $\mathbb{R}^n$ (independentes ou não) não cabem mais de m vetores independentes. Ou seja, todo conjunto formado por $m+1$ (ou mais) vetores $\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_{m+1}$ de $\text{Span}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m)$ é dependente.*

Demonstração. A prova é por indução sobre a cardinalidade m . O caso $m = 1$ é simples. Se $\mathbf{x}_1$ é o vetor nulo, o resultado é trivial. Se $\mathbf{x}_1$ não é nulo, e $\mathbf{y}_1 = t\mathbf{x}_1$ e $\mathbf{y}_2 = s\mathbf{x}_1$ são dois vetores não nulos, então é claro que $s\mathbf{y}_1 - t\mathbf{y}_2 = \mathbf{0}$, e portanto $\mathbf{y}_1$ e $\mathbf{y}_2$ são dependentes.

Assumimos então o teorema válido para $m-1$, e consideramos $m+1$ vetores $\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_{m+1}$ no espaço gerado pelos m vetores $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m$, assim que

$$\mathbf{y}_i = \sum_{j=1}^m a_{ij} \mathbf{x}_j \quad i = 1, 2, \dots, m+1 \quad (4.1)$$

com certos coeficientes a_{ij} . Se em todas estas expressões falta (ou seja, aparece sempre com coeficiente a_{im} nulo) $\mathbf{x}_m$, então os $\mathbf{y}_i$'s pertencem de fato ao espaço gerado pelos $m-1$ vetores $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{m-1}$. Pela hipótese indutiva são então dependentes.

Se, por outro lado, na expressão de um dos $\mathbf{y}$'s, que podemos considerar, a menos de reordenar os vetores, ser o último $\mathbf{y}_{m+1}$, aparece o $\mathbf{x}_m$ com coeficiente não nulo, então podemos dividir por este coeficiente e representar

$$\mathbf{x}_m = b\mathbf{y}_{m+1} + \sum_{j=1}^{m-1} c_{ij} \mathbf{x}_j$$

Podemos então substituir esta expressão nas restantes (4.1), e assim representar

$$\mathbf{y}_i = d_i \mathbf{y}_{m+1} + \sum_{j=1}^{m-1} e_{ij} \mathbf{x}_j \quad i = 1, 2, \dots, m$$

Isto significa que os m vetores $\mathbf{y}_i - d_i \mathbf{y}_{m+1}$, com $i = 1, 2, \dots, m$, pertencem ao espaço gerado pelos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{m-1}$. Pela hipótese indutiva são dependentes, e isto claramente implica que os $\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_{m+1}$ são também dependentes. $\square$

Uma *base* de $\mathbb{R}^n$ é um conjunto livre de geradores de $\mathbb{R}^n$, que convém pensar como uma lista ordenada de vetores. Em geral, uma base do subespaço linear $\mathbf{V} \subset \mathbb{R}^n$ é um conjunto livre de geradores de $\mathbf{V}$.

De acordo com o teorema 4.1, uma base de $\mathbb{R}^n$ é portanto uma lista $\mathcal{B} = (\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_m)$ de vetores $\mathbf{b}_k \in \mathbb{R}^n$ tal que cada vetor $\mathbf{x} \in \mathbb{R}^n$ admite uma única representação

$$\mathbf{x} = t_1 \mathbf{b}_1 + t_2 \mathbf{b}_2 + \dots + t_m \mathbf{b}_m$$

como combinação linear dos vetores de $\mathcal{B}$. Os coeficientes t_k são ditos *coordenadas/componentes* de $\mathbf{x}$ relativamente à base $\mathcal{B}$.

É tautológico que a *base canónica*, o conjunto formado pelos n vetores

$$\mathbf{e}_1 = (1, 0, \dots, 0) \quad \mathbf{e}_2 = (0, 1, 0, \dots, 0) \quad \dots \quad \mathbf{e}_n = (0, \dots, 0, 1),$$

é uma base de $\mathbb{R}^n$. De fato, todo vetor $\mathbf{x} = (x_1, x_2, \dots, x_n)$ é uma combinação linear única $\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_n \mathbf{e}_n$. Pelo teorema 4.2, qualquer outra base deve ser composta por $m \leq n$ vetores. Por outro lado, se uma base de $\mathbb{R}^n$ é formada por m vetores, sempre o teorema 4.2 implica que $n \leq m$, pois os vetores da base canónica são n vetores independentes gerados pela base. Consequentemente,

Teorema 4.3. *Toda a base de $\mathbb{R}^n$ é composta de n vetores.*

Da mesma forma, é evidente que toda a base de um subespaço vetorial $\mathbf{V} \subset \mathbb{R}^n$ é composta pelo mesmo número de vetores. A cardinalidade de uma (e portanto de qualquer) base de um subespaço $\mathbf{V} \subset \mathbb{R}^n$ é chamada *dimensão* do subespaço, e denotada por $\dim(\mathbf{V})$. Em particular, a dimensão do espaço $\mathbb{R}^n$ é $\dim(\mathbb{R}^n) = n$. A dimensão de um subespaço $\mathbf{V} \subset \mathbb{R}^n$ gerado por m vetores independentes é $\dim(\mathbf{V}) = m$. Por exemplo, a dimensão de uma reta passando pela origem é um, a dimensão de um plano passando pela origem é dois, ...

Teorema 4.4. *Todo conjunto formado por n vetores independentes é uma base de $\mathbb{R}^n$.*

Demonstração. Sejam $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ vetores linearmente independentes de $\mathbb{R}^n$. Se não geram o espaço total, então existe um vetor $\mathbf{x}_{n+1}$ que não é combinação linear dos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$. É claro que então os $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n, \mathbf{x}_{n+1}$ são independentes, mas isto contradiz o teorema 4.2, pois $\mathbb{R}^n$ é gerado pelos n vetores da base canónica. $\square$

Em particular, um subespaço de dimensão n de $\mathbb{R}^n$ é necessariamente o próprio $\mathbb{R}^n$. Finalmente, é útil a possibilidade de completar qualquer sistema independente até formar uma base.

Teorema 4.5. *Qualquer conjunto linearmente independente é um subconjunto de uma base de $\mathbb{R}^n$. Ou seja, se $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m$ são independentes e $m < n$, então existem vetores $\mathbf{x}_{m+1}, \mathbf{x}_{m+2}, \dots, \mathbf{x}_n$ tais que os $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m, \mathbf{x}_{m+1}, \dots, \mathbf{x}_n$ formam uma base de $\mathbb{R}^n$.*

Demonstração. Se o espaço gerado pelos vetores independentes $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m$ não é todo $\mathbb{R}^n$, então existe um vetor não nulo $\mathbf{x}_{m+1}$ que não é combinação linear dos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m$. É claro que então os $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m, \mathbf{x}_{m+1}$ são independentes. Se não geram todo $\mathbb{R}^n$, então existe um vetor $\mathbf{x}_{m+2}$ que não é combinação linear dos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{m+1}$... E assim a seguir. Esta construção deve terminar quando são atingidos n vetores independentes, pois apenas cabem n vetores independentes em $\mathbb{R}^n$, e estes formam claramente uma base. $\square$

Um corolário natural é que todo subespaço de $\mathbb{R}^n$ admite uma base.

ex: Os vetores $(1, 1)$ e $(1, -1)$ formam uma base de $\mathbb{R}^2$?

- ex:** Determine uma base de $\mathbb{R}^2$ contendo o vetor $(2, 3)$.
- ex:** Determine uma base de $\mathbb{R}^3$ contendo os vetores $(0, 1, 1)$ e $(1, 1, 1)$.
- ex:** Verifique se os vetores $(1, 0)$ e $(1, 1)$ formam uma base de $\mathbb{R}^2$.
- ex:** Verifique se os vetores $(0, 1, 2)$, $(-1, 0, 3)$ e $(1, 1, -1)$ formam uma base de $\mathbb{R}^3$.
- ex:** Calcule as coordenadas do vetor $(1, 1)$ relativamente a base de $\mathbb{R}^2$ formada pelos vetores $(1, 2)$ e $(-1, 2)$.
- ex:** Verifique se os vetores $(1, 0, 0, 0)$, $(1, 1, 0, 0)$, $(1, 1, 1, 0)$ e $(1, 1, 1, 1)$ formam uma base de $\mathbb{R}^4$.
- ex:** [Ap69] 12.15.

4.4 Bases e ortogonalidade

No espaço euclidiano, a ortogonalidade implica a independência.

Sistemas ortonormados. No espaço euclidiano $\mathbb{R}^n$, munido do produto escalar (2.1), há uma relação simples entre ortogonalidade e independência. Uma família (finita ou infinita) de vetores não nulos $\mathbf{v}_k$'s dois a dois ortogonais, ou seja, tais que $\mathbf{v}_i \cdot \mathbf{v}_j = 0$ se $i \neq j$, é dita *família/sistema ortogonal*.

Teorema 4.6. *Uma família ortogonal de vetores não nulos do espaço euclidiano $\mathbb{R}^n$ é independente.*

Demonstração. Sejam $\mathbf{v}_1, \dots, \mathbf{v}_n$ vetores não nulos e (dois a dois) ortogonais. Se

$$t_1 \mathbf{v}_1 + t_2 \mathbf{v}_2 + \dots + t_n \mathbf{v}_n = \mathbf{0}$$

então, ao calcular os produtos escalares com os vetores $\mathbf{v}_k$, obtemos

$$\begin{aligned} 0 &= \mathbf{v}_k \cdot (t_1 \mathbf{v}_1 + t_2 \mathbf{v}_2 + \dots + t_m \mathbf{v}_m) \\ &= t_k \|\mathbf{v}_k\|^2 \end{aligned}$$

(pois todos os outros produtos escalares são nulos pela ortogonalidade). Sendo os $\|\mathbf{v}_k\| > 0$, todos os coeficientes t_k são nulos. $\square$

Em particular, todo o conjunto ortogonal de n vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$ não nulos no espaço euclidiano $\mathbb{R}^n$ é uma base. Os vetores “normalizados” $\mathbf{u}_k = \mathbf{v}_k / \|\mathbf{v}_k\|$ formam então uma *base ortonormada*, ou seja, uma base formada por vetores ortogonais e unitários.

Por exemplo, a base canônica do espaço euclidiano $\mathbb{R}^n$, formada pelos vetores $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$, é ortonormada. De fato é possível construir muitas bases ortogonais, e portanto ortonormadas, de acordo com o seguinte algoritmo.

Teorema 4.7 (ortogonalização de Gram-Schmidt). *Seja $\mathbf{v}_1, \mathbf{v}_2, \dots$, um conjunto independente de vetores do espaço euclidiano $\mathbb{R}^n$. Então existe um conjunto ortogonal $\mathbf{u}_1, \mathbf{u}_2, \dots$ tal que, para todos $m = 1, 2, \dots$, o espaço $\text{Span}(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m)$ gerados pelo primeiros m vetores $\mathbf{v}$'s coincide com o espaço $\text{Span}(\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_m)$ gerado pelos primeiros n vetores $\mathbf{u}$'s.*

Demonstração. Basta definir o conjunto $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_n, \dots$ recursivamente por

$$\begin{aligned} \mathbf{u}_1 &= \mathbf{v}_1 & \mathbf{u}_2 &= \mathbf{v}_2 - \frac{\mathbf{v}_2 \cdot \mathbf{u}_1}{\|\mathbf{u}_1\|^2} \mathbf{u}_1 & \dots \\ \mathbf{u}_{k+1} &= \mathbf{v}_{k+1} - \left(\frac{\mathbf{v}_{k+1} \cdot \mathbf{u}_1}{\|\mathbf{u}_1\|^2} \mathbf{u}_1 + \frac{\mathbf{v}_{k+1} \cdot \mathbf{u}_2}{\|\mathbf{u}_2\|^2} \mathbf{u}_2 + \dots + \frac{\mathbf{v}_{k+1} \cdot \mathbf{u}_n}{\|\mathbf{u}_n\|^2} \mathbf{u}_n \right) & \dots \end{aligned}$$

Ou seja, $\mathbf{u}_{k+1}$ é obtido retirando de $\mathbf{v}_{k+1}$ a soma das suas projeções sobre os $\mathbf{u}_1, \dots, \mathbf{u}_k$. É imediato verificar que $\mathbf{u}_{k+1}$ não é nulo (caso contrário $\mathbf{v}_{k+1}$ seria uma combinação linear dos $\mathbf{v}_1, \dots, \mathbf{v}_k$) e é ortogonal ao subespaço $\text{Span}(\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_k)$, e portanto a todos os $\mathbf{u}_i$ com $i \leq k$. $\square$

Um corolário é que para todo subespaço $\mathbf{V} \subset \mathbb{R}^n$ de dimensão $\dim(\mathbf{V}) = m$ existe uma base ortonormada $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_m$ de $\mathbb{R}^n$ tal que $\mathbf{V} = \text{Span}(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_m)$. É então imediato verificar que os vetores $\mathbf{e}_{m+1}, \mathbf{e}_{m+2}, \dots, \mathbf{e}_n$ geram o subespaço ortogonal $\mathbf{V}^\perp$. Consequentemente, as dimensões de um subespaço e do subespaço ortogonal são complementares ou seja,

$$\dim(\mathbf{V}) + \dim(\mathbf{V}^\perp) = n \quad (4.2)$$

Coefficientes de Fourier. Se $\{\mathbf{s}_1, \mathbf{s}_2, \dots, \mathbf{s}_n\}$ é uma base ortogonal de $\mathbb{R}^n$, então cada vetor $\mathbf{x} \in \mathbb{R}^n$ é uma combinação linear única

$$\mathbf{x} = t_1 \mathbf{s}_1 + t_2 \mathbf{s}_2 + \dots + t_n \mathbf{s}_n$$

Se calculamos o produto escalar desta expressão com cada $\mathbf{s}_k$ obtemos

$$\mathbf{x} \cdot \mathbf{s}_k = t_k \|\mathbf{s}_k\|^2$$

pois os outros produtos escalares são nulos. Consequentemente, as componentes do vetor $\mathbf{x}$ são

$$t_k = \frac{\mathbf{x} \cdot \mathbf{s}_k}{\|\mathbf{s}_k\|^2}$$

Em particular, se a base é ortonormada, então as coordenadas de $\mathbf{x}$ são os “coeficientes de Fourier”

$$\boxed{t_k = \mathbf{x} \cdot \mathbf{s}_k}$$

assim que

$$\mathbf{x} = (\mathbf{x} \cdot \mathbf{s}_1) \mathbf{s}_1 + (\mathbf{x} \cdot \mathbf{s}_2) \mathbf{s}_2 + \dots + (\mathbf{x} \cdot \mathbf{s}_n) \mathbf{s}_n.$$

Por exemplo, a base canônica $\mathbf{i}, \mathbf{j}, \mathbf{k}$ é uma base ortonormada do espaço euclidiano $\mathbb{R}^3$. As coordenadas do vetor $\mathbf{r} = (x, y, z)$ relativamente a esta base são

$$x = \mathbf{r} \cdot \mathbf{i} \quad y = \mathbf{r} \cdot \mathbf{j} \quad z = \mathbf{r} \cdot \mathbf{k}$$

ex: Verifique se o conjunto formado pelos vetores $(0, \sqrt{3}/2, 1/2)$, $(0, -1/2, \sqrt{3}/2)$ e $(1, 0, 0)$ é ortogonal e ortonormado.

ex: Determine uma base ortogonal de $\mathbb{R}^2$ contendo o vetor $(1, 1)$.

ex: Determine uma base ortonormada de $\mathbb{R}^2$ contendo o vetor $(1/\sqrt{2}, 1/\sqrt{2})$.

ex: [Ap69] 12.15.

5 Produto vetorial, área e volume

ref: [Ap69] Vol 1, 13.9-17

20 out 2022

5.1 Independência no plano e determinante.

A independência de dois vetores do plano pode ser testada ao calcular apenas um número, chamado determinante, que representa uma área orientada.

Independência no plano e determinante. Os vetores $\mathbf{v} = (a, c)$ e $\mathbf{w} = (b, d)$ são independentes sse a única combinação linear $x\mathbf{v} + y\mathbf{w}$ nula é a combinação linear trivial, com $x = 0$ e $y = 0$, ou seja, sse a única solução do “sistema linear homogêneo”

$$\begin{cases} ax + by = 0 \\ cx + dy = 0 \end{cases}$$

é a solução trivial. Ao retirar b vezes a segunda equação de d vezes a primeira equação, e depois ao retirar c vezes a primeira equação de a vezes a segunda equação, temos que

$$\begin{cases} ax + by = 0 \\ cx + dy = 0 \end{cases} \Rightarrow \begin{cases} (ad - bc)x = 0 \\ (ad - bc)y = 0 \end{cases}$$

As quatro coordenadas dos dois vetores do plano podem ser arranjadas numa “matriz” (do latim MATRIX, madre ou útero, e portanto a/o que dá origem, que gera), neste caso uma tabela de duas linhas e duas colunas,

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

cujas colunas são as componentes dos vetores (a, c) e (b, d) (ou cujas linhas são as componentes dos vetores (a, b) e (c, d)).

O *determinante* desta matriz 2×2 é o número

$$\text{Det} \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{vmatrix} a & b \\ c & d \end{vmatrix} := ad - bc$$

Observe que trocar linhas com colunas não altera o determinante, ou seja,

$$\begin{vmatrix} a & b \\ c & d \end{vmatrix} = \begin{vmatrix} a & c \\ b & d \end{vmatrix}$$

Por outro lado, trocar a ordem das linhas ou das colunas produz uma mudança de sinal no determinante, ou seja,

$$\begin{vmatrix} a & b \\ c & d \end{vmatrix} = - \begin{vmatrix} b & a \\ d & c \end{vmatrix} = - \begin{vmatrix} c & d \\ a & b \end{vmatrix}$$

A conclusão da discussão anterior é que

Teorema 5.1. Os vetores (a, c) e (b, d) são independentes sse $\begin{vmatrix} a & b \\ c & d \end{vmatrix} \neq 0$.

ex: Diga se os vetores $(-1, 4)$ e $(3, -12)$ são independentes.

ex: Diga se os vetores $(5, 7)$ e $(2, 9)$ são independentes.

ex: Determine os valores do parâmetro real t tais que os vetores $(t, 1 - t)$ e $(1 + t, t)$ são independentes.

Determinante e área. O *paralelogramo* gerado/definido pelos vetores $\mathbf{v} = (a, c)$ e $\mathbf{w} = (b, d)$ de $\mathbb{R}^2$ é o conjunto $P = \{t\mathbf{v} + s\mathbf{w} \mid 0 \leq t, s \leq 1\} \subset \mathbb{R}^2$. A sua área é igual ao módulo do determinante da matriz $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$, ou seja,

$$\boxed{\text{Área}(P) = \left| \text{Det} \begin{pmatrix} a & b \\ c & d \end{pmatrix} \right| = |ad - bc|}$$

Naturalmente, não temos nesta altura uma definição rigorosa de área de uma região do plano, ou de volume de uma região do espaço, ... Apenas sabemos que a área de um quadrado de lados unitário é igual a 1, e portanto sabemos calcular áreas de retângulos e de triângulos retângulos usando as propriedades naturais que esperamos tenha uma área, aditividade e homogeneidade.

ex: Prove a fórmula acima para a área do paralelogramo (retire da área do retângulo de base $a + b$ e altura $c + d$ as áreas dos retângulos e triângulos que sobram ...).

ex: Calcule a área do paralelogramo definido pelos vetores $(0, 1)$ e $(1, 1)$, e do paralelogramo definido pelos vetores $(5, -2)$ e $(-3, 1)$.

ex: Calcule a área do triângulo de vértices $(3, 2)$, $(6, -4)$ e $(8, 8)$.

5.2 Produto vetorial

Produto vetorial. No espaço euclidiano $\mathbb{R}^3$ (e apenas no espaço desta dimensão, por acaso o espaço onde vivemos!), munido da base ortonormada canónica $\mathbf{i}, \mathbf{j}, \mathbf{k}$ (e a ordem é importante), é possível definir uma operação binária natural que associa a um par ordenado de vetores um terceiro vetor com um significado geométrico interessante. O *produto vetorial*, ou *externo* (em inglês *cross product*) dos vetores $\mathbf{r} = (x, y, z)$ e $\mathbf{r}' = (x', y', z')$ é o vetor

$$\boxed{\mathbf{r} \times \mathbf{r}' := (yz' - zy', -xz' + zx', xy' - yx')}$$
 (5.1)

Sendo cada coordenada do produto vetorial uma função linear das outras duas coordenadas dos dois vetores, é claro que o produto vetorial é distributivo sobre a adição e compatível com a multiplicação escalar. Em outras palavras, o produto vetorial é “bilinear”, ou seja, é linear em cada variável, no sentido em que satisfaz

$$(\lambda \mathbf{r} + \mu \mathbf{r}') \times \mathbf{r}'' = \lambda(\mathbf{r} \times \mathbf{r}'') + \mu(\mathbf{r}' \times \mathbf{r}'')$$

$$\mathbf{r} \times (\lambda \mathbf{r}' + \mu \mathbf{r}'') = \lambda(\mathbf{r} \times \mathbf{r}') + \mu(\mathbf{r} \times \mathbf{r}'')$$

Também evidente é que o produto vetorial é “anti-simétrico”, ou seja,

$$\mathbf{r} \times \mathbf{r}' = -\mathbf{r}' \times \mathbf{r}$$

Em particular, isto implica que o produto vetorial de um vetor com si próprio é o vetor nulo, ou seja, $\mathbf{r} \times \mathbf{r} = \mathbf{0}$.

Um cálculo mostra que os produtos vetoriais entre os vetores da base canónica $\mathbf{i} = (1, 0, 0)$, $\mathbf{j} = (0, 1, 0)$ e $\mathbf{k} = (0, 0, 1)$ são

$$\boxed{\mathbf{i} \times \mathbf{j} = \mathbf{k} \quad \mathbf{j} \times \mathbf{k} = \mathbf{i} \quad \text{e} \quad \mathbf{k} \times \mathbf{i} = \mathbf{j}} \quad (5.2)$$

(para decorar estas fórmulas basta observar a ordem cíclica dos vetores). De fato, a própria definição (5.1) é uma consequência destes três produtos básicos, assumindo a bilinearidade e a anti-simetria.

Mais uns cálculos elementares (mas chatos!) mostram que o produto vetorial satisfaz a *identidade de Jacobi*

$$\boxed{\mathbf{a} \times (\mathbf{b} \times \mathbf{c}) + \mathbf{b} \times (\mathbf{c} \times \mathbf{a}) + \mathbf{c} \times (\mathbf{a} \times \mathbf{b}) = \mathbf{0}} \quad (5.3)$$

(observe a ordem cíclica dos vetores também nesta fórmula) e a *identidade de Lagrange*

$$\|\mathbf{r} \times \mathbf{r}'\|^2 = \|\mathbf{r}\|^2 \|\mathbf{r}'\|^2 - (\mathbf{r} \cdot \mathbf{r}')^2 \quad (5.4)$$

Esta última fórmula revela o significado geométrico do produto vetorial. Diz que a norma do produto vetorial mede a discrepância entre as quantidades que entram na desigualdade de Schwarz 2.1. Em particular, a norma é nula, e portanto o próprio produto vetorial é nulo, sse os vetores são dependentes, ou seja, proporcionais.

Por outro lado, quando os vetores $\mathbf{r}$ e $\mathbf{r}'$ são independentes, e portanto geram um plano $\mathbb{R}\mathbf{r} + \mathbb{R}\mathbf{r}'$, então o produto vetorial não é nulo, e um cálculo elementar mostra que o vetor $\mathbf{r} \times \mathbf{r}'$ é ortogonal a este plano. Em particular, $\|\mathbf{r} \times \mathbf{r}'\| = \|\mathbf{r}\| \|\mathbf{r}'\|$ sse $\mathbf{r}$ e $\mathbf{r}'$ são ortogonais.

ex: Mostre que o produto vetorial é bilinear e anti-comutativo.

ex: Se $\mathbf{v} \times \mathbf{k} = 0$, o que pode concluir sobre o vetor $\mathbf{v}$? E se $\mathbf{v} \times \mathbf{i} = \mathbf{v} \times \mathbf{j} = 0$?

ex: Mostre que $\mathbf{r} \times \mathbf{r}' = \mathbf{0}$ sse $\mathbf{r}$ e $\mathbf{r}'$ são dependentes, ou seja, proporcionais (use a identidade de Lagrange e a desigualdade de Schwarz).

ex: Calcule $\mathbf{r} \cdot (\mathbf{r} \times \mathbf{r}')$ e $\mathbf{r}' \cdot (\mathbf{r} \times \mathbf{r}')$ e deduza que $\mathbf{r} \times \mathbf{r}'$ é ortogonal ao plano $\mathbb{R}\mathbf{r} + \mathbb{R}\mathbf{r}'$.

ex: Considere a equação paramétrica $\mathbf{r} = \mathbf{a} + t\mathbf{v} + s\mathbf{w}$ do ponto genérico do plano passando por $\mathbf{a} \in \mathbb{R}^3$ e gerado pelo vetores independentes $\mathbf{v}$ e $\mathbf{w}$. Calcule o seu produto escalar com o vetor $\mathbf{n} = \mathbf{v} \times \mathbf{w}$, normal ao plano, e deduza a equação cartesiana do plano $(\mathbf{r} - \mathbf{a}) \cdot \mathbf{n} = 0$.

ex: O produto vetorial não é associativo! Por exemplo, $\mathbf{i} \times (\mathbf{i} \times \mathbf{j}) \neq (\mathbf{i} \times \mathbf{i}) \times \mathbf{j}$.

ex: Verifique a identidade de (5.4). Ou seja, calcule explicitamente os dois polinômios de grau 4 nas coordenadas dos dois vetores, $\|\mathbf{r} \times \mathbf{r}'\|^2$ e $\|\mathbf{r}\|^2 \|\mathbf{r}'\|^2 - (\mathbf{r} \cdot \mathbf{r}')^2$, e verifique que são iguais.

ex: Verifique a *fórmula de Lagrange*

$$\mathbf{a} \times (\mathbf{b} \times \mathbf{c}) = \mathbf{b}(\mathbf{a} \cdot \mathbf{c}) - \mathbf{c}(\mathbf{a} \cdot \mathbf{b}) \quad (5.5)$$

ex: Calcule $\mathbf{r} \times \mathbf{r}'$ quando $\mathbf{r} = (1, -1, 1)$ e $\mathbf{r}' = (2, -2, 2)$ ou quando $\mathbf{r} = (-2, -1, 3)$ e $\mathbf{r}' = (\pi, -\pi, 0)$.

Produto vetorial e determinante. Cada coordenada do produto vetorial $\mathbf{r} \times \mathbf{r}'$ é igual, a menos de um sinal, ao determinante da matriz 2×2 obtida usando como linhas as duas outras coordenadas dos vetores $\mathbf{r}$ e $\mathbf{r}'$. De fato, uma representação formal do produto vetorial entre os vetores $\mathbf{r} = (x, y, z)$ e $\mathbf{r}' = (x', y', z')$ é

$$\mathbf{r} \times \mathbf{r}' = \text{“Det} \begin{pmatrix} \mathbf{i} & \mathbf{j} & \mathbf{k} \\ x & y & z \\ x' & y' & z' \end{pmatrix} \text{”} := \mathbf{i} \begin{vmatrix} y & z \\ y' & z' \end{vmatrix} - \mathbf{j} \begin{vmatrix} x & z \\ x' & z' \end{vmatrix} + \mathbf{k} \begin{vmatrix} x & y \\ x' & y' \end{vmatrix}$$

onde a última expressão define o determinante da “matriz” 3×3 formada usando como primeira linha os três vetores da base canônica e como segunda e terceira linhas as coordenadas dos dois vetores $\mathbf{r}$ e $\mathbf{r}'$ (naturalmente, esta tabela não deve ser chamada matriz, pois estamos a misturar nela elementos de natureza diferente).

ex: Calcule $\mathbf{r} \times \mathbf{r}'$ quando $\mathbf{r} = (3, -2, 8)$ e $\mathbf{r}' = (1, 1, 1)$ ou quando $\mathbf{r} = (-\pi, e, 10)$ e $\mathbf{r}' = (7, 5, 3)$

ex: [Ap69] 13.11.

Produto vetorial e área. A norma do produto vetorial $\mathbf{r} \times \mathbf{r}'$ é

$$\|\mathbf{r} \times \mathbf{r}'\| = \|\mathbf{r}\| \|\mathbf{r}'\| |\sin \theta| \quad (5.6)$$

onde θ é o ângulo entre os vetores $\mathbf{r}$ e $\mathbf{r}'$. De fato, se $\mathbf{r}$ e $\mathbf{r}'$ são independentes então a identidade de Lagrange (5.4) e a definição (2.12) implicam

$$\begin{aligned} \|\mathbf{r} \times \mathbf{r}'\|^2 &= \|\mathbf{r}\|^2 \|\mathbf{r}'\|^2 - (\mathbf{r} \cdot \mathbf{r}')^2 \\ &= \|\mathbf{r}\|^2 \|\mathbf{r}'\|^2 (1 - \cos^2 \theta) \end{aligned}$$

Se, por outro lado, $\mathbf{r}$ e $\mathbf{r}'$ são dependentes, a igualdade é trivial.

Portanto, o produto vetorial $\mathbf{r} \times \mathbf{r}'$ é um vetor ortogonal ao plano $\mathbb{R}\mathbf{r} + \mathbb{R}\mathbf{r}'$ cuja norma é igual a área do paralelogramo definido pelos vetores $\mathbf{r}$ e $\mathbf{r}'$, ou seja,

$$\text{Área}(\{\mathbf{r} + s\mathbf{r}' \text{ com } 0 \leq s \leq 1\}) = \|\mathbf{r} \times \mathbf{r}'\|$$

E interessante observar que esta fórmula diz que a área de um paralelogramo no espaço é igual à raiz quadrada

$$\sqrt{\left| \begin{matrix} y & z \\ y' & z' \end{matrix} \right|^2 + \left| \begin{matrix} x & z \\ x' & z' \end{matrix} \right|^2 + \left| \begin{matrix} x & y \\ x' & y' \end{matrix} \right|^2}$$

da soma dos quadrados das áreas das projeções do paralelogramo nos três planos y - z , x - z e x - y . Uma espécie de teorema de Pitágoras para áreas!

ex: Calcule a área do paralelogramo de lados $\vec{OP}$ e $\vec{OQ}$, onde $P = (2, 4, -1)$ e $Q = (1, -3, 1)$.

ex: Calcule a área do triângulo de vértices $(1, 2, 0)$, $(2, 3, 4)$ e $(-1, 0, 0)$.

Orientação. Por outro lado, o sinal do produto vetorial depende da “orientação” escolhida pela base canônica $\mathbf{i}, \mathbf{j}, \mathbf{k}$. Na orientação “destra”, o produto vetorial $\mathbf{r} \times \mathbf{r}'$ tem a orientação do polegar da nossa mão direita se o indicador representa o vetor $\mathbf{r}$ e o dedo médio representa $\mathbf{r}'$ (assim que $\mathbf{r}$ pode ser transformado no vetor $\mathbf{r}'$ por meio de uma rotação de um ângulo agudo em torno de $\mathbf{r} \times \mathbf{r}'$).

Produto vetorial e vetor normal. Se $\mathbf{u}$ e $\mathbf{v}$ são dois vetores linearmente independentes de $\mathbb{R}^3$, então o produto $\mathbf{n} = \mathbf{u} \times \mathbf{v}$ é um vetor não nulo ortogonal ao plano $\mathbb{R}\mathbf{u} + \mathbb{R}\mathbf{v}$. É claro então que os vetores $\mathbf{u}, \mathbf{v}, \mathbf{n}$ são independentes. De fato, se $x\mathbf{u} + y\mathbf{v} + z\mathbf{n} = \mathbf{0}$, e calculamos o produto escalar desta expressão pelo vetor $\mathbf{n}$, obtemos $z\|\mathbf{n}\|^2 = 0$ (pois $\mathbf{n}$ é ortogonal a $\mathbf{u}$ e $\mathbf{v}$), e portanto $z = 0$. Mas a independência dos $\mathbf{u}$ e $\mathbf{v}$ então implica que também $x = y = 0$. Consequentemente, pelo teorema 4.4,

Teorema 5.2. Se $\mathbf{u}$ e $\mathbf{v}$ são linearmente independentes e $\mathbf{n} = \mathbf{u} \times \mathbf{v}$, então o conjunto $\{\mathbf{u}, \mathbf{v}, \mathbf{n}\}$ é uma base de $\mathbb{R}^3$.

Em particular, se $\mathbf{u}$ e $\mathbf{v}$ são vetores unitários e ortogonais, então também $\mathbf{n} = \mathbf{u} \times \mathbf{v}$ é unitário, e o conjunto $\{\mathbf{u}, \mathbf{v}, \mathbf{n}\}$ é uma base ortonormada de $\mathbb{R}^3$.

Uma consequência é que, como a intuição sugere (não cabe mais que uma reta ortogonal a um plano no espaço de dimensão 3)

Teorema 5.3. Todo vetor $\mathbf{w}$ ortogonal aos vetores independentes $\mathbf{u}$ e $\mathbf{v}$ é proporcional ao produto vetorial $\mathbf{n} = \mathbf{u} \times \mathbf{v}$.

Demonstração. Sendo $\mathbf{u}, \mathbf{v}, \mathbf{n}$ uma base de $\mathbb{R}^3$, todo vetor do espaço é $\mathbf{w} = x\mathbf{u} + y\mathbf{v} + z\mathbf{n}$ para alguns coeficientes x, y e z . Se $\mathbf{w}$ é ortogonal a $\mathbf{u}$ e $\mathbf{v}$, então $\|\mathbf{w}\|^2 = z\mathbf{n} \cdot \mathbf{w}$. Por outro lado, como também $\mathbf{n}$ é ortogonal a $\mathbf{u}$ e $\mathbf{v}$, então $\mathbf{n} \cdot \mathbf{w} = z\|\mathbf{n}\|^2$. Consequentemente,

$$\|\mathbf{w}\|^2 \|\mathbf{n}\|^2 = z(\mathbf{n} \cdot \mathbf{w}) \|\mathbf{n}\|^2 = (z\mathbf{n}) \cdot \mathbf{w}$$

e, pela desigualdade de Schwarz 2.1, caso da igualdade, os vetores $\mathbf{w}$ e $\mathbf{n}$ são proporcionais. $\square$

ex: Mostre que se $\mathbf{u}$ e $\mathbf{v}$ são não nulos e ortogonais, e $\mathbf{n} = \mathbf{u} \times \mathbf{v}$, então valem as relações

$$(\mathbf{u} \times \mathbf{v}) \times \mathbf{u} = \mathbf{v} \quad \text{e} \quad \mathbf{v} \times (\mathbf{u} \times \mathbf{v}) = \mathbf{u}$$

(use a fórmula de Lagrange (5.5)).

ex: O plano gerado pelos vetores linearmente independentes $\mathbf{u}$ e $\mathbf{v}$ no espaço $\mathbb{R}^3$ é

$$\mathbb{R}\mathbf{u} + \mathbb{R}\mathbf{v} = (\mathbf{u} \times \mathbf{v})^\perp = \{\mathbf{r} \in \mathbb{R}^3 \text{ t.q. } \mathbf{r} \cdot (\mathbf{u} \times \mathbf{v}) = 0\}$$

ex: Determine um vetor normal aos vetores $(2, 3, -1)$ e $(5, 2, 4)$.

ex: Determine uma base de $\mathbb{R}^3$ contendo os vetores $(0, 1, 1)$ e $(1, 1, 1)$.

ex: Determine um vetor normal ao plano que passa pelos pontos $(0, 1, 0)$, $(1, 1, 0)$ e $(0, 2, 3)$

ex: Determine uma equação cartesiana do plano gerado pelos vetores $(-3, 1, 2)$ e $(1, 5, -2)$.

ex: Determine uma equação cartesiana do plano que passa pelos pontos $(0, 0, 0)$, $(1, 0, 0)$ e $(0, 1, 0)$.

ex: Determine uma equação cartesiana do plano $\{(1+t+s, t-s, 5t) \text{ com } (t, s) \in \mathbb{R}^2\}$.

ex: [Ap69] 13.11.

5.3 Produto misto

Produto misto/triplo escalar e determinante. O *produto misto/triplo* dos vetores $\mathbf{r} = (x, y, z)$, $\mathbf{r}' = (x', y', z')$ e $\mathbf{r}'' = (x'', y'', z'')$ do espaço euclidiano $\mathbb{R}^3$, nesta ordem, é o escalar $\mathbf{r} \cdot (\mathbf{r}' \times \mathbf{r}'')$. A fórmula explícita,

$$\begin{aligned} \mathbf{r} \cdot (\mathbf{r}' \times \mathbf{r}'') &= x(y'z'' - z'y'') - y(x'z'' - z'x'') + z(x'y'' - y'x'') \\ &= xy'z'' - xy''z' + x''yz' - x'y'z'' + x'y''z - x''y'x \end{aligned}$$

não diz grande coisa, e não é fácil de decorar. No entanto, é útil observar que é uma soma de todos os possíveis produtos de coordenadas diferentes dos três vetores, com sinais positivo ou negativo dependendo se a permutação é par ou ímpar, a começar pela permutação inicial $xy'z''$. Mais importante é que a representação do produto vetorial como determinante formal produz uma maneira “visual” de calcular o produto misto. O produto misto é igual ao *determinante* da matriz 3×3 cujas linhas são as componentes dos três vetores, que é definido pela seguinte fórmula, à custa do determinante de uma matriz 2×2 ,

$$\mathbf{r} \cdot (\mathbf{r}' \times \mathbf{r}'') = \text{Det} \begin{pmatrix} x & y & z \\ x' & y' & z' \\ x'' & y'' & z'' \end{pmatrix} := x \begin{vmatrix} y' & z' \\ y'' & z'' \end{vmatrix} - y \begin{vmatrix} x' & z' \\ x'' & z'' \end{vmatrix} + z \begin{vmatrix} x' & y' \\ x'' & y'' \end{vmatrix}$$

Outra notação para o determinante de uma matriz 3×3 é

$$\text{Det} \begin{pmatrix} x & y & z \\ x' & y' & z' \\ x'' & y'' & z'' \end{pmatrix} = \begin{vmatrix} x & y & z \\ x' & y' & z' \\ x'' & y'' & z'' \end{vmatrix}$$

(esta notação pode gerar confusões quando tentamos calcular o módulo do determinante ...)

O determinante é uma função dos três vetores que formam as linhas da matriz. É fácil verificar que o determinante não muda se trocamos as linhas pelas colunas. É também fácil verificar que mudar a ordem das colunas ou das linhas na matriz (ou seja, dos vetores no produto misto) só

pode alterar o sinal do determinante, e o sinal muda quando trocamos as posições de apenas duas linhas ou de duas colunas. Ou seja, o sinal do determinante, pensado como função das três linhas ou das três colunas, depende da “paridade” da permutação destes três vetores. Por exemplo,

$$\begin{vmatrix} x & y & z \\ x' & y' & z' \\ x'' & y'' & z'' \end{vmatrix} = \begin{vmatrix} x & x' & x'' \\ y & y' & y'' \\ z & z' & z'' \end{vmatrix} = - \begin{vmatrix} x' & y' & z' \\ x & y & z \\ x'' & y'' & z'' \end{vmatrix} = \begin{vmatrix} x' & y' & z' \\ x'' & y'' & z'' \\ x & y & z \end{vmatrix} = \dots$$

ou, em termos do produto misto,

$$\mathbf{r} \cdot (\mathbf{r}' \times \mathbf{r}'') = \mathbf{r}' \cdot (\mathbf{r}'' \times \mathbf{r}) = \mathbf{r}'' \cdot (\mathbf{r} \times \mathbf{r}') = -\mathbf{r} \cdot (\mathbf{r}' \times \mathbf{r}'') = -\mathbf{r}' \cdot (\mathbf{r} \times \mathbf{r}'') = \dots$$

Consequentemente, é claro que o produto misto é nulo se pelo menos dois dos três vetores forem iguais, ou proporcionais. É também nulo se um dos vetores pertence ao plano gerado pelos outros dois. De fato,

Teorema 5.4. *Os vetores $\mathbf{a}, \mathbf{b}, \mathbf{c}$ do espaço euclidiano $\mathbb{R}^3$ são independentes (e portanto formam uma base) sse $\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c}) \neq 0$.*

Demonstração. Se $\mathbf{a}, \mathbf{b}, \mathbf{c}$ são dependentes, então existe uma combinação linear não trivial tal que $x\mathbf{a} + y\mathbf{b} + z\mathbf{c} = \mathbf{0}$. A menos de trocar os nomes dos vetores, podemos assumir que $x \neq 0$. Isto quer dizer que $\mathbf{a} = -(y/x)\mathbf{b} - (z/x)\mathbf{c}$, ou seja que $\mathbf{a}$ pertence ao plano gerado por $\mathbf{b}$ e $\mathbf{c}$. Então

$$\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c}) = -(y/x)\mathbf{b} \cdot (\mathbf{b} \times \mathbf{c}) - (z/x)\mathbf{c} \cdot (\mathbf{b} \times \mathbf{c}) = 0$$

pois $\mathbf{b}$ e $\mathbf{c}$ são ortogonais a $\mathbf{b} \times \mathbf{c}$.

Vice-versa, assumimos que $\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c}) = 0$. Se $\mathbf{b} \times \mathbf{c} = \mathbf{0}$, então os vetores $\mathbf{b}$ e $\mathbf{c}$ são proporcionais, e a fortiori também os vetores $\mathbf{a}, \mathbf{b}, \mathbf{c}$ são dependentes. Se, por outro lado, $\mathbf{n} = \mathbf{b} \times \mathbf{c} \neq \mathbf{0}$, então os vetores $\mathbf{b}, \mathbf{c}, \mathbf{n}$ formam uma base de $\mathbb{R}^3$, pelo teorema 5.2. Isto implica que $\mathbf{a} = x\mathbf{n} + y\mathbf{b} + z\mathbf{c}$, e consequentemente

$$0 = \mathbf{a} \cdot (\mathbf{b} \times \mathbf{c}) = (x\mathbf{n} + y\mathbf{b} + z\mathbf{c}) \cdot (\mathbf{b} \times \mathbf{c}) = x\|\mathbf{n}\|^2$$

porque $\mathbf{b} \times \mathbf{c}$ é ortogonal a $\mathbf{b}$ e $\mathbf{c}$. Como $\mathbf{n} \neq \mathbf{0}$, isto implica que $x = 0$, e portanto que $\mathbf{a}$ é uma combinação linear de $\mathbf{b}$ e $\mathbf{c}$, ou seja, que $\mathbf{a}, \mathbf{b}, \mathbf{c}$ são dependentes. $\square$

ex: Calcule o produto misto $\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})$ quando $\mathbf{a} = (1, 1, 0)$, $\mathbf{b} = (1, 3, 1)$ e $\mathbf{c} = (0, 1, 1)$.

ex: Calcule o produto misto $\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})$ quando $\mathbf{a} = (-2, 5, 1)$, $\mathbf{b} = (0, 3, 0)$ e $\mathbf{c} = (6, 7, -3)$.

ex: Diga se os vetores $(7, 2, 3)$, $(-1, -5, 3)$ e $(0, 1, -3)$ são independentes.

ex: Diga se os vetores $(1, 0, 1)$, $(0, 1, 0)$ e $(1, 1, 1)$ são independentes.

ex: Determine os valores de t para os quais os vetores $(t, 1, 0)$, $(1, t, 1)$ e $(0, 1, t)$ são independentes.

ex: [Ap69] 13.14.

Determinantes e volumes no espaço. A fórmula (5.6) e a definição (2.12) implicam que o módulo do produto misto é

$$|\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})| = \|\mathbf{a}\| \|\mathbf{b}\| \|\mathbf{c}\| |\sin(\theta) \cos(\phi)|$$

onde θ é o ângulo entre $\mathbf{b}$ e $\mathbf{c}$, e ϕ é o ângulo entre $\mathbf{a}$ e $\mathbf{b} \times \mathbf{c}$. Por outro lado, o sinal do produto misto depende da orientação dos vetores, e é positivo se $\mathbf{a}, \mathbf{b}, \mathbf{c}$ são orientados como os vetores $\mathbf{i}, \mathbf{j}, \mathbf{k}$ da base canônica, pois $\mathbf{i} \cdot (\mathbf{j} \times \mathbf{k}) = \|\mathbf{i}\|^2 = 1$.

O *paralelepípedo* gerado/definido pelos vetores $\mathbf{a}$, $\mathbf{b}$ e $\mathbf{c}$ de $\mathbb{R}^3$ é o conjunto

$$P = \{t\mathbf{a} + s\mathbf{b} + u\mathbf{c} \text{ com } 0 \leq t, s, u \leq 1\}.$$

Considerações elementares dizem que seu volume deve ser igual produto da área da base, por exemplo o paralelogramo gerado por $\mathbf{b}$ e $\mathbf{c}$, que é igual a $\|\mathbf{b}\|\|\mathbf{c}\|\sin\theta$, vezes a norma da projeção de $\mathbf{a}$ sobre o vetor unitário normal ao plano gerado por $\mathbf{b}$ e $\mathbf{c}$, que é igual a $\|\mathbf{a}\|\cos\phi$. Consequentemente, o volume do paralelepípedo é

$$\text{Volume}(\{t\mathbf{a} + s\mathbf{b} + u\mathbf{c} \text{ com } 0 \leq t, s, u \leq 1\}) = |\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})| = \left| \text{Det} \begin{pmatrix} a_1 & a_2 & a_3 \\ b_1 & b_2 & b_3 \\ c_1 & c_2 & c_3 \end{pmatrix} \right|$$

ex: Calcule o volume do paralelepípedo definido pelos vetores $\mathbf{i} + \mathbf{j}$, $\mathbf{j} + \mathbf{k}$ e $\mathbf{k} + \mathbf{i}$.

ex: Calcule o volume do paralelepípedo gerado pelos vetores $(3, 3, 1)$, $(2, 1, 2)$ e $(5, 1, 1)$.

ex: Calcule o volume do paralelepípedo gerado pelos vetores $(0, 0, 1)$, $(5, 7, -3)$ e $(-9, 0, 0)$.

Regra de Cramer. Resolver o sistema de três equações lineares

$$\begin{cases} a_1x + b_1y + c_1z = d_1 \\ a_2x + b_2y + c_2z = d_2 \\ a_3x + b_3y + c_3z = d_3 \end{cases}$$

nas três incógnitas x , y e z , significa representar o vetor $\mathbf{d} = (d_1, d_2, d_3)$ como combinação linear

$$x\mathbf{a} + y\mathbf{b} + z\mathbf{c} = \mathbf{d}$$

dos vetores $\mathbf{a} = (a_1, a_2, a_3)$, $\mathbf{b} = (b_1, b_2, b_3)$ e $\mathbf{c} = (c_1, c_2, c_3)$ com coeficientes x , y e z .

É claro que o sistema admite sempre uma solução, para qualquer vetor $\mathbf{d}$, sse os vetores $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$ são independentes, e portanto formam uma base de $\mathbb{R}^3$. Pelo teorema 5.4, isto acontece sse o produto misto $\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})$ é diferente de zero. Menos evidente é que o produto misto permite determinar uma fórmula para a solução, em termos dos vetores $\mathbf{a}$, $\mathbf{b}$, $\mathbf{c}$ e $\mathbf{d}$. De fato, se calculamos o produto escalar desta representação de $\mathbf{d}$ pelos vetores $\mathbf{b} \times \mathbf{c}$, depois $\mathbf{c} \times \mathbf{a}$ e finalmente $\mathbf{a} \times \mathbf{b}$, obtemos as três identidades

$$\begin{aligned} x\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c}) &= \mathbf{d} \cdot (\mathbf{b} \times \mathbf{c}) \\ y\mathbf{b} \cdot (\mathbf{c} \times \mathbf{a}) &= \mathbf{d} \cdot (\mathbf{c} \times \mathbf{a}) \\ z\mathbf{c} \cdot (\mathbf{a} \times \mathbf{b}) &= \mathbf{d} \cdot (\mathbf{a} \times \mathbf{b}) \end{aligned}$$

(observe que este truque generaliza o que foi utilizado para provar que dois vetores do plano são independentes sse o determinante da matriz 2×2 formada pelas coordenadas é diferente de zero, mas que agora estamos a eliminar duas variáveis de cada vez!) O coeficiente de x , y e z , é sempre o mesmo. Portanto, se o produto misto $\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})$ é diferente de zero (ou seja, se o determinante da matriz 3×3 cujas colunas são os vetores $\mathbf{a}$, $\mathbf{b}$ e $\mathbf{c}$ é diferente de zero, ou seja, se os três vetores são independentes), então o sistema admite uma única solução (x, y, z) , dada pela *regra de Cramer* seguinte:

$$x = \frac{\mathbf{d} \cdot (\mathbf{b} \times \mathbf{c})}{\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})} \quad y = \frac{\mathbf{c} \cdot (\mathbf{d} \times \mathbf{a})}{\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})} \quad z = \frac{\mathbf{a} \cdot (\mathbf{b} \times \mathbf{d})}{\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})}$$

Observe que o denominador é o determinante da matriz 3×3 cujas colunas são os vetores $\mathbf{a}$, $\mathbf{b}$ e $\mathbf{c}$, e o numerador da i -ésima coordenada é obtido ao substituir, nesta matriz, a i -ésima coluna pelo vetor $\mathbf{d}$.

ex: Resolva os sistemas

$$\begin{cases} x + y + z = 6 \\ 2x + y + 3z = 9 \\ 3x - y - 5z = -17 \end{cases} \quad \begin{cases} x + 2y - 3z = 4 \\ 2x - 3y + z = 1 \\ 3x - y - 2z = 5 \end{cases}$$

Força magnética. A *força de Lorentz* que experimenta uma partícula com carga eléctrica q e velocidade $\mathbf{v}$ num campo eléctrico $\mathbf{E}$ e magnético $\mathbf{B}$ é (nas unidades do S.I.)

$$\mathbf{F} = q(\mathbf{E} + \mathbf{v} \times \mathbf{B})$$

ex: Mostre que num referencial inercial em que o campo eléctrico é nulo, i.e. $\mathbf{E} = \mathbf{0}$, e portanto a única força é força magnética $q\mathbf{v} \times \mathbf{B}$, a energia cinética $K := \frac{1}{2}m\|\mathbf{v}\|^2$ é conservada, calculando a derivada

$$\frac{d}{dt} \left(\frac{1}{2}m\|\mathbf{v}\|^2 \right) = m\dot{\mathbf{v}} \cdot \mathbf{v}$$

e utilizando a equação de Newton $m\dot{\mathbf{v}} = \mathbf{F}$.

Momento angular e torque. O *momento angular* (relativo à origem do referencial) de uma partícula de massa $m > 0$ colocada na posição $\mathbf{r} \in \mathbb{R}^3$ com momento linear $\mathbf{p} = m\dot{\mathbf{r}}$, é o produto vetorial

$$\mathbf{L} := \mathbf{r} \times \mathbf{p}$$

A derivada do momento angular de uma partícula sujeita à lei de Newton $\dot{\mathbf{p}} = \mathbf{F}$ é igual ao *binário* (ou *torque*) $\mathbf{T} := \mathbf{r} \times \mathbf{F}$, pois

$$\begin{aligned} \dot{\mathbf{L}} &= \dot{\mathbf{r}} \times \mathbf{p} + \mathbf{r} \times \dot{\mathbf{p}} \\ &= \mathbf{r} \times \mathbf{F}. \end{aligned}$$

sendo $\dot{\mathbf{r}} \times \mathbf{p} = \mathbf{0}$.

ex: O momento angular do sistema de n partículas de massas m_i , colocadas nas posições $\mathbf{r}_i \in \mathbb{R}^3$ com momentos lineares $\mathbf{p}_i := m_i\dot{\mathbf{r}}_i$, com $i = 1, 2, \dots, n$, é

$$\mathbf{L} := \sum_{i=1}^n \mathbf{r}_i \times \mathbf{p}_i$$

Sejam $\mathbf{R} := \frac{1}{M} \sum_{i=1}^n m_i \mathbf{r}_i$ o centro de massa do sistema e $\mathbf{P} := M\dot{\mathbf{R}} = \sum_{i=1}^n m_i \dot{\mathbf{r}}_i$ o momento linear do centro de massa, onde $M := \sum_{i=1}^n m_i$ denota a massa total. Mostre que

$$\mathbf{L} = \mathbf{R} \times \mathbf{P} + \mathbf{L}'$$

onde $\mathbf{L}' := \sum_{i=1}^n \mathbf{r}'_i \times \mathbf{p}'_i$, com $\mathbf{r} = \mathbf{R} + \mathbf{r}'$, é o momento angular relativo ao centro de massa.

Angular momentum of Kepler orbits and Hamilton's theorem. If two point masses with positions $\mathbf{r}_1$ and $\mathbf{r}_2$ move around their center of masses according to Newton's gravitational law, then the difference vector $\mathbf{r} = \mathbf{r}_1 - \mathbf{r}_2$ obeys the law

$$\ddot{\mathbf{r}} = -\frac{\mathbf{r}}{\|\mathbf{r}\|^3} \quad (5.7)$$

(if we set to one both the reduced mass of the system and the universal gravitational constant). If $\mathbf{v} = \dot{\mathbf{r}}$ denotes the velocity vector, then a computation shows that the angular momentum

$$\mathbf{L} := \mathbf{r} \times \mathbf{v}$$

is a constant of the motion. If at some (initial) time the vectors $\mathbf{r}$ and $\mathbf{v}$ are not parallel, then $\mathbf{L} \neq \mathbf{0}$. We may therefore choose a reference Cartesian system in which $\mathbf{L} = \ell \mathbf{k}$ for some $\ell > 0$, and write the position vector as $\mathbf{r}(t) = r \cos(\theta) \mathbf{i} + r \sin(\theta) \mathbf{j}$ for some time-dependent angle θ and length $r = \|\mathbf{r}\|$. A computation shows that

$$\ell = r^2 \dot{\theta}, \quad (5.8)$$

expressing *Kepler second law*¹², according to which “the position vector from the sun to a planet sweeps out equal areas in equal times”. If we eliminate dt from Newton equation (5.7) (written for $\dot{\mathbf{v}}$) and equation (5.8), we get

$$d\mathbf{v}/d\theta = -\ell^{-1}(\cos(\theta)\mathbf{i} + \sin(\theta)\mathbf{j})$$

and, after integration,

$$\mathbf{v} = \mathbf{v}_0 + \ell^{-1}(\sin(\theta)\mathbf{i} - \cos(\theta)\mathbf{j}),$$

for some constant vector $\mathbf{v}_0$ (on the $z = 0$ plane). Thus, “the velocity vector moves along a circle orthogonal to the angular velocity vector with radius inversely proportional to its length”^{13 14}.

Nabla, gradiente, divergência, rotacional e laplaciano. Os produtos escalares e vetorial são também usados para definir certos operadores diferenciais importantes da física-matemática. O operador *nabla* é o operador diferencial formal

$$\nabla = \mathbf{i} \frac{\partial}{\partial x} + \mathbf{j} \frac{\partial}{\partial y} + \mathbf{k} \frac{\partial}{\partial z}$$

que pode agir sobre campos escalares ou vetoriais definidos numa região do espaço euclidiano $\mathbb{R}^3$. Por exemplo, o *gradiente* do campo escalar $f(x, y, z)$ é o produto formal

$$\text{grad}(f) := \nabla f = \frac{\partial f}{\partial x} \mathbf{i} + \frac{\partial f}{\partial y} \mathbf{j} + \frac{\partial f}{\partial z} \mathbf{k}$$

A *divergência* do campo vetorial $\mathbf{F}(x, y, z) = (A(x, y, z), B(x, y, z), C(x, y, z))$ é o produto escalar formal

$$\text{div}(\mathbf{F}) := \nabla \cdot \mathbf{F} = \frac{\partial A}{\partial x} + \frac{\partial B}{\partial y} + \frac{\partial C}{\partial z}$$

O *rotacional* do campo vetorial $\mathbf{F} = (A, B, C)$ é o produto vetorial formal

$$\text{curl}(\mathbf{F}) := \nabla \times \mathbf{F} = \left(\frac{\partial C}{\partial y} - \frac{\partial B}{\partial z} \right) \mathbf{i} + \left(\frac{\partial A}{\partial z} - \frac{\partial C}{\partial x} \right) \mathbf{j} + \left(\frac{\partial B}{\partial x} - \frac{\partial A}{\partial y} \right) \mathbf{k}$$

A *divergência do gradiente* do campo escalar f é

$$\Delta f := \nabla \cdot \nabla f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} + \frac{\partial^2 f}{\partial z^2}$$

O operador $\Delta = \nabla \cdot \nabla$, ou também ∇^2 , é chamado *laplaciano*, ou *operador de Laplace*, e é um dos operadores mais importante da física matemática, pois aparece nas equações da corda vibrante, das ondas, do calor, de Poisson, de Schrödinger, ...

ex: Verifique qua a fórmula de Lagrange (5.5) implica

$$\nabla \times (\nabla \times \mathbf{F}) = \nabla (\nabla \cdot \mathbf{F}) - \Delta \mathbf{F}$$

(onde o laplaciano do campo vetorial é definido coordenada por coordenada).

¹²J. Kepler, *Astronomia Nova*, Prague, 1609.

¹³W.R. Hamilton, The hodograph or a new method of expressing in symbolic language the Newtonian law of attraction, *Proc. Roy. Irish Acad.* **3** (1846), 344-353.

¹⁴J. Milnor, On the geometry of the Kepler problem. *Amer. Math. Monthly* **90** (1983), 353-365.

6 Números complexos

ref: [Ap69] Vol. 1, 9.1-10

6.1 O corpo dos números complexos

Os números complexos formam um corpo que estende o corpo dos números reais, e cujas operações refletem a geometria euclidiana, ou melhor, conforme, do plano.

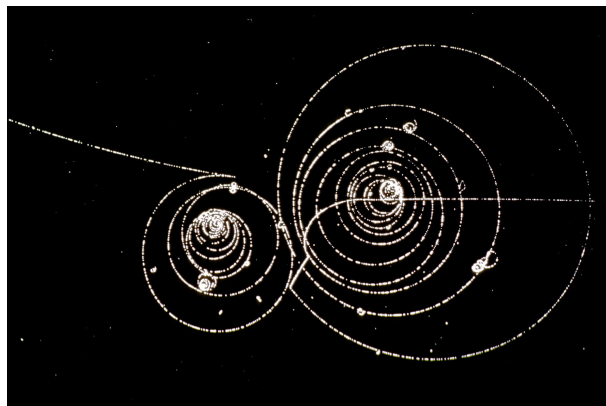
História muito breve. Os números complexos foram inventados/descobertos no século XVI como um truque “sofístico” para resolver polinômios do gênero $x^3 + px + q = 0$. Hoje em dia, fazem parte da formulação das leis fundamentais da Natureza, como, por exemplo, a *equação de Schrödinger*

$$i\hbar \frac{\partial \psi}{\partial t} = -\frac{\hbar^2}{2m} \Delta \psi + V \psi$$

da mecânica quântica, ou os *integrals de Feynman*

$$\int_{\text{paths}} e^{iS[x]/\hbar} \mathcal{D}x$$

da teoria quântica dos campos.



Nas palavras de Roger Penrose [Pe05],

... complex numbers, as much as reals, and perhaps even more, find a unity with nature that is truly remarkable. It is as though Nature herself is as impressed by the scope and consistency of the complex-number system as we are ourselves, and has entrusted to these numbers the precise operations of her world at its minutest scales.

O corpo dos números complexos. Do ponto de vista algébrico/abstrato (para um matemático que sabe o que é uma “extensão” de um corpo), o corpo dos *números complexos* é $\mathbb{C} := \mathbb{R}(i)$, onde $i^2 = -1$.

Mais compreensível (para um matemático que sabe o que é um “anel”) é dizer que é o quociente

$$\mathbb{C} := \mathbb{R}[X] / \langle X^2 + 1 \rangle$$

do anel $\mathbb{R}[X]$ dos polinômios na incógnita X com coeficientes reais módulo o ideal gerado pelo polinômio $X^2 + 1$. Um polinômio arbitrário na incógnita X é uma expressão formal do gênero $p(X) = a_n X^n + \dots + a_1 X + a_0$, com coeficientes reais $a_k \in \mathbb{R}$. Somas e produtos entre polinômios são definidos da forma natural. No quociente, dois polinômios $p(X)$ e $q(X)$ são identificados se diferem por um polinômio da forma $f(X)(X^2 + 1)$, onde $f(X)$ é um polinômio arbitrário. Isto significa que

podemos simplificar e substituir cada fator $X^2 + 1$ por 0, ou seja, cada segunda potência X^2 por -1 . Mas então também podemos substituir $X^3 = XX^2 = -X$, depois $X^4 = XX^3 = -X^2 = 1$, É claro portanto que todo polinômio pode ser identificado, no quociente, com um polinômio de grau apenas um, do gênero $a + bX$, com $a, b \in \mathbb{R}$. O número complexo que tradicionalmente chamamos i , cujo quadrado satisfaz $i^2 + 1 = 0$, é precisamente a imagem de X no quociente.

Na prática, para seres humanos, $\mathbb{C}$ é o conjunto das expressões formais

$$z = x + iy$$

com $x, y \in \mathbb{R}$, que chamamos (e pensamos) “números complexos”, munido das operações binárias “soma” e “multiplicação” definidas usando as usuais regras algébrica dos polinômios com coeficientes reais na incógnita “ i ” e no fim usando a substituição $i^2 = -1$. O resultado é que a soma é definida por

$$(x_1 + iy_1) + (x_2 + iy_2) = (x_1 + x_2) + i(y_1 + y_2) \quad (6.1)$$

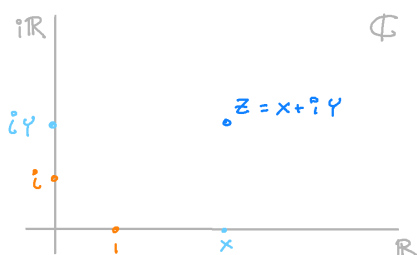
e a multiplicação é definida por

$$\begin{aligned} (x_1 + iy_1) \cdot (x_2 + iy_2) &= x_1x_2 + ix_1y_2 + iy_1x_2 + i^2y_1y_2 \\ &= (x_1x_2 - y_1y_2) + i(x_1y_2 + y_1x_2). \end{aligned} \quad (6.2)$$

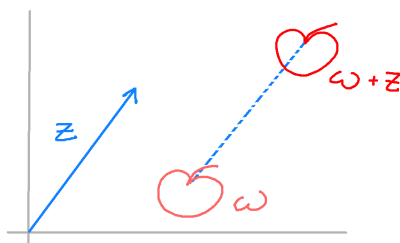
É subentendido que dois números complexos $x_1 + iy_1$ e $x_2 + iy_2$ são (considerados) iguais sse $x_1 = x_2$ e $y_1 = y_2$. É também conveniente denotar simplesmente $x + i0 = x$ e $0 + iy = iy$. Em particular, $x \mapsto x + i0$ define uma inclusão $\mathbb{R} \subset \mathbb{C}$, e as operações definidas acima são as usuais operações no corpo dos reais.

Se $i := 0 + i \cdot 1 \in \mathbb{C}$, então $i \cdot i = -1$, ou seja, $\pm i$ são (as únicas) “raízes quadradas de -1 ”. De fato (e esta é a origem das fórmulas acima, que portanto não devem ser decoradas), somas e multiplicações entre números complexos podem ser manipuladas como as correspondentes operações entre números reais (ou seja, usando as propriedades associativas, comutativas e distributivas), e depois substituindo $i \cdot i$ por -1 .

É natural identificar os números complexos $z = x + iy$ com os pontos/vetores (x, y) do plano $\mathbb{R}^2$, e denotar a correspondência com $x + iy \approx (x, y)$. A reta real $\mathbb{R} \subset \mathbb{C}$ é naturalmente identificada com o eixo dos x 's em $\mathbb{R}^2$, e a reta “imaginária” $i\mathbb{R} \subset \mathbb{C}$ é naturalmente identificada com o eixo dos y 's.



Então a soma $z_1 + z_2$ corresponde à soma dos vetores $z_1 \approx (x_1, y_1)$ e $z_2 \approx (x_2, y_2)$ do plano. O conjunto $\mathbb{C}$, munido da operação $+$ definida em (6.1), é um grupo abeliano aditivo, cujo elemento neutro é $0 := 0 + i0$. O oposto do número complexo $z = x + iy$ é o número complexo $-z = (-x) + i(-y)$ (denotado simplesmente por $-z = -x - iy$), que verifica $z + (-z) = 0$. Somar um número complexo z , i.e. fazer $w \mapsto w + z$, corresponde a fazer uma translação no plano.



Todo $z = x + iy \neq 0$ admite um único inverso multiplicativo, um número complexo $1/z$ tal que $z \cdot (1/z) = (1/z) \cdot z = 1$, dado por

$$\frac{1}{z} = \frac{x}{x^2 + y^2} - i \frac{y}{x^2 + y^2} \quad (6.3)$$

como é fácil verificar (observe que $z \neq 0$ sse $x^2 + y^2 > 0$). Portanto, o conjunto $\mathbb{C}^\times := \mathbb{C} \setminus \{0\}$, munido da operação $\cdot$ definida em (6.2), é um grupo abeliano, o grupo multiplicativo dos números complexos invertíveis, cujo elemento neutro é $1 := 1 + i0$.

As potências inteiras de um número complexo são definidas por recorrência: $z^{n+1} := z \cdot z^n$, se $n \geq 1$, sendo $z^0 := 1$. Se $z \in \mathbb{C}^\times$, então as potências negativas são definidas por $z^{-n} := (1/z)^n$. Por exemplo, $i^2 = -1$, e $i^{-1} = 1/i = -i$.

A propriedade distributiva $(z_1 + z_2) \cdot z_3 = z_1 z_3 + z_2 z_3$, que implicitamente foi usada na definição da multiplicação e que portanto é válida para todos triplos de números complexos, mostra finalmente que $\mathbb{C}$ é um corpo. Em particular, um produto de dois números é nulo, ou seja, $zw = 0$, sse pelo menos um dos fatores é nulo, ou seja, $z = 0$ ou $w = 0$.

O corpo dos números complexos contém, como subcorpos, o corpo dos reais $\mathbb{R}$, que por sua vez contém o corpo $\mathbb{Q}$ dos racionais. No entanto, não é possível estender a ordem de $\mathbb{R}$ a uma ordem de $\mathbb{C}$ que seja compatível com as operações: o corpo dos números complexos não é um corpo ordenado.

ex: Calcule

$$(2 + i3) + (3 - i2) \quad (1 - i) \cdot (2 - i) \quad (1 + i) + (1 - i) \cdot (2 - i5)$$

ex: Na identificação $\mathbb{C} \approx \mathbb{R}^2$ definida por $x + iy \approx (x, y)$, o produto $(a + ib)(x + iy)$ é dado por

$$\begin{pmatrix} a & -b \\ b & a \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix}$$

Em particular, se $a + ib \neq 0$, então o produto é

$$\sqrt{a^2 + b^2} \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix}$$

onde θ é um ângulo tal que $\cos \theta = a/\sqrt{a^2 + b^2}$ e $\sin \theta = b/\sqrt{a^2 + b^2}$. Esta fórmula revela o significado geométrico da multiplicação: a multiplicação por um número complexo $a + ib$ diferente de zero corresponde a uma rotação de um ângulo θ e uma homotetia de razão $\sqrt{a^2 + b^2}$.

ex: Represente na forma $x + iy$ os seguintes números complexos

$$i^3 \quad \frac{1}{1+i} \quad \frac{2-i}{1+i} \quad \frac{1-i}{1+i} \cdot \frac{i}{2+i} \quad (1-i3)^2 \quad i^{17} \quad (2 \pm i)^3$$

ex: Verifique que, dado $z = x + iy$, a única solução de $(x + iy)(a + ib) = 1$ é dada pela fórmula (6.3), ou seja, $a = x/(x^2 + y^2)$ e $b = -y/(x^2 + y^2)$ (por exemplo, resolvendo um sistema de 2 equações lineares nas incógnitas a e b).

ex: Ao acrescentar a raiz quadrada de -1 aos números reais acontece um primeiro milagre: todo número complexo (por exemplo, real) admite umas raízes quadradas. Verifique que o quadrado de

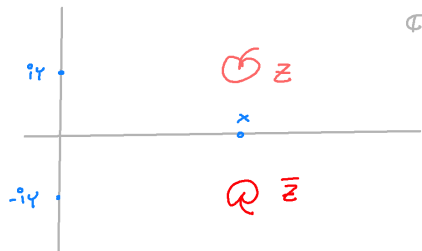
$$\sqrt{\frac{1}{2} \left(x + \sqrt{x^2 + y^2} \right)} + i \sqrt{\frac{1}{2} \left(-x + \sqrt{x^2 + y^2} \right)}$$

e do seu oposto é igual a $z = x + iy$.

Conjugação. Os corpo números complexos também admitem uma involução (uma transformação que é a própria inversa) que respeita as operações algébricas. O *conjugado* de $z = x + iy$ é

$$\bar{z} := x - iy,$$

ou seja, a imagem do ponto $x + iy \approx (x, y)$ pela reflexão na reta $y = 0$ do plano $\mathbb{R}^2 \approx \mathbb{C}$. Em particular, i e o seu conjugado $\bar{i} = -i$ são as duas raízes de -1 .



A conjugação é um “automorfismo” de $\mathbb{C}$, ou seja, respeita soma e produtos:

$$\overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2 \quad \text{e} \quad \overline{z_1 \cdot z_2} = \bar{z}_1 \cdot \bar{z}_2 \quad (6.4)$$

(a segunda identidade não é óbvia, mas um “milagre” que relaciona multiplicação e geometria euclidiana do plano). Observe também que a conjugação é uma involução, ou seja, $\bar{\bar{z}} = z$.

Os números reais

$$x = \Re(z) := \frac{z + \bar{z}}{2} \quad \text{e} \quad y = \Im(z) := \frac{z - \bar{z}}{2i}$$

são ditos *parte real* e *parte! imaginária* do número complexo $z = x + iy$, respetivamente. Observe que $z = \bar{z}$ sse z é real, i.e. sse $\Im(z) = 0$.

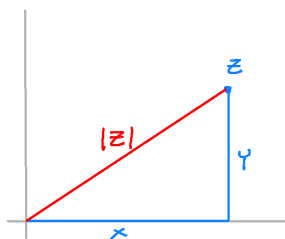
Norma e módulo. A conjugação permite definir a *norma* (no sentido da teoria de números) do número complexo $z = x + iy$ como

$$N(z) := z\bar{z} = x^2 + y^2 \quad (6.5)$$

que, sendo uma soma de quadrados de números reais, é um número real não-negativo. O *módulo*, ou *valor absoluto*, de $z = x + iy$ é a raiz quadrada de $N(z)$, ou seja,

$$|z| := \sqrt{z\bar{z}} = \sqrt{x^2 + y^2} \quad (6.6)$$

De acordo com o teorema de Pitágoras, é igual à norma euclidiana do vetor $(x, y) \in \mathbb{R}^2$. Em particular, $|z| = 0$ sse $z = 0$.



É claro que $|z| = |\bar{z}|$. Menos evidente, consequência da segunda das (6.4), é que o valor absoluto é multiplicativo, ou seja,

$$|zw| = |z| |w| \quad (6.7)$$

Consequentemente, $|1/z| = 1/|z|$ se $z \neq 0$, e, mais em geral, $|z/w| = |z|/|w|$ se $w \neq 0$. O inverso multiplicativo de um número complexo $z \neq 0$ é então

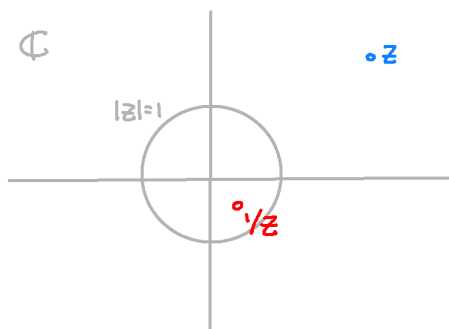
$$1/z = \bar{z}/|z|^2.$$

Os números complexos “unitários”, ou seja, de módulo igual a um, definem a *circunferência unitária*

$$\mathbb{S} := \{z \in \mathbb{C}; |z| = 1\} \subset \mathbb{C}$$

Pela (6.7), o produto de dois números complexos unitários é ainda unitário. Assim, $\mathbb{S}$ é um subgrupo do grupo multiplicativo $\mathbb{C}^\times$, isomorfo ao grupo $U(1)$ das transformações unitárias do espaço euclidiano complexo $\mathbb{C}$.

Portanto, a inversão $z \mapsto 1/z$ de um número complexo não nulo corresponde a uma “reflexão” na circunferência unitária, a transformação $z \mapsto z^* := z/|z|^2$, seguida por uma conjugação $z^* \mapsto \overline{z^*}$, uma reflexão na reta real.



ex: Verifique a seguinte identidade entre números reais

$$(a^2 + b^2)(c^2 + d^2) = (ac - bd)^2 + (ad + bc)^2$$

e deduza a (6.7). Outra possibilidade, menos misteriosa, é partir pelas fatorizações

$$a^2 + b^2 = (a + ib)(a - ib) \quad c^2 + d^2 = (c + id)(c - id)$$

multiplicar as duas, rearranjar os fatores à direita, multiplicar os primeiros dois e os últimos dois

$$\begin{aligned} (a^2 + b^2)(c^2 + d^2) &= (a + ib)(c + id)(a - ib)(c - id) \\ &= ((ac - bd) + i(ad + bc)) ((ac - bd) - i(ad + bc)) = \dots \end{aligned}$$

até observar que ...

Inteiros gaussianos, somas de quadrados e ternos pitagóricos. Os *inteiros gaussianos* são os números complexos da forma $z = a + ib$ com a e b inteiros. Formam o subconjunto $\mathbb{Z} + i\mathbb{Z} \subset \mathbb{C}$, identificado ao reticulado $\mathbb{Z}^2 \subset \mathbb{R}^2$ dos pontos do plano com coordenadas inteiras. É evidente, pelas fórmulas (6.1) e (6.2), que somas e produtos entre inteiros gaussianos também são inteiros gaussianos, ou seja, que os inteiros gaussianos formam um anel. É também evidente, pela (6.5), que a norma, ou seja, o quadrado do módulo de um inteiro gaussiano é um número inteiro, soma de dois quadrados.

Sejam $z = a + ib$ e $w = c + id$ dois inteiros gaussianos. Se calculamos o quadrado em (6.7), e usamos a própria definição (6.6), obtemos

$$(a^2 + b^2)(c^2 + d^2) = (ac - bd)^2 + (ad + bc)^2 \quad (6.8)$$

Esta é conhecida como *identidade de Diofanto*, ou de *Brahmagupta-Fibonacci*, que diz que “um produto de duas somas de dois quadrados é também uma soma de dois quadrados”.

Em particular, quando $z = w = a + ib$, temos (lendo a fórmula acima de direita para esquerda)

$$(a^2 - b^2)^2 + (2ab)^2 = (a^2 + b^2)^2$$

Esta é conhecida como *fórmula de Euclides*¹⁵ e, ao variar os inteiros $a > b \geq 1$, produz uma infinidade de “ternos pitagóricos” $m^2 + n^2 = p^2$ (e, de fato, todos os ternos pitagóricos “primitivos”, ou seja, os ternos (m, n, p) sem fatores comuns!).

Raízes de polinômios reais de grau dois. Um polinômio de grau dois com coeficiente reais é uma função do gênero $f(x) = ax^2 + bx + c$, com $a, b, c \in \mathbb{R}$ e $a \neq 0$. Para calcular as suas raízes, ou seja, os pontos onde $f(x) = 0$, podemos dividir por a , e considerar o polinômio mônico (ou seja, tal que o termo de grau maior tem coeficiente unitário)

$$f(z) = z^2 + 2\alpha z + \beta$$

com $2\alpha = b/a$ e $\beta = c/a$ (coloquei o fator 2 para simplificar os cálculos seguintes). Ao “completar o quadrado”, observamos que

$$\begin{aligned} z^2 + 2\alpha z + \beta &= z^2 + 2\alpha z + \alpha^2 - \alpha^2 + \beta \\ &= (z + \alpha)^2 + (\beta - \alpha^2) \end{aligned}$$

e portanto as raízes são soluções de

$$(z + \alpha)^2 = \alpha^2 - \beta.$$

O número $\delta := \alpha^2 - \beta$ é chamado *discriminante* do polinômio. Se $\delta \geq 0$, temos duas raízes reais $z_{\pm} = -\alpha \pm \sqrt{\delta}$, eventualmente coincidentes quando $\delta = 0$. Em termos dos coeficientes originais a, b, c , esta é a famosa *fórmula resolvente*

$$z_{\pm} = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$$

Se $\delta < 0$, e portanto $\delta = -\omega^2$ para algum $\omega > 0$, o polinômio não admite raízes reais. No entanto, podemos observar que $(\pm i\omega)^2 = \delta$. Então temos duas raízes complexas e conjugadas

$$z_{\pm} = -\alpha \pm i\omega.$$

Nos dois casos, o polinômio mônico fatoriza como produto

$$f(z) = (z - z_+)(z - z_-)$$

de duas raízes, iguais se $\delta = 0$, e simétricas em relação ao eixo real se $\delta < 0$.

ex: Resolva as seguintes equações

$$z^2 - 2z + 2 = 0 \quad z^2 + z + 1 = 0$$

6.2 Representação polar e raízes

A fórmula de Euler revela a relação entre o produto entre números complexo e as homotetias e as rotações do plano.

Representação polar. A *representação polar* do número complexo $z = x + iy \approx (x, y) \in \mathbb{R}^2$ é

$$z = \rho e^{i\theta}$$

onde $\rho = |z| = \sqrt{x^2 + y^2} \geq 0$ é o módulo de z , $\theta = \arg(z) \in \mathbb{R}/2\pi\mathbb{Z}$ é um *argumento* de z , ou seja, um ângulo tal que $x = \rho \cos(\theta)$ e $y = \rho \sin(\theta)$ (logo definido a menos de múltiplos inteiros de 2π), e o número complexo unitário $e^{i\theta} \in \mathbb{S}$ é (provisoriamente) definido pela *fórmula de Euler*

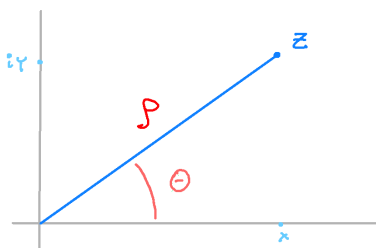
$$e^{i\theta} := \cos(\theta) + i \sin(\theta). \quad (6.9)$$

¹⁵Euclides, *Elementos*, Livro X, Proposição XXIX.

à custa, portanto, das funções trigonométricas $\cos$ e $\sin$, supostas definidas anteriormente (e.g. num curso de cálculo). Observe que

$$\cos(\theta) = \frac{e^{i\theta} + e^{-i\theta}}{2} \quad \text{e} \quad \sin(\theta) = \frac{e^{i\theta} - e^{-i\theta}}{2i}$$

Pode ser útil escolher um valor do argumento, e chamar *argumento principal* de um número z o único argumento que satisfaz $\text{Arg}(z) \in (-\pi, \pi]$.



Produto em representação polar. Se $z_1 = \rho_1 e^{i\theta_1}$ e $z_2 = \rho_2 e^{i\theta_2}$, então as fórmulas de adição para seno e cosseno mostram que

$$z_1 z_2 = \rho_1 \rho_2 e^{i(\theta_1 + \theta_2)} \quad (6.10)$$

Por exemplo, o quadrado de um número complexo $z = \rho e^{i\theta}$ é $z^2 = \rho^2 e^{i2\theta}$. Também, se $z_2 \neq 0$,

$$\frac{z_1}{z_2} = \frac{\rho_1}{\rho_2} e^{i(\theta_1 - \theta_2)}.$$

Estas fórmulas mostram que o grupo multiplicativo dos complexos diferentes de zero é um produto $\mathbb{C}^\times = \mathbb{R}_+^\times \times \mathbb{S}$ do grupo multiplicativo dos reais positivos e da circunferência unitária. Também revelam, mais uma vez, o significado geométrico da multiplicação entre números complexos.

Uma primeira consequência é que o inverso do número complexo $z = \rho e^{i\theta}$, com $\rho > 0$, é $z^{-1} = \rho^{-1} e^{-i\theta}$. Outra é que a multiplicação por $z = \rho e^{i\theta} \neq 0$, no plano $\mathbb{C} \approx \mathbb{R}^2$, ou seja, a transformação $w \mapsto zw$, corresponde a uma homotetia $w \mapsto \rho w$ de razão $|z| = \rho > 0$ (uma dilatação ou contração se $\rho \neq 1$) e uma rotação $w \mapsto e^{i\theta} w$ de um ângulo θ .



Em particular, a multiplicação por um número complexo de módulo um, ou seja, da forma $e^{i\theta}$ com θ real, corresponde a uma rotação anti-horária de um ângulo θ . Por exemplo, a multiplicação por $i = e^{i\pi/2}$ é uma “raiz quadrada” da rotação $z \mapsto e^{i\pi} z = -z$ de um ângulo π , logo uma rotação de um ângulo $\pi/2$ (chamada “rotação de Wick” pelos físicos).

Exponencial. A fórmula de Euler (6.9) e a propriedade (6.10) permitem definir o *exponencial* de um número complexo arbitrário $z = x + iy$ como

$$e^z := e^x e^{iy} = e^x (\cos y + i \sin y) \quad (6.11)$$

Assim, o módulo de e^z é igual ao número real e^x , que é estritamente positivo, e o argumento de e^z é igual a y , a parte imaginária de z . Em particular, $e^z \neq 0$. É imediato então verificar, usando a

equação funcional do exponencial real e as fórmulas de adição trigonométricas, que o exponencial complexo satisfaz a regra do produto

$$e^{z+w} = e^z e^w$$

ou seja, define um homomorfismo do grupo aditivo $\mathbb{C}$ no grupo multiplicativo $\mathbb{C}^\times$. Em particular, $e^{-z} = 1/e^z$.

ex: Descreva as imagens das retas horizontais e verticais, e em particular dos eixos real e imaginário, pela função exponencial.

Raízes. Se $n = 1, 2, 3, \dots$, então cada número complexo $w \neq 0$ possui exatamente n raízes n -ésimas, ou seja, n números complexos z que resolvem

$$z^n = w.$$

Para encontrar estas raízes, é conveniente usar as coordenadas polares. Se $w = \rho e^{i\theta}$, com $\rho \neq 0$, então $z = r e^{i\varphi}$ satisfaz $z^n = w$ se

$$r^n e^{in\varphi} = \rho e^{i\theta}$$

Isto quer dizer que r é a raiz n -ésima (positiva) do número positivo ρ , e que $n\varphi = \theta + 2\pi k$, onde k é um inteiro arbitrário. Os valores de k que diferem por múltiplos de n dão origem às mesmas soluções. Consequentemente, as raízes n -ésimas de $w = \rho e^{i\theta}$ são os n números complexos

$$z_k = \sqrt[n]{\rho} e^{i(\theta+2\pi k)/n}$$

com $k = 0, 1, 2, \dots, n-1$.

Particularmente importante é o caso $w = 1$. Os números complexos

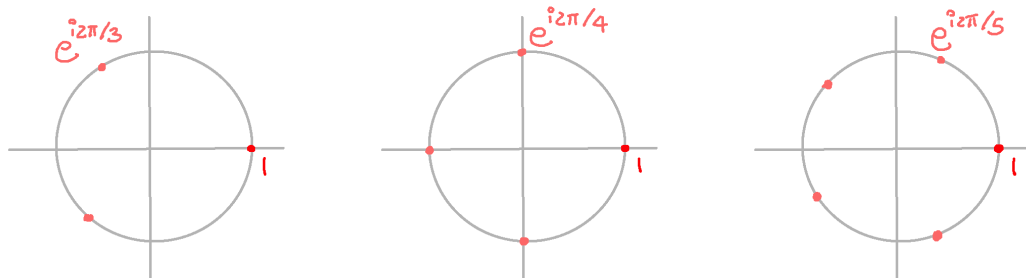
$$\zeta_k := e^{i2\pi k/n}$$

com $k = 0, 1, 2, \dots, n-1$, que resolvem $(\zeta_k)^n = 1$, são chamados *raízes n -ésimas da unidade*. Formam os vértices de um polígono regular de n lados, inscrito na circunferência unitária. Observe que $\zeta_k = (\zeta_1)^k$, onde $\zeta_1 = e^{i2\pi/n}$ é uma raiz “primitiva”. Também, $\zeta_k \zeta_{n-k} = 1$, ou seja, $\zeta_{n-k} = 1/\zeta_k$. Consequentemente, as raízes n -ésima da unidade formam um subgrupo de ordem n do grupo multiplicativo $\mathbb{S}$. A identidade elementar

$$(z-1)(z^{n-1} + z^{n-2} + \dots + z + 1) = z^n - 1$$

mostra que as raízes n -ésimas da unidade, que permitem “construir” um polígono regular de n lados, são raízes também do polinômio mônico

$$z^{n-1} + z^{n-2} + \dots + z + 1$$



ex: Represente na forma polar os seguintes números complexos:

$$-i \quad i-1 \quad 1+i \quad 3-4i$$

ex: Verifique que o conjugado de $z = \rho e^{i\theta}$ é $\bar{z} = \rho e^{-i\theta}$.

ex: Calcule

$$e^{i\pi} \quad e^{-i\pi/2} \quad \sqrt{i} \quad \sqrt{-i} \quad \sqrt{1+i} \quad \sqrt[4]{i}$$

ex: Resolva as equações $z^3 = 1$, $z^5 = 1$ e $z^3 = 81$.

ex: Verifique que $(1 + z + z^2 + \dots + z^n)(1 - z) = 1 - z^{n+1}$, e portanto, se $z \neq 1$,

$$1 + z + z^2 + \dots + z^n = \frac{1 - z^{n+1}}{1 - z}$$

Considere $z = e^{i\theta}$ com $\theta \neq 2\pi\mathbb{Z}$ e real, calcule a parte real e deduza

$$1 + \cos(\theta) + \cos(2\theta) + \dots + \cos(n\theta) = \frac{1}{2} + \frac{\sin((n+1/2)\theta)}{2 \sin(\theta/2)}$$

ex: Mostre que se ω é uma raiz n -ésima não trivial da unidade (ou seja, $\omega^n = 1$ e $\omega \neq 1$) então

$$1 + \omega + \omega^2 + \omega^3 + \dots + \omega^{n-1} = 0.$$

ex: Use a representação polar e a fórmula de Euler (6.9) para provar a *fórmula de de Moivre*

$$(\cos(\theta) + i \sin(\theta))^n = \cos(n\theta) + i \sin(n\theta). \quad (6.12)$$

Deduza fórmulas

$$\cos(n\theta) = \dots \quad \text{e} \quad \sin(n\theta) = \dots$$

para valores pequenos de n .

ex: Deduza que existem polinômios algébricos $T_n(x)$ de grau $n \geq 0$, (chamados *polinômios de Chebyshev*) tais que

$$\cos(n\theta) = T_n(\cos \theta)$$

(observe que as potências pares de $\sin \theta$ podem ser substituídas por potências pares de $\cos \theta$ usando a identidade trigonométrica), e calcule os primeiros.

ex: Calcule

$$\left(\frac{1}{\sqrt{2}} - \frac{1}{\sqrt{2}}i\right)^{13} \quad \left(\frac{\sqrt{3}}{2} + \frac{1}{2}i\right)^{17}$$

Norma e métrica euclidiana. É evidente que a parte real e a parte imaginária de um número complexo $z = x + iy$ são limitadas pelo módulo, ou seja, $|x| \leq |z|$ e $|y| \leq |z|$. Por outro lado, um cálculo direto mostra que $|z \pm w|^2 = |z|^2 + |w|^2 \pm 2\Re(z\bar{w})$. Portanto, sendo $|\Re(z\bar{w})| \leq |z||w|$, o módulo satisfaz a *desigualdade do triângulo*

$$\boxed{|z + w| \leq |z| + |w|} \quad (6.13)$$

A desigualdade do triângulo diz que $|z|$ é uma norma, e portanto

$$d(z, w) := |z - w|$$

é uma métrica no plano complexo. Ou seja, é positiva quando $z \neq w$, nula sse $z = w$, e satisfaz a desigualdade do triângulo

$$d(z, w) \leq d(z, p) + d(p, w).$$

De fato, como já observado, é a métrica euclidiana de $\mathbb{C} \approx \mathbb{R}^2$, definida pelo produto escalar euclidiano

$$\langle (x, y), (x', y') \rangle = xx' + yy' = \Re(z\bar{z}')$$

Um caso particular da desigualdade do triângulo é $|x + iy| \leq |x| + |y|$, e diz que o módulo de um número complexo é limitado pela soma dos módulos das suas partes real e imaginária.

ex: Diga quando vale a igualdade na (6.13).

ex: Verifique que o produto escalar euclidiano entre os vetores z e w de $\mathbb{C} \approx \mathbb{R}^2$ é igual a

$$\langle z, w \rangle = \Re(z, \bar{w})$$

Deduzza que coseno e seno do ângulo θ entre tais vetores (supostos não nulos) são

$$\cos \theta = \frac{\Re(z, \bar{w})}{|z||w|} \quad \text{e} \quad \sin \theta = \frac{\Re(z, i\bar{w})}{|z||w|}$$

ex: Mostre que

$$|z + w|^2 + |z - w|^2 = 2(|z|^2 + |w|^2)$$

ex: Prove a desigualdade

$$|z \pm w| \geq ||z| - |w||$$

ex: O *norma do supremo* no plano $\mathbb{C} \approx \mathbb{R}^2$ é definida por $\|x + iy\|_\infty := \max\{|x|, |y|\}$. Mostre que as normas $\|\cdot\|_\infty$ e $|\cdot|$ são equivalentes, ou seja,

$$\|z\|_\infty \leq |z| \leq \sqrt{2} \|z\|_\infty \quad (6.14)$$

(é caso particular de um teorema mais geral, que diz que todas as normas de um espaço vetorial de dimensão finita são equivalentes). Estas desigualdades dizem que $|z|$ é pequeno quando $\|z\|_\infty$ é pequeno, e vice-versa. Portanto, as topologias (i.e. as noções de aberto, de limite de uma sucessão, ...) geradas pelas duas normas são equivalentes.

6.3 Polinômios e fatorização

O “teorema fundamental da álgebra” afirma que um polinômio de grau n possui exatamente n raízes no plano complexo, se contadas corretamente.

Polinômios. Usando repetidamente somas e multiplicações, é possível construir os *polinômios*

$$f(z) = a_n z^n + a_{n-1} z^{n-1} + \cdots + a_1 z + a_0 \quad (6.15)$$

com coeficientes $a_k \in \mathbb{C}$, sendo pelos menos um deles não nulo. O *grau* de um polinômio é o maior dos índices dos seus coeficientes não nulos. Assim, o grau do polinômio (6.15) é $\deg(f) = n$ se $a_n \neq 0$. Por exemplo, uma função constante $f(z) = \alpha$ é um polinômio de grau 0, uma função afim $f(z) = \alpha z + \beta$ com $\alpha \neq 0$ é um polinômio de grau 1, uma função quadrática $f(z) = \alpha z^2 + \beta z + \gamma$ com $\alpha \neq 0$ é um polinômio de grau 2, ...

Somas e produtos finitos de polinômios são polinômios. Portanto, o espaço Pol dos polinômios forma um anel comutativo, cuja identidade é o polinômio constante $f(z) = 1$. É claro que todo polinômio de grau n é proporcional a um polinômio *mônico* de grau n , um polinômio cujo coeficiente de grau máximo é igual a $b_n = 1$, ou seja, do gênero

$$g(z) = z^n + b_{n-1} z^{n-1} + \cdots + b_1 z + b_0.$$

Observe também que $\deg(fg) = \deg(f) + \deg(g)$, desde que f e g sejam diferentes do polinômio nulo $h(z) = 0$ (que portanto não convém chamar “polinômio”).

ex: É possível dizer alguma coisa sobre o grau de uma soma de polinômios, ou seja, $\deg(f + g)$?

ex: E sobre o grau de uma composição de polinômios, ou seja, $\deg(f \circ g)$?

Zeros e fatorização. Os zeros, ou raízes, do polinómio (6.15) são os pontos p onde $f(p) = 0$, que formam o conjunto $Z(f) := \{ p \in \mathbb{C} \text{ t.q. } f(p) = 0 \}$. É claro que as raízes de $f(z)$ são também raízes do polinómio mónico $f(z)/a_n$, e vice-versa.

Polinómios com coeficientes reais podem não ter raízes reais, como no caso do polinómio quadrático $x^2 + 1$. Por outro lado, em 1799 Gauss (aos 22 anos) deu pela primeira vez umas provas razoavelmente corretas do que hoje é chamado “teorema fundamental da álgebra” (que seria mais correto chamar, como também sugerido por van der Werden [Wa91], “teorema fundamental da teoria dos números complexos”), que os matemáticos resimem ao dizer que “o corpo dos números complexos é algebricamente fechado”. Apesar do nome tradicional, deve ser considerado um teorema de análise, pois as suas provas mais elementares usam pelo menos o teorema de Weierstrass sobre a existência de máximos de funções contínuas definidas em compactos, assim como a bi-dimensionalidade do plano complexo.

Teorema 6.1 (Gauss). *Todo polinómio de grau $n \geq 1$ admite (pelo menos) uma raiz em $\mathbb{C}$.*

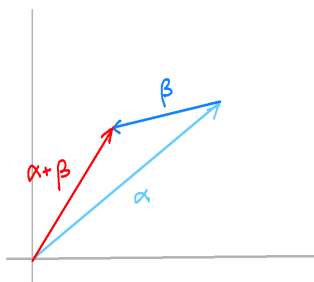
Demonstração. É claro que, a menos de dividir pelo coeficiente de grau máximo, basta provar o teorema para um polinómio mónico $f(z) = z^n + a_{n-1}z^{n-1} + \dots + a_1z + a_0$ com $n \geq 1$. É também claro que se $|z| = R$ é suficientemente grande então também $|f(z)| \simeq R^n$ é grande, por exemplo superior ao valor de $|f(0)| = |a_0|$. Pelo teorema de Weierstrass, a função contínua $z \mapsto |f(z)|$ atinge um mínimo no disco fechado $|z| \leq R$, e, pela observação anterior, este mínimo é atingido num ponto p com $|p| < R$, e é também um mínimo absoluto de $|f(z)|$ no plano complexo. A menos de uma translação (ou seja de considerar o polinómio $f(z - p)$) podemos assumir que o mínimo é atingido na origem, e portanto que o polinómio tem a forma

$$f(z) = \alpha + z^m(b_0 + b_1z + \dots + b_kz^k)$$

com $b_0 \neq 0$, $m \geq 1$, e que o seu mínimo módulo é $|f(0)| = |\alpha|$. Queremos provar que $\alpha = 0$, ou seja, que a origem é uma raiz do polinómio $f(z)$. Seja então $\alpha = \rho e^{i\theta} \neq 0$. O polinómio é da forma $\alpha + \beta(z)$, onde $\beta(z)$ é um polinómio de grau $\geq m \geq 1$ que se anula na origem. Se o módulo de z é muito pequeno, então

$$\beta(z) = z^m(b_0 + b_1z + \dots + b_kz^k) \simeq z^m b_0.$$

Ao variar z numa circunferência suficientemente pequena à volta da origem, $\beta(z)$ dá pelo menos uma volta (de fato, m voltas) em torno da origem, e portanto necessariamente o seu argumento assume valores próximos do oposto do argumento de α . Então, se $\alpha \neq 0$ e β é pequeno, é claro que $|\alpha + \beta| < |\alpha|$, como mostra a figura seguinte.



Mais precisamente, num ponto onde $\beta(z) \simeq \varepsilon e^{-i\theta}$ com $\varepsilon \ll \rho$,

$$|f(z)| = |\alpha + \beta(z)| \simeq |\rho e^{i\theta} + \varepsilon e^{-i\theta}| = \rho - \varepsilon < |\alpha|,$$

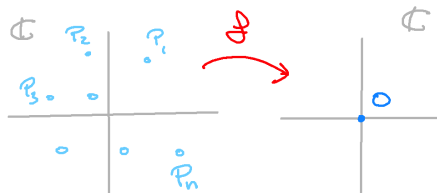
o que contradiz o fato de $|\alpha|$ ser o mínimo absoluto de $|f(z)|$. $\square$

Se p é uma raiz do polinómio $f(z)$ de grau $n \geq 1$, então $f(z) = (z - p)g(z)$ onde $g(z)$ é um polinómio de grau $n - 1$ (exercício). O teorema de Gauss 6.15 então implica o seguinte

Teorema 6.2 (teorema fundamental da álgebra). *Um polinómio $f(z) = a_nz^n + \dots + a_1z + a_0$ de grau $n \geq 1$ fatoriza no produto*

$$f(z) = a_n(z - p_1)(z - p_2) \dots (z - p_n), \quad (6.16)$$

onde $p_1, p_2, \dots, p_n$ são as suas n raízes (não necessariamente distintas).



A fatorização (6.16) é única, a menos de permutações dos fatores. Vice-versa, é evidente que existe um único polinómio mónico de grau n cujas raízes são n números complexos distintos $p_1, p_2, \dots, p_n$, pois basta multiplicar os $(z - p_k)$'s.

Se uma raiz p é repetida exatamente k vezes, ou seja, se

$$f(z) = (z - p)^k g(z)$$

com $g(p) \neq 0$, então o inteiro $k \in \mathbb{N}$ é chamado *multiplicidade* da raiz p , ou também “*ordem* de f no ponto p ”, e denotado por $\text{ord}(f, p) = k$. É natural chamar os pontos p onde $f(p) \neq 0$ pontos de ordem $\text{ord}(f, p) = 0$. A maneira correcta de “contar” o número de zeros de um polinómio é

$$|Z(f)| := \sum_{p \in \mathbb{C}} \text{ord}(f, p)$$

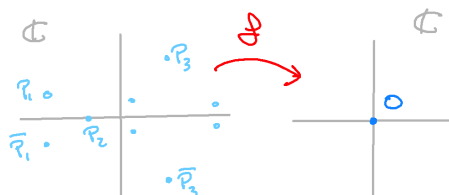
(observe que a soma é finita porque apenas um número finito de pontos têm ordem $\neq 0$). O teorema fundamental da álgebra diz então que esta soma é igual a $|Z(f)| = \deg(f)$.

ex: Verifique que ordem do produto pontual de dois polinómios é igual a soma das ordens dos fatores, ou seja, $\text{ord}(f \cdot g, p) = \text{ord}(f, p) + \text{ord}(g, p)$.

Raízes de polinómios reais. Se p é uma raiz do polinómio $f(z) = a_n z^n + a_{n-1} z^{n-1} + \dots + a_1 z + a_0$, então $\bar{p}$ é uma raiz do polinómio $\bar{f}(z) := \bar{a}_n z^n + \dots + \bar{a}_1 z + \bar{a}_0$, obtido de $f(z)$ trocando cada coeficiente pelo seu conjugado. Se os coeficientes de $f(z)$ são reais (i.e. $a_k = \bar{a}_k$), então $\bar{f} = f$. Consequentemente,

Teorema 6.3. *As raízes não reais de um polinómio com coeficientes reais ocorrem em pares de números complexos conjugados, p e $\bar{p}$.*

Ou seja, o conjunto $Z(f)$ das raízes de um polinómio real é simétrico em relação ao eixo real do plano complexo.



ex: Se p uma raiz do polinómio (6.15) de grau $n \geq 1$, então

$$f(z) - f(p) = a_n(z^n - p^n) + \dots + a_1(z - p).$$

Use a identidade

$$z^k - p^k = (z - p)(z^{k-1} + z^{k-2}p + \dots + zp^{k-2} + p^{k-1})$$

e deduza que $f(z) = (z - p)g(z)$ onde g é um polinómio de grau $n - 1$.

7 Espaços lineares

ref: [Ap69] Vol. 2, 1.1-10 ; [La97] Ch. III

7.1 Espaços lineares

As propriedades abstratas de $\mathbb{R}^n$ são convenientemente axiomatizadas na noção de “espaço linear/vetorial”, fundamental em matemática assim como em física.

27 out 2022

Espaços lineares/vetoriais. Um *espaço linear/vetorial real*, ou espaço linear definido sobre o corpo $\mathbb{R}$ dos números reais, é um conjunto $\mathbf{V}$, formado por objetos $\mathbf{v}$'s chamados *vetores*, munido de duas operações: a *soma de vetores*, que envia todo par de vetores $\mathbf{v}$ e $\mathbf{w}$ no vetor

$$\mathbf{v} + \mathbf{w},$$

que satisfaz os axiomas

EL1 (*propriedade associativa*) $\mathbf{u} + (\mathbf{v} + \mathbf{w}) = (\mathbf{u} + \mathbf{v}) + \mathbf{w}$,

EL2 (*propriedade comutativa*) $\mathbf{v} + \mathbf{w} = \mathbf{w} + \mathbf{v}$

EL3 (*existência do elemento neutro*) existe um vetor $\mathbf{0}$ tal que $\mathbf{0} + \mathbf{v} = \mathbf{v}$ para todo $\mathbf{v}$

EL4 (*existência do oposto*) para todo $\mathbf{v}$ existe um vetor $\mathbf{v}'$ tal que $\mathbf{v} + \mathbf{v}' = \mathbf{0}$

e a *multiplicação/produto por escalares*, que envia um vetor $\mathbf{v}$ e um número $\lambda \in \mathbb{R}$ no vetor

$$\lambda \mathbf{v}$$

que satisfaz os axiomas

EL5 (*existência da identidade*) $1\mathbf{v} = \mathbf{v}$

EL6 (*propriedade associativa*) $(\lambda\mu)\mathbf{v} = \lambda(\mu\mathbf{v})$

EL7 (*propriedade distributiva para a adição escalar*) $(\lambda + \mu)\mathbf{v} = \lambda\mathbf{v} + \mu\mathbf{v}$

EL8 (*propriedade distributiva para a adição vetorial*) $\lambda(\mathbf{v} + \mathbf{w}) = \lambda\mathbf{v} + \lambda\mathbf{w}$

Se substituirmos o corpo $\mathbb{R}$ dos números reais pelo corpo $\mathbb{C}$ dos números complexos, obtemos a definição de *espaço linear/vetorial complexo*. Os números reais ou complexos, dependendo do caso, são chamados *escalares* do espaço vetorial. É também possível e interessante definir espaços vetoriais sobre outros corpos, como corpos finitos ...

O próprio $\mathbb{R}$, munido das operações soma e produto, é um espaço vetorial real, o mais simples. Da mesma forma, o corpo complexo $\mathbb{C}$ é um espaço vetorial complexo.

Os axiomas EL1-8 implicam uma série de propriedades naturais, que costumamos usar quase sem pensar, e cujas provas podem ser exercícios úteis. O elemento neutro é único. O oposto de cada vetor $\mathbf{v}$ é único, e é igual a $-\mathbf{v} := (-1)\mathbf{v}$. É imediato verificar que $(-\lambda)\mathbf{v} = -(\lambda\mathbf{v}) = \lambda(-\mathbf{v})$ e que $-(\mathbf{v} + \mathbf{w}) = (-\mathbf{v}) + (-\mathbf{w})$. É então conveniente simplificar estas fórmulas usando a notação $\mathbf{v} - \mathbf{w} := \mathbf{v} + (-\mathbf{w})$. Os produtos $0\mathbf{v}$ e $\lambda\mathbf{0}$ de um vetor arbitrário por 0 ou de um escalar arbitrário pelo vetor nulo $\mathbf{0}$ são iguais ao vetor nulo $\mathbf{0}$ (e, por esta razão, o escalar 0 é por vezes confundido com o vetor nulo $\mathbf{0}$, sem perigo de ambiguidade). A propriedade distributiva também implica que $\mathbf{v} + \mathbf{v} = 2\mathbf{v}$, que $\mathbf{v} + \mathbf{v} + \mathbf{v} = 3\mathbf{v}$, ... e portanto que

$$\sum_{k=1}^n \mathbf{v} = \underbrace{\mathbf{v} + \mathbf{v} + \cdots + \mathbf{v}}_{n \text{ vezes}} = n\mathbf{v}$$

Mais importantes são as leis do corte seguintes. Se $\lambda\mathbf{v} = \mathbf{0}$, então $\mathbf{v} = \mathbf{0}$ ou $\lambda = 0$. Consequentemente, se $\lambda\mathbf{v} = \mu\mathbf{v}$ então $\mathbf{v} = \mathbf{0}$ ou $\lambda = \mu$, ou também, se $\lambda\mathbf{v} = \lambda\mathbf{w}$, então $\mathbf{v} = \mathbf{w}$ ou $\lambda = 0$.

O arquétipo de um espaço vetorial real é, naturalmente, o espaço vetorial $\mathbb{R}^n$ das n -úplas $\mathbf{x} = (x_1, x_2, \dots, x_n)$ de números reais $x_k \in \mathbb{R}$, munido das operações “adição” e “produto por um escalar” definidas na seção 1. Em particular, são espaços vetoriais o “plano” $\mathbb{R}^2$ da geometria de Euclides e o “espaço” $\mathbb{R}^3$ da mecânica Newtoniana.

ex: Mostre que o elemento neutro $\mathbf{0}$ de um espaço vetorial é único, e que $0\mathbf{v} = \mathbf{0}$ para todo $\mathbf{v}$.

ex: Mostre que o oposto de cada vetor $\mathbf{v}$ de um espaço vetorial é único, e igual a $-\mathbf{v} := (-1)\mathbf{v}$ (o que justifica a notação).

ex: [Ap69] Vol 2 1.5.

O espaço linear complexo $\mathbb{C}^n$. O protótipo de um espaço vetorial complexo é o espaço

$$\mathbb{C}^n := \underbrace{\mathbb{C} \times \mathbb{C} \times \cdots \times \mathbb{C}}_{n \text{ vezes}}$$

das n -uplas $\mathbf{z} = (z_1, z_2, \dots, z_n)$ de números complexos $z_k \in \mathbb{C}$, munido da *adição*, definida por

$$\mathbf{z} + \mathbf{w} := (z_1 + w_1, z_2 + w_2, \dots, z_n + w_n)$$

e da *multiplicação por um escalar*, definida por

$$\lambda \mathbf{z} := (\lambda z_1, \lambda z_2, \dots, \lambda z_n)$$

onde agora os escalares λ são números complexos. O vetor nulo é o vetor $\mathbf{0} := (0, 0, \dots, 0)$. O oposto do vetor $\mathbf{z} = (z_1, z_2, \dots, z_n)$ é o vetor $-\mathbf{z} := (-1)\mathbf{z} = (-z_1, -z_2, \dots, -z_n)$. A multiplicação pelo escalar 1 transforma um vetor em si próprio, ou seja, $1\mathbf{z} = \mathbf{z}$, e a multiplicação pelo escalar “zero” produz o vetor nulo, ou seja, $0\mathbf{z} = \mathbf{0}$. É imediato verificar que as operações assim definidas em $\mathbb{C}^n$ satisfazem os axiomas EL1-8.

Também é possível definir a *conjugação*, a involução que envia o vetor $\mathbf{z} = (z_1, z_2, \dots, z_n)$ no vetor

$$\bar{\mathbf{z}} := (\bar{z}_1, \bar{z}_2, \dots, \bar{z}_n)$$

É então imediato verificar que a conjugação respeita as somas, ou seja,

$$\overline{\mathbf{z} + \mathbf{w}} = \bar{\mathbf{z}} + \bar{\mathbf{w}}$$

mas é anti-linear, no sentido em que

$$\overline{\lambda \mathbf{z}} = \bar{\lambda} \bar{\mathbf{z}}$$

se $\lambda \in \mathbb{C}$ é um escalar. Os vetores que satisfazem $\bar{\mathbf{z}} = \mathbf{z}$ formam o subconjunto $\mathbb{R}^n \subset \mathbb{C}^n$, formado por n -úplas de números reais.

Subespaços e geradores. Seja $\mathbf{V}$ um espaço vetorial, real ou complexo. A *combinação linear (finita)* dos vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m \in \mathbf{V}$ com *coeficientes* os escalares $\lambda_1, \lambda_2, \dots, \lambda_m$ é o vetor

$$\lambda_1 \mathbf{v}_1 + \lambda_2 \mathbf{v}_2 + \cdots + \lambda_m \mathbf{v}_m = \sum_{k=1}^m \lambda_k \mathbf{v}_k$$

Pelas propriedades associativa e comutativa, EL1 e EL2, podemos não usar parêntesis na notação, e este vetor não depende da ordem dos adendos.

Um subconjunto não vazio $W \subset \mathbf{V}$ que, munido das operações definidas em $\mathbf{V}$, é ele próprio um espaço vetorial (ou seja, tal que $\mathbf{w} + \mathbf{w}' \in W$ e $\lambda \mathbf{w} \in W$ para todos os $\mathbf{w}, \mathbf{w}' \in W$ e $\lambda \in \mathbb{R}$) é dito *subespaço (linear/vetorial)* de $\mathbf{V}$.

É claro que um subconjunto $W \subset \mathbf{V}$ é um subespaço sse contém todas as combinações lineares

$$\lambda \mathbf{v} + \mu \mathbf{w}$$

dos seus vetores $\mathbf{v}, \mathbf{w} \in W$ com coeficientes arbitrários λ, μ , e, consequentemente, as combinações lineares finitas dos seus vetores com coeficientes arbitrários. Em particular, todo subespaço contém o vetor nulo $\mathbf{0}$ (a combinação linear trivial), e o subespaço minimal de $\mathbf{V}$ é o subespaço trivial $\{\mathbf{0}\}$, também denotado simplesmente por 0 . O subespaço maximal é o próprio $\mathbf{V}$. É claro também que

a interseção $\cap_k V_k$ de uma família (finita ou não) de subespaços $V_k \subset \mathbf{V}$ é também um subespaço de $\mathbf{V}$.

Se $S \subset \mathbf{V}$ é um subconjunto arbitrário (finito ou não) de $\mathbf{V}$, o conjunto $\text{Span}(S)$ das combinações lineares finitas

$$\lambda_1 \mathbf{v}_1 + \lambda_2 \mathbf{v}_2 + \cdots + \lambda_m \mathbf{v}_m$$

com $\mathbf{v}_k \in S$ e coeficientes λ_k escalares arbitrários, é um subespaço de $\mathbf{V}$, dito subespaço *gerado* por S , os *expansão linear* do subconjunto S . Os vetores de S são então chamados *geradores* do subespaço $\text{Span}(S)$.

É claro que se $W \subset \mathbf{V}$ é um subespaço, então $\text{Span}(W)$ é o próprio W .

Transformações lineares e isomorfismos. É importante observar que muitas propriedades interessantes dos espaços vetoriais não dependem da “natureza” dos seus elementos (pontos, vetores, sequências, funções, ...) mas apenas da estrutura algébrica descrita nos axiomas EL1-8 e nas suas consequências. As noções naturais de “morfismos” (ou seja, transformações que respeitam a estrutura) entre espaços lineares e portanto de “equivalências” são as seguintes.

Uma *transformação linear* entre os espaços lineares $\mathbf{V}$ e $\mathbf{W}$ (sobre o mesmo corpo, real ou complexo) é uma aplicação $T : \mathbf{V} \rightarrow \mathbf{W}$ que respeita as operações de soma e produto por escalares, ou seja, tal que para cada dois vetores $\mathbf{v}, \mathbf{v}' \in \mathbf{V}$ e cada escalar λ

$$T(\mathbf{v} + \mathbf{v}') = T(\mathbf{v}) + T(\mathbf{v}') \quad \text{e} \quad T(\lambda \mathbf{v}) = \lambda T(\mathbf{v})$$

É imediato verificar que uma transformação linear envia o elemento neutro no elemento neutro, e o oposto de um vetor no oposto da sua imagem. O estudo das transformações lineares é o objeto dos restantes capítulos destas notas. Nesta altura, é apenas importante para definir a seguinte noção.

Um *isomorfismo (linear)* entre os espaços lineares $\mathbf{V}$ e $\mathbf{W}$ (sobre o mesmo corpo) é uma transformação linear bijetiva $\Phi : \mathbf{V} \rightarrow \mathbf{W}$. A aplicação inversa $\Phi^{-1} : \mathbf{W} \rightarrow \mathbf{V}$ (que existe porque Φ é injetiva e sobrejetiva), definida por $\Phi^{-1}(\mathbf{w}) := \mathbf{v}$ se $\Phi(\mathbf{v}) = \mathbf{w}$, também é linear. De fato, se $\Phi(\mathbf{v}) = \mathbf{w}$ e $\Phi(\mathbf{v}') = \mathbf{w}'$ e λ é um escalar, então $\Phi(\mathbf{v} + \mathbf{v}') = \mathbf{w} + \mathbf{w}'$ e $\Phi(\lambda \mathbf{v}) = \lambda \mathbf{w}$, e portanto

$$\Phi^{-1}(\mathbf{w} + \mathbf{w}') = \mathbf{v} + \mathbf{v}' = \Phi^{-1}(\mathbf{w}) + \Phi^{-1}(\mathbf{w}') \quad \text{e} \quad \Phi^{-1}(\lambda \mathbf{w}) = \lambda \mathbf{v} = \lambda \Phi^{-1}(\mathbf{w})$$

É claro que a existência de um isomorfismo entre dois espaços vetoriais é uma relação de equivalência, pois a identidade é um isomorfismo de um espaço em si próprio, e a composição de dois isomorfismos é também um isomorfismo. Dois espaços vetoriais isomorfos são indistinguíveis enquanto espaços lineares. Uma notação para indicar que dois espaços lineares são isomorfos, sem necessariamente especificar o isomorfismo ϕ (que não é único se os espaços não são triviais!), é $\mathbf{V} \approx \mathbf{W}$.

ex: Verifique que a composição de dois isomorfismos é um isomorfismo.

e.g. O plano complexo enquanto espaço vetorial real. O plano complexo $\mathbb{C}$ pode ser considerado um espaço vetorial real. O número complexo $z = x + iy$ pode ser representado como uma soma formal $x\mathbf{1} + y\mathbf{i}$. A soma entre $z \approx x\mathbf{1} + y\mathbf{i}$ e $z' \approx x'\mathbf{1} + y'\mathbf{i}$ é então definida por

$$(x\mathbf{1} + y\mathbf{i}) + (x'\mathbf{1} + y'\mathbf{i}) := (x + x')\mathbf{1} + (y + y')\mathbf{i}$$

e o produto do escalar $\lambda \in \mathbb{R}$ vezes o vetor $z \approx x\mathbf{1} + y\mathbf{i}$ é definido por

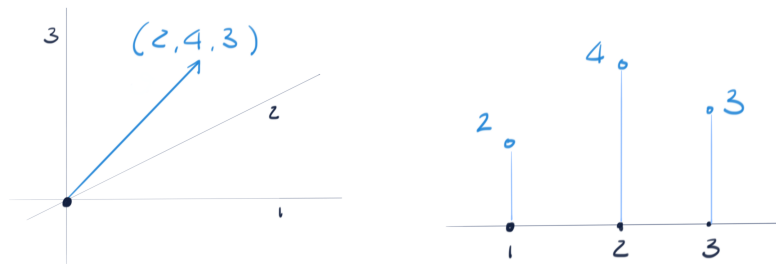
$$\lambda(x\mathbf{1} + y\mathbf{i}) := (\lambda x)\mathbf{1} + (\lambda y)\mathbf{i}$$

É imediato verificar que $x\mathbf{1} + y\mathbf{i} \mapsto (x, y)$ define um isomorfismo entre o espaço vetorial real $\mathbb{C}$ e o espaço vetorial real $\mathbb{R}^2$. É por esta razão que o corpo dos números complexos é chamado “plano complexo”.

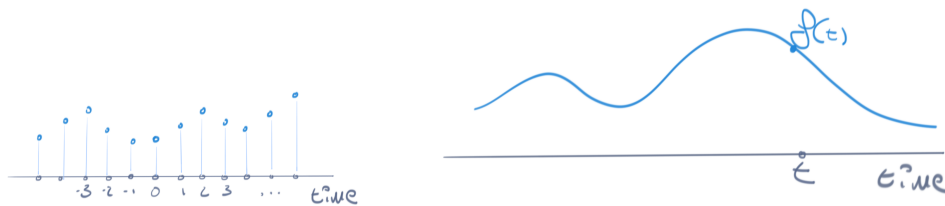
7.2 Espaços de funções

Os espaços interessantes em análise, em física e em engenharia, são espaços de funções, chamados “espaços funcionais”.

Vetores & funções. Os espaços vetoriais mais familiares são, naturalmente, o plano e o espaço da geometria de Euclides. Consequentemente, visualizamos vetores como setas saindo da origem. No entanto, as coordenadas do vetor $(2, 4, 3)$ podem representar os valores da temperatura medida em três diferentes dias, por exemplo, domingo, segunda e terça-feira. Uma visualização mais natural deste vetor, sendo que o nosso tempo tem uma orientação natural, é então a segunda, mostrada no desenho. Tecnicamente, este é o gráfico de uma função definida no conjunto finito formado pelos pontos/dias 1, 2, 3, com valores $x[1] = 2$, $x[2] = 4$ e $x[3] = 3$.



Se observamos temperaturas durante muitos, ou até idealmente infinitos, dias, então a imagem natural do vetor é o gráfico de uma sucessão de números $(\dots, 2, 4, 3, \dots)$, dependendo de um tempo parametrizado por números inteiros. Ou seja, um valor $x[n]$ para cada tempo n , que para os engenheiros é um “sinal discreto”. Vista de muito longe, ou seja, idealizando o tempo como um contínuo, esta parece o gráfico de uma função $f(t)$ de uma variável real, que os engenheiros chamam “sinal contínuo”.



Vice-versa, um sinal contínuo $f(t)$ pode ser observado apenas em múltiplos inteiros de um “tempo de amostragem” $\tau > 0$, e transformado num sinal discreto $x[n] := f(n\tau)$, com $n \in \mathbb{Z}$ ou $\mathbb{N}$, que é uma sequência. Um problema natural é tentar “reconstruir” o sinal contínuo a partir da sua amostragem ...

Espaços de funções. Seja X um conjunto arbitrário. Denotamos por $\mathcal{F}(X, \mathbb{R})$, ou mais simplesmente por $\mathbb{R}^X$, o espaço

$$\mathcal{F}(X, \mathbb{R}) := \{\text{funções } f : X \rightarrow \mathbb{R}\}$$

de todas as funções com valores reais definidas em X . Da mesma forma, denotamos por $\mathcal{F}(X, \mathbb{C})$ ou $\mathbb{C}^X$ o espaço de todas as funções com valores complexos definidas em X . Os elementos destes espaços são também denotados brevemente por $x \mapsto f(x)$, ou simplesmente por $f(x)$ ou apenas f , quando o contexto permite. São espaços vetoriais reais e complexos, respetivamente, se munidos das operações “adição” e “produto por um escalar” definidas por

$$(f + g)(x) := f(x) + g(x) \quad \text{e} \quad (\lambda f)(x) := \lambda f(x)$$

O elemento neutro é a função identicamente nula, $0(x) = 0$ para todos $x \in X$. O oposto da função $f(t)$ é a função definida por $(-f)(x) = -f(x)$. Dois vetores deste espaço, ou seja, duas funções f e g , são iguais sse $f(x) = g(x)$ para todos os $x \in X$ (devem pensar que os valores $f(x)$ são as “coordenadas” do vetor f , uma para cada ponto x de X).

Também interessantes são os “campos vetoriais”, funções com valores num espaço vetorial como $\mathbb{R}^m$ (campos de força, de velocidade, campo eletro-magnético, ...). Mais em geral, se $\mathbf{V}$ é um espaço vetorial, então o espaço $\mathbf{V}^X = \mathcal{F}(X, \mathbf{V})$ tem uma estrutura natural de espaço vetorial (real ou complexo, dependendo de $\mathbf{V}$).

Engenheiros e físicos estão por exemplo interessados em “sinais” $f(t)$ (a intensidade de uma onda de som, onde t é o tempo), ou “funções de onda” $\psi(\mathbf{r}, t)$ (em mecânica quântica), ou outros “campos” $u(\mathbf{r}, t)$ (um deslocamento, um campo de velocidades, o campo eletro-magnético, ...) que resolvem certas equações diferenciais parciais como a equação de onda, de calor, de Laplace, de Schrödinger, ...

e.g. Espaços de funções sobre conjuntos finitos. Se X é um conjunto finito composto por n elementos, por exemplo $X = \{1, 2, \dots, n\}$, então $\mathbb{R}^X$ é simplesmente o espaço vetorial $\mathbb{R}^n$. As n -úplas $(x_1, x_2, \dots, x_n)$ de números reais podem ser pensadas como os valores $x_n = f(n)$ das funções $f : X \rightarrow \mathbb{R}$. De fato, é esta a maneira mais natural (“functorial”, de acordo com Gromov) de ver os espaços vetoriais de dimensão finita.

Funções contínuas e diferenciáveis. Seja $X \subset \mathbb{R}$ um intervalo da reta real, como por exemplo $(0, 1)$, ou $[0, 1]$, ou o própria reta real $\mathbb{R}$. São espaços vetoriais, reais ou complexos, o espaço $\mathcal{C}^0(X)$ das funções contínuas $f : X \rightarrow \mathbb{R}$, os espaços $\mathcal{C}^k(X)$ das funções com k derivadas contínuas, o espaço $\mathcal{C}^\infty(X)$ das funções infinitamente deriváveis. As inclusões são

$$\mathcal{C}^\infty(X) \subset \dots \subset \mathcal{C}^{k+1}(X) \subset \mathcal{C}^k(X) \subset \dots \subset \mathcal{C}^0(X) \subset \mathbb{R}^X$$

Também é possível definir espaços de funções diferenciáveis definidas em regiões $X \subset \mathbb{R}^n$, ou em objetos mais interessantes chamados “variedades diferenciáveis” ...

Polinômios. O espaço $\text{Pol}(\mathbb{R}) := \mathbb{R}[t]$ dos polinômios $p(t) = a_n t^n + \dots + a_1 t + a_0$ com coeficientes reais na variável t (e grau n arbitrário) é o subespaço de $\mathcal{C}^\infty(\mathbb{R})$ gerado pela família enumerável dos monômios $1, t, t^2, t^3, \dots$. O espaço $\text{Pol}_n(\mathbb{R})$ dos polinômios reais de grau $\leq n$ é um subespaço de $\text{Pol}(\mathbb{R})$. Também, $\text{Pol}_n(\mathbb{R}) \subset \text{Pol}_m(\mathbb{R})$ se $n \leq m$.

Da mesma forma, é útil considerar os espaços lineares complexos $\text{Pol}_n(\mathbb{C}) \subset \text{Pol}(\mathbb{C})$ dos polinômios com coeficientes complexos ...

Sucessões. O espaço $\mathbb{R}^\infty := \mathbb{R}^{\mathbb{N}}$ é o espaço das sucessões reais, cujos elementos podem ser pensados como “vetores infinitos” $\mathbf{x} = (x_1, x_2, \dots, x_n, \dots)$, com $x_n \in \mathbb{R}$. Subespaços naturais são o espaço b das sucessões limitadas, o espaço c das sucessões convergentes, o espaço c_0 das sucessões que convergem para 0, e o espaço ℓ das sucessões com suporte compacto (i.e. tais que $x_n = 0$ se n é suficientemente grande) são subespaços vetoriais de $\mathbb{R}^\infty$. As inclusões são

$$\ell \subset c_0 \subset c \subset b \subset \mathbb{R}^\infty.$$

Da mesma forma, é possível considerar o espaço $\mathbb{Z}$ das sucessões “bi-infinitas”, cujos elementos são $\mathbf{x} = (\dots x_{-2}, x_{-1}, x_0, x_1, x_2, \dots)$, e os seus subespaços naturais.

Também é útil considerar espaços de sucessões complexas, $\mathbb{C}^\infty$ ou $\mathbb{C}^\mathbb{Z}$, e seus subespaços significativos.

ex: O conjunto dos polinômios de grau exatamente igual a n é um subespaço vetorial de $\text{Pol}(\mathbb{R})$?

ex: O conjunto das funções não negativas, ou seja, tais que $f(t) \geq 0$ pra todo t , é um subespaço do espaço $\mathbb{R}^\mathbb{R}$?

ex: O conjunto das funções $f : \mathbb{R} \rightarrow \mathbb{R}$ tais que $f(1) = 0$ é um subespaço de $\mathbb{R}^\mathbb{R}$? E o conjunto das funções tais que $f(1) = 1$?

ex: Verifique que

$$a_0 + a_1 t + a_2 t^2 + \dots + a_n t^n \mapsto (a_0, a_1, a_2, \dots, a_n)$$

define um isomorfismo entre os espaços lineares $\text{Pol}_n(\mathbb{R})$ e $\mathbb{R}^{n+1}$.

ex: Considere o espaço vetorial $\mathcal{C}^\infty(\mathbb{R})$ das funções reais infinitamente deriváveis definidas na reta real. Diga se são subespaços vetoriais os seguintes subconjuntos

$$\{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f(0) = 0\} \quad \{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f(0) = 1\} \quad \{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f(0) = f(1)\}$$

$$\{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f(0) - f(1) + f(2) = 0\} \quad \{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f(1) + f(2) + f(3) = 4\}$$

$$\{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f'(0) = 0\} \quad \{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f''(0) = 0\}$$

$$\{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f'(0) - f''(0) = 0\}$$

$$\{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f''(t) - f'(t) + f(t) = 0 \quad \forall t\} \quad \{f \in \mathcal{C}^\infty(\mathbb{R}) \text{ t.q. } f'' + 2f' + 3f = t \quad \forall t\}$$

7.3 Independência, bases e dimensão

A independência linear num espaço vetorial é definida como no caso de $\mathbb{R}^n$.

Conjuntos livres/linearmente independentes. Seja $\mathbf{V}$ um espaço linear. O conjunto $S \subset \mathbf{V}$ (finito ou não) é *livre*/(linearmente) *independente* se gera cada vetor de $\text{Span}(S)$ duma única maneira, e portanto (pela mesma prova do teorema 4.1) sse gera o vetor nulo duma única maneira, ou seja, se a única combinação linear finita nula

$$\lambda_1 \mathbf{s}_1 + \lambda_2 \mathbf{s}_2 + \cdots + \lambda_m \mathbf{s}_m = \mathbf{0}$$

com $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m \in S$ é a combinação trivial com coeficientes $\lambda_1 = \lambda_2 = \cdots = \lambda_m = 0$.

Caso contrário, o conjunto é dito (linearmente) *dependente*. Isto significa que (pelo menos) um dos vetores é dependente dos outros, ou seja, pertence ao subespaço gerado pelos outros. É claro que todo conjunto que contém o vetor nulo é dependente (por razões triviais).

ex: Verifique se $\cos t$ e $\sin t$ são linearmente independentes no espaço $\mathcal{C}^\infty(\mathbb{R})$. E os vetores $\cos^2 t$, $\sin^2 t$ e $1/2$?

ex: Verifique se os vetores $(1, -1, 1, -1, 1, -1, \dots)$ e $(1, 1, 1, 1, 1, 1, \dots)$ são linearmente independentes no espaço $\mathbb{R}^\infty$ das sucessões reais.

Dimensão finita e bases. O espaço linear $\mathbf{V}$, por exemplo um subespaço de outro espaço vetorial, tem *dimensão finita*, ou é *finitamente gerado*, se admite um conjunto finito de geradores.

Seja $\mathbf{V}$ um espaço linear de dimensão finita, real ou complexo. Uma *base* de $\mathbf{V}$ é um conjunto livre de geradores de $\mathbf{V}$. É conveniente pensar nas bases como conjuntos ordenados, ou seja, listas. Assim uma base (ordenada) é uma lista $\mathcal{B} = (\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n)$ de vetores $\mathbf{b}_k \in \mathbf{V}$ tal que cada vetor $\mathbf{v} \in \mathbf{V}$ admite uma e uma única representação

$$\mathbf{v} = v_1 \mathbf{b}_1 + v_2 \mathbf{b}_2 + \cdots + v_n \mathbf{b}_n$$

como combinação linear dos $\mathbf{b}_k$'s. Os escalares v_i são as *componentes*, ou *coordenadas*, do vetor $\mathbf{v}$ relativamente à base $\mathcal{B}$. É claro, pela prova do teorema 4.1, que as coordenadas de um vetor relativamente a uma base são únicas.

O teorema 4.2 e as suas consequências têm provas que apenas dependem da estrutura de $\mathbb{R}^n$ enquanto espaço vetorial, e portanto continuam válidos para espaços vetoriais arbitrários. Em particular, se o espaço linear $\mathbf{V}$ (ou um subespaço de um espaço linear) admite um conjunto finito de geradores, então a sua *dimensão* $\dim(\mathbf{V})$ (também denotada por $\dim_{\mathbb{R}} \mathbf{V}$ ou $\dim_{\mathbb{C}} \mathbf{V}$, dependendo se o espaço é real ou complexo) pode ser definida como sendo o número de elementos de uma base. Num espaço de dimensão n , todo conjunto livre de n elementos é uma base, e todo conjunto livre de $m \leq n$ elementos é um subconjunto de uma base. Também é claro que se X é um subespaço de $\mathbf{V}$ então $\dim(X) \leq \dim(\mathbf{V})$.

Se $(\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n)$ é uma base do espaço linear $\mathbf{V}$, então todo vetor $\mathbf{v} \in \mathbf{V}$ admite uma única representação $\mathbf{v} = v_1 \mathbf{b}_1 + v_2 \mathbf{b}_2 + \dots + v_n \mathbf{b}_n$. É imediato verificar que a aplicação

$$\mathbf{v} \mapsto (v_1, v_2, \dots, v_n)$$

define um isomorfismo $\mathbf{V} \approx \mathbb{R}^n$ ou $\mathbb{C}^n$ (dependendo se o espaço linear é real ou complexo). O isomorfismo inverso é $(v_1, v_2, \dots, v_n) \mapsto v_1 \mathbf{b}_1 + v_2 \mathbf{b}_2 + \dots + v_n \mathbf{b}_n$. Assim,

Teorema 7.1. *Todo espaço linear de dimensão finita n é isomorfo a $\mathbb{R}^n$ ou $\mathbb{C}^n$, dependendo se real ou complexo.*

Naturalmente, o isomorfismo não é único, pois depende da escolha de uma base.

ex: Verifique que

$$1 \quad t \quad t^2 \quad t^3 \quad \dots \quad t^n$$

é uma base do espaço $\text{Pol}_n(\mathbb{R})$ dos polinômios de grau $\leq n$, que portanto tem dimensão $n+1$.

ex: Determine as coordenadas do polinômio $f(t) = (1-t)^2$ relativamente à base ordenada $(1, t, t^2)$ de $\text{Pol}_2(\mathbb{R})$.

ex: [Ap69] Vol 2 1.10.

Polinômio interpolador de Lagrange. Se $p_0, p_1, p_2, \dots, p_n$ são $n+1$ pontos distintos da reta real (ou da reta complexa), então

$$\begin{aligned} f(z) &= (z - p_0)(z - p_1)(z - p_2) \dots (z - p_n) \\ &= z^{n+1} - (p_0 + p_1 + p_2 + \dots + p_n) z^n + \dots + (-1)^n (p_0 p_1 p_2 \dots p_n) \end{aligned}$$

é um polinômio mônico de grau n que se anula (apenas) nos p_k 's. Fixado $k = 0, 1, \dots, n$, o polinômio

$$f_k(z) = f(z)/(z - p_k) = \prod_{i \neq k} (z - p_i)$$

é um polinômio mônico de grau n que se anula nos p_j com $j \neq k$ e que assume um valor $f_k(p_k) \neq 0$ no ponto p_k . (e também em todos os outros pontos). Então os

$$\ell_k(z) := \frac{f_k(z)}{f_k(p_k)}$$

são polinômios mônicos de grau n que se anulam nos p_j com $j \neq k$ e valem $\ell_k(p_k) = 1$. Se $\alpha_0, \alpha_1, \alpha_2, \dots, \alpha_n$ são $n+1$ números reais (ou complexos) arbitrários, não necessariamente distintos, então

$$g(z) := \sum_{k=0}^n \alpha_k \ell_k(z)$$

é um polinômio de grau $\leq n$ tal que $g(p_k) = \alpha_k$ para todo $k = 0, 1, 2, \dots, n$, chamado *polinômio interpolador de Lagrange*.

ex: Verifique que os ℓ_k 's (ou os f_k 's) são linearmente independentes (calcule uma combinação linear nos pontos p_k 's), e portanto formam uma base de $\text{Pol}_n(\mathbb{C})$.

ex: Deduza que o polinômio interpolador de Lagrange $g(z)$ é o único polinômio de grau $\leq n$ cujo gráfico passa pelos $n+1$ pontos $(p_k, \alpha_k) \in \mathbb{C} \times \mathbb{C}$.

ex: Dado um polinômio genérico $f(z) = a_n z^n + \dots + a_1 z + a_0$ de grau $\leq n$, calcule as suas coordenadas na base formada pelos ℓ_k 's.

Quaterniões. Numa tentativa de estender o corpo dos números complexos e assim representar os pontos do espaço de dimensão 3, Hamilton descobriu¹⁶ que era necessário prescindir da comutatividade do produto e acrescentar mais uma dimensão. O resultado é um espaço vetorial real de dimensão 4, denotado por $\mathbb{H}$ em sua homenagem, munido de um produto associativo, mas não comutativo, que admite um inverso de cada vetor não nulo (os matemáticos dizem uma “álgebra associativa com divisão”). Uma base deste espaço é formada por objetos que denotamos $\mathbf{1}, \mathbf{i}, \mathbf{j}, \mathbf{k}$. Os *quaterniões* são então expressões formais

$$x = x_0\mathbf{1} + x_1\mathbf{i} + x_2\mathbf{j} + x_3\mathbf{k}$$

com coeficientes $x_k \in \mathbb{R}$. A soma e o produto por um escalar são definidos da maneira natural,

$$(x_0\mathbf{1} + x_1\mathbf{i} + x_2\mathbf{j} + x_3\mathbf{k}) + (y_0\mathbf{1} + y_1\mathbf{i} + y_2\mathbf{j} + y_3\mathbf{k}) := (x_0 + y_0)\mathbf{1} + (x_1 + y_1)\mathbf{i} + (x_2 + y_2)\mathbf{j} + (x_3 + y_3)\mathbf{k}$$

e

$$\lambda(x_0\mathbf{1} + x_1\mathbf{i} + x_2\mathbf{j} + x_3\mathbf{k}) := (\lambda x_0)\mathbf{1} + (\lambda x_1)\mathbf{i} + (\lambda x_2)\mathbf{j} + (\lambda x_3)\mathbf{k}$$

É útil (e isto é o espírito das intenções da Hamilton e da notação) identificar os quaterniões $\mathbf{i}, \mathbf{j}, \mathbf{k}$ com os vetores homônimos da base canônica de $\mathbb{R}^3$, e o quaternião $\mathbf{1}$ com o escalar 1. Desta forma, um quaternião é uma soma formal $x = x_0 + \mathbf{x}$ de um escalar $x_0 \in \mathbb{R}$, chamado “parte real”, e um vetor $\mathbf{x} = x_1\mathbf{i} + x_2\mathbf{j} + x_3\mathbf{k} \in \mathbb{R}^3$, chamado “parte vetorial”. Soma e produto por um escalar são simplesmente

$$(x_0 + \mathbf{x}) + (y_0 + \mathbf{y}) = (x_0 + y_0) + (\mathbf{x} + \mathbf{y}) \quad \lambda(x_0 + \mathbf{x}) = (\lambda x_0) + \lambda \mathbf{x}$$

Os quaterniões “escalares/reaís”, do género $t = t_0\mathbf{1}$ com $t_0 \in \mathbb{R}$, formam um subespaço isomorfo a $\mathbb{R}$.

O produto entre dois quaterniões é definido declarando que $\mathbf{1}$ é a identidade, logo satisfaz $1x = x1 = x$ para todo $x \in \mathbb{H}$, que os quaterniões escalares comutam com todos os outros, que os produtos entre os outros elementos da base são

$$\mathbf{i}\mathbf{j} = -\mathbf{j}\mathbf{i} = \mathbf{k} \quad \mathbf{j}\mathbf{k} = -\mathbf{k}\mathbf{j} = \mathbf{i} \quad \mathbf{k}\mathbf{i} = -\mathbf{i}\mathbf{k} = \mathbf{j} \quad \mathbf{i}\mathbf{i} = \mathbf{j}\mathbf{j} = \mathbf{k}\mathbf{k} = \mathbf{i}\mathbf{j}\mathbf{k} = -1 \quad (7.1)$$

(algumas derivam das precedentes), e finalmente estendido usando a propriedade distributiva. O produto que assim resulta é associativo mas não comutativo. A fórmula final para o produto entre dois quaterniões é simplificada se observamos que as primeiras três destas relações (7.1) correspondem aos produtos vetoriais (5.2) entre os vetores da base canônica de $\mathbb{R}^3$. Na notação vetorial, o produto entre dois quaterniões é portanto definido por

$$(x_0 + \mathbf{x})(y_0 + \mathbf{y}) = (x_0y_0 - \mathbf{x} \cdot \mathbf{y}) + (x_0\mathbf{y} + y_0\mathbf{x} + \mathbf{x} \times \mathbf{y}) \quad (7.2)$$

(os físicos podem reconhecer na parte real do produto a métrica de Minkowski do espaço-tempo da relatividade restrita). Os quaterniões escalares formam um corpo isomorfo a $\mathbb{R}$. Também é fácil verificar que os quaterniões “complexos”, do género $t_0\mathbf{1} + t_1\mathbf{i}$ com $t_0, t_1 \in \mathbb{R}$, formam um corpo isomorfo a $\mathbb{C}$. Por outro lado, o produto entre dois quaterniões com parte escalar nula, logo essencialmente dois vetores de $\mathbb{R}^3$, é um quaternião

$$(0 + \mathbf{x})(0 + \mathbf{y}) = \mathbf{x} \cdot \mathbf{y} + \mathbf{x} \times \mathbf{y}$$

cuja parte escalar é o produto escalar entre os vetores, e cuja parte vetorial é o produto vetorial entre os dois vetores (e esta é a origem dos nomes destes dois produtos).

O *conjugado* do quaternião $x = x_0 + \mathbf{x}$ é o quaternião $\bar{x} := x_0 - \mathbf{x}$. A conjugação é claramente uma involução, mas acontece que não respeita exatamente os produtos, pois

$$\overline{xy} = \bar{y}\bar{x} \quad (7.3)$$

como consequência da (7.2) e da anti-simetria do produto vetorial. No entanto, o produto $x\bar{x} = \bar{x}x$ de um quaternião com o seu conjugado, em qualquer ordem, é um escalar, logo um número real, não negativo

$$x\bar{x} = x_0^2 + \mathbf{x} \cdot \mathbf{x}$$

¹⁶W.R. Hamilton, *On Quaternions; or on a new System of Imaginaries in Algebra*. Letter to John T. Graves (17 October 1843).

A sua raiz quadrada é chamada *norma* de x , e denotada por $\|x\| = \sqrt{x\bar{x}}$. É claro que um quaternião é não nulo sse a sua norma é diferente de zero, logo positiva. Isto permite calcular o inverso multiplicativo de todo quaternião não nulo x pela mesma fórmula que define o inverso de um número complexo não nulo:

$$x^{-1} = \frac{\bar{x}}{\|x\|^2}$$

Naturalmente, o “quociente” entre dois quaterniões x e y , com $x \neq 0$, é qualquer uma das duas expressões $x^{-1}y$ ou yx^{-1} , em geral distintas.

Hoje sabemos ¹⁷ que as únicas álgebra associativas com divisão de dimensão finita sobre os reais são $\mathbb{R}$, $\mathbb{C}$ e $\mathbb{H}$.

ex: Considere dois quaterniões $a = a_0 + \mathbf{a}$ e $b = b_0 + \mathbf{b}$, por exemplo com coordenadas inteiras. O escalar $b\bar{b}$ comuta com todos os quaterniões, e portanto, pela (7.3),

$$\|a\|^2 \|b\|^2 = a\bar{a}b\bar{b} = a b\bar{b}\bar{a} = (ab)\overline{(ab)} = \|ab\|^2$$

Calcule as três normas, e deduza a *identidade dos quatros quadrados de Euler*

$$\begin{aligned} (a_0^2 + a_1^2 + a_2^2 + a_3^2) + (b_0^2 + b_1^2 + b_2^2 + b_3^2) &= (a_0b_0 - a_1b_1 - a_2b_2 - a_3b_3)^2 \\ &\quad + (a_0b_1 + a_1b_0 + a_2b_3 - a_3b_2)^2 \\ &\quad + (a_0b_2 - a_1b_3 + a_2b_0 + a_3b_1)^2 \\ &\quad + (a_0b_3 + a_1b_2 - a_2b_1 + a_3b_0)^2 \end{aligned}$$

que diz que um produto de duas somas de quatros quadrados (por exemplo, de números inteiros) é também uma soma de quatros quadrados (de números inteiros), logo estende a identidade de Diofanto ou de Brahmagupta-Fibonacci (6.8).

Dimensão infinita. Um espaço linear que não é finitamente gerado, ou seja, que não admite um conjunto finito de geradores, é dito espaço de *dimensão infinita*. Esta expressão quer apenas dizer que o espaço não tem dimensão finita. Acontece que tentar definir uma dimensão quando o espaço não é finitamente gerado não é útil, por exemplo não permite provar um resultado que estende o teorema 8.10.

Por exemplo, o espaço $\text{Pol}(\mathbb{R})$ dos polinómios com coeficientes reais na incógnita t tem dimensão infinita. De fato, o conjunto infinito e numerável dos monómios $\mathbf{e}_n(t) = t^n$, com $n = 0, 1, 2, 3, \dots$ é um conjunto independente (dois polinómios são iguais sse todos os coeficientes são iguais). Por outro lado, um conjunto finito de polinómios não pode gerar o espaço, pois tendo um grau máximo finito, não pode gerar polinómios de grau arbitrariamente grande.

A fortiori, têm dimensão infinita os espaços $\mathcal{C}^k(\mathbb{R})$ das funções com k derivadas contínuas definidas na reta real (ou num intervalo, ...), onde $k = 0, 1, 2, \dots, \infty$. Da mesma forma, têm dimensão infinita os espaços $\ell \subset c_0 \subset c \subset b \subset \mathbb{R}^\infty$ ou $\mathbb{C}^\infty$ das sucessões, reais ou complexas.

Usando o axioma da escolha, é possível mostrar que todo espaço linear admite uma base, ou seja, um conjunto livre de geradores. No entanto, em particular nos problemas práticos da física e da engenharia, é tipicamente mais útil tentar aproximar vetores/funções genéricos com combinações lineares finitas de um conjunto numerável de vetores independentes particularmente “simples” ou significativos. Este é o espírito, por exemplo, da teoria das séries de Fourier, que estuda a possibilidade de aproximar uma função $f(t)$ periódica de período 2π por meio de “polinómios trigonométricos”, combinações lineares finitas $\sum_{n=-N}^N e^{int}$ de harmónicas $e^{int} \dots$

ex: Verifique que o conjunto numerável dos monómios

$$\mathbf{e}_0(t) = 1 \quad \mathbf{e}_1(t) = t \quad \mathbf{e}_2(t) = t^2 \quad \mathbf{e}_3(t) = t^3 \dots$$

é linearmente independentes no espaço $\text{Pol}(\mathbb{R})$ dos polinómios com coeficientes reais. Sendo também um conjunto de geradores, é uma base: todo polinómio é uma combinação finita única dos $\mathbf{e}_k$'s.

¹⁷F.G. Frobenius, Über lineare Substitutionen und bilineare Forme, *Journal für die reine und angewandte Mathematik* **84** (1878), 1-63.

ex: Verifique que o conjunto numerável das “harmônicas” $\mathbf{h}_n(t) = e^{int}$, com $n \in \mathbb{Z}$, é um conjunto independente no espaço das funções periódicas de período 2π com valores complexos.

ex: Verifique que o conjunto não numerável dos exponenciais $e^{\lambda t}$, com $\lambda \in \mathbb{R}$, é um conjunto independente no espaço linear $\mathcal{C}^0(\mathbb{R})$ das funções contínuas de uma variável real.

Superposition principle in quantum mechanics. If a physical system can be prepared in each of the states $|1\rangle, |2\rangle, |3\rangle, \dots$, for example corresponding to certain values $\lambda_1, \lambda_2, \lambda_3, \dots$ of a certain observable L , then the most general state is a *superposition* ¹⁸

$$|\psi\rangle = \psi_1 |1\rangle + \psi_2 |2\rangle + \psi_3 |3\rangle + \dots$$

with complex coefficients $\psi_n \in \mathbb{C}$. The observation of the observable L in the state $|\psi\rangle$ will then give one of the value λ_n (and not a value in between!) with relative frequency (if the experience is repeated a large number of times) proportional to $|\psi_n|^2$. States of a quantum system therefore live in a complex linear space $\mathbf{H}$ (which typically is infinite dimensional). Actually, proportional states must be considered as equal, which amounts to say that we may only consider normalized states, those with $|\psi_1|^2 + |\psi_2|^2 + |\psi_3|^2 + \dots = 1$, and do not distinguish between those which differ by a multiplicative phase $e^{i\phi}$ (although this ambiguity is then central in quantum field theory). Therefore the state space of a quantum system is a projective space $\mathbb{P}\mathbf{H}$ of a complex linear space $\mathbf{H}$.

Rational linear independence on the line. The real numbers $x_1, x_2, \dots, x_n$ are said *linear independent over the field of rationals* if the only solution of

$$k_1 x_1 + k_2 x_2 + \dots + k_n x_n = 0$$

with $k_i \in \mathbb{Z}$ is the trivial solution $k_1 = k_2 = \dots = k_n = 0$.

ex: Show that (any finite subset of) the sequence

$$\log 2, \log 3, \log 5, \log 7, \dots, \log p, \dots$$

of natural logarithms of prime numbers is linear independent over the rationals (use the unique factorization of an integer into prime factors) ¹⁹.

7.4 Produtos e quocientes

Somas e somas diretas. Se X e Y são subespaços de $\mathbf{V}$, então a *soma*

$$X + Y := \{\mathbf{x} + \mathbf{y} \text{ com } \mathbf{x} \in X, \mathbf{y} \in Y\}$$

é um subespaço de $\mathbf{V}$. Da mesma forma é definida a soma $X_1 + X_2 + \dots + X_m = \sum_{k=1}^m X_k$ de uma família finita de subespaços $X_k \subset \mathbf{V}$, o conjunto das possíveis combinações lineares $\mathbf{v}_1 + \mathbf{v}_2 + \dots + \mathbf{v}_m$ com $\mathbf{v}_k \in X_k$.

Usando sucessivamente o teorema 4.5 é possível ver que

¹⁸“Any state may be considered as the result of a superposition of two or more other states, and indeed in an infinite number of ways. Conversely any two or more states may be superposed to give a new state ... The non-classical nature of the superposition process is brought out clearly if we consider the superposition of two states, A and B, such that there exists an observation which, when made on the system in state A, is certain to lead to one particular result, a say, and when made on the system in state B is certain to lead to some different result, b say. What will be the result of the observation when made on the system in the superposed state? The answer is that the result will be sometimes a and sometimes b, according to a probability law depending on the relative weights of A and B in the superposition process. It will never be different from both a and b [i.e, either a or b]. The intermediate character of the state formed by superposition thus expresses itself through the probability of a particular result for an observation being intermediate between the corresponding probabilities for the original states, not through the result itself being intermediate between the corresponding results for the original states.”

P.A.M. Dirac [Di47]

¹⁹H. Bohr, 1910.

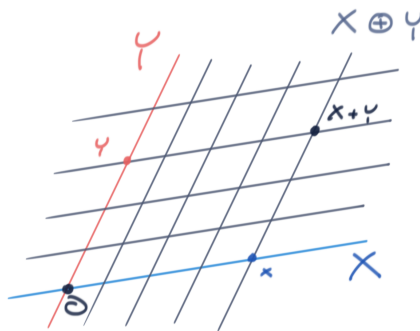
Teorema 7.2. Se X e Y são subespaços do espaço linear de dimensão finita V , então

$$\dim(X + Y) = \dim(X) + \dim(Y) - \dim(X \cap Y)$$

Sejam X, Y subespaços do espaço vetorial V . Se cada vetor $v \in X + Y$ admite uma única representação $v = x + y$ com $x \in X$ e $y \in Y$, então a soma $X + Y$ dos subespaços X e Y é chamada *soma direta*, e denotada por

$$X \oplus Y$$

É claro que isto acontece quando $X \cap Y = 0$. Se o espaço total é uma soma direta $V = X \oplus Y$ de dois subespaços X e Y , então X e Y são chamados *subespaços complementares*, ou também o subespaço Y é dito *complementar* do subespaço X , e vice-versa.



Se o espaço total V tem dimensão finita, então todo subespaço X admite um complementar Y , embora não univocamente determinado. De fato, se $e_1, e_2, \dots, e_m$ é uma base de X , e os vetores $e_{m+1}, e_{m+2}, \dots, e_n$ são escolhidos de maneira tal que a reunião $e_1, e_2, \dots, e_n$ forma uma base de V , então é fácil ver que o subespaço Y gerado pelos $e_{m+1}, e_{m+2}, \dots, e_n$ é um complementar de X .

Em geral, um subespaço $X \subset V$ é uma soma direta $X = \bigoplus_{k=1}^m X_k$ de uma família finita de subespaços X_k 's se todo vetor $v \in X$ é representado de uma única maneira como soma

$$v = v_1 + v_2 + \dots + v_m$$

de vetores $v_k \in X_k$.

e.g. Por exemplo, o plano $\mathbb{R}^2$ é, por definição, uma soma direta das retas $\mathbb{R}i$ e $\mathbb{R}j$. Também é uma soma direta das retas $\mathbb{R}i$ e $\mathbb{R}(i+j)$. De fato, um complementar da reta $\mathbb{R}i$ no plano é qualquer reta não horizontal, do gênero $\mathbb{R}(\alpha i + \beta j)$ com $\beta \neq 0$.

O espaço $\mathbb{R}^3$ é uma soma dos planos/subespaços $z = 0$ e $x = 0$, mas não é uma soma direta deles, pois $(x, y, z) = (x, y, 0) + (0, 0, z) = (x, 0, 0) + (0, y, z) = (x, y/2, 0) + (0, y/2, z) = \dots$. Por outro lado, $\mathbb{R}^3$ é uma soma direta do plano $z = 0$ e da reta $\mathbb{R}k$. Também é uma soma direta do plano $z = 0$ e da reta $\mathbb{R}(1, 2, 3) \dots$

e.g. Funções pares ou ímpares. Uma função $f : \mathbb{R} \rightarrow \mathbb{R}$ é dita *par* se $f(-t) = f(t)$ para todos $t \in \mathbb{R}$, e é dita *ímpar* se $f(-t) = -f(t)$ para todo $t \in \mathbb{R}$. Os conjuntos

$$\mathbb{R}_{\pm}^{\mathbb{R}} := \{ f : \mathbb{R} \rightarrow \mathbb{R} \quad \text{t.q.} \quad f(-t) = \pm f(t) \quad \forall t \}$$

das funções pares ou ímpares são subespaços do espaço vetorial real $\mathbb{R}^{\mathbb{R}}$ de todas as funções reais de uma variável real. A interseção é o subespaço trivial formado pela função identicamente nula. Toda função $f(t)$ é uma soma única de uma função par e uma função ímpar, pois $f(t) = f_+(t) + f_-(t)$ se

$$f_{\pm}(t) = \frac{1}{2}(f(t) \pm f(-t))$$

e é claro que $f_{\pm} \in \mathbb{R}_{\pm}^{\mathbb{R}}$. Portanto, o espaço das funções reais de uma variável real é uma soma direta $\mathbb{R}^{\mathbb{R}} = \mathbb{R}_{+}^{\mathbb{R}} \oplus \mathbb{R}_{-}^{\mathbb{R}}$.

Produtos. Também é possível definir de forma coerente a soma direta dois espaços lineares que não são necessariamente subespaços de um espaço dado. A *soma direta* (também chamada *biproduto*, ou *produto cartesiano*) dos espaços lineares $\mathbf{V}$ e $\mathbf{W}$, definidos sobre o mesmo corpo, é o espaço linear $\mathbf{V} \oplus \mathbf{W}$ cujos vetores são os pares ordenados $(\mathbf{v}, \mathbf{w}) \in \mathbf{V} \times \mathbf{W}$ de vetores $\mathbf{v} \in \mathbf{V}$ e $\mathbf{w} \in \mathbf{W}$, munido das operações

$$(\mathbf{v}, \mathbf{w}) + (\mathbf{v}', \mathbf{w}') := (\mathbf{v} + \mathbf{v}', \mathbf{w} + \mathbf{w}') \quad \text{e} \quad \lambda(\mathbf{v}, \mathbf{w}) := (\lambda\mathbf{v}, \lambda\mathbf{w})$$

O vetor nulo é $(\mathbf{0}, \mathbf{0})$ (o par ordenado formado pelos vetores nulos de cada um dos dois espaços). Então $\mathbf{v} \mapsto (\mathbf{v}, \mathbf{0})$ e $\mathbf{w} \mapsto (\mathbf{0}, \mathbf{w})$ definem isomorfismos de $\mathbf{V}$ e $\mathbf{W}$ em subespaços V e W de $\mathbf{V} \oplus \mathbf{W}$, respetivamente, e é claro que $\mathbf{V} \oplus \mathbf{W} = V \oplus W$.

e.g. Somas diretas e produtos cartesianos. Naturalmente, a soma direta $\mathbb{R} \oplus \mathbb{R}$ é o plano cartesiano $\mathbb{R}^2$, a soma direta $\mathbb{R}^2 \oplus \mathbb{R}$ é o espaço $\mathbb{R}^3$, ...

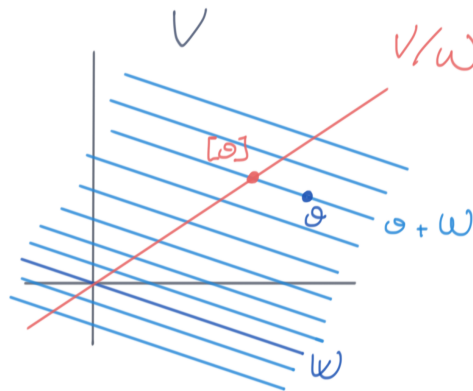
Espaço quociente e co-dimensão. Seja W um subespaço do espaço linear $\mathbf{V}$. O *espaço quociente* $\mathbf{V}/W$ é o conjunto das classes de equivalência pela seguinte relação de equivalência: dois vetores $\mathbf{v}$ e $\mathbf{v}'$ estão na mesma classe se diferem por um vetor de W , ou seja, se $\mathbf{v}' = \mathbf{v} + \mathbf{w}$ com $\mathbf{w} \in W$. Uma notação natural para a classe de $\mathbf{v}$ é

$$[\mathbf{v}] := \mathbf{v} + W$$

(a classe de $\mathbf{v}$ é o “plano paralelo a W que passa por $\mathbf{v}$ ”). O espaço quociente tem uma estrutura natural de espaço linear, se soma e produto por um escalar são definidos por

$$[\mathbf{v}] + [\mathbf{v}'] = [\mathbf{v} + \mathbf{v}'] \quad \text{e} \quad \lambda[\mathbf{v}] = [\lambda\mathbf{v}]$$

respetivamente (a prova que esta definição não depende dos representantes das classes é um exercício). A correspondência $\mathbf{v} \mapsto [\mathbf{v}]$ define uma transformação linear $\pi : \mathbf{V} \rightarrow \mathbf{V}/W$, chamada “projeção”.



Se o espaço quociente $\mathbf{V}/W$ tem dimensão finita, então a sua dimensão é chamada *co-dimensão* de W (enquanto subespaço de $\mathbf{V}$), e denotada por

$$\text{codim}(W) := \dim(\mathbf{V}/W)$$

Se $\mathbf{V}$ tem dimensão finita, então também W tem dimensão finita, e a dimensão do espaço total é uma soma

$$\dim(\mathbf{V}) = \dim(\mathbf{V}/W) + \dim(W)$$

assim que $\text{codim}(W) = \dim(\mathbf{V}) - \dim(W)$. De fato, se os $\mathbf{e}_1, \dots, \mathbf{e}_m$ formam uma base de W , então é possível completar o sistema acrescentando os vetores $\mathbf{e}_{m+1}, \mathbf{e}_{m+2}, \dots, \mathbf{e}_n$ até obter uma base de $\mathbf{V}$. É depois imediato verificar que as classes $[\mathbf{e}_{m+1}], [\mathbf{e}_{m+2}], \dots, [\mathbf{e}_n]$ formam uma base de $\mathbf{V}/W$.

Se o espaço linear é uma soma direta $\mathbf{V} = X \oplus Y$, então o espaço quociente $\mathbf{V}/Y$ é naturalmente isomorfo a X . De fato, cada vetor é uma soma única $\mathbf{v} = \mathbf{x} + \mathbf{y}$ de um vetor $\mathbf{x} \in X$ e um vetor $\mathbf{y} \in Y$, assim que a classe $[\mathbf{v}]$ pode ser identificada de maneira única com o vetor $\mathbf{x} \in X$, e vice-versa. Em particular,

$$\dim(X \oplus Y) = \dim(X) + \dim(Y)$$

8 Transformações lineares

ref: [Ap69] Vol 2, 2.1-8 ; [La87] Ch. IV-V

8.1 Formas lineares e espaço dual

3 nov 2022

Linearidade. Se cada kilo de ♡ custa A euros e cada kilo de ♠ custa B euros, então a kilos de ♡ e b kilos de ♠ custam $aA + bB$ euros. Ou seja, a função “preço” P satisfaz

$$P(a♡ + b♠) = a \cdot P(♡) + b \cdot P(♠)$$

Esta propriedade é chamada *linearidade*.

Por outro lado, a superfície e o volume de um cubo de lado 2ℓ são 4 e 8 vezes a superfície e o volume de um cubo de lado ℓ , respetivamente (e esta é uma das razões pela existência de dimensões típicas de animais e plantas, como explicado por D’Arcy Thompson²⁰). São funções não lineares.

ex: Dê exemplos de funções lineares e de funções não lineares.

Formas lineares. Seja $\mathbf{V}$ um espaço linear real (ou complexo). Uma função real $f : \mathbf{V} \rightarrow \mathbb{R}$ (ou complexa $f : \mathbf{V} \rightarrow \mathbb{C}$) é dita *aditiva* se $\forall \mathbf{v}, \mathbf{w} \in \mathbf{V}$

$$f(\mathbf{v} + \mathbf{w}) = f(\mathbf{v}) + f(\mathbf{w})$$

e é dita *homogénea* se $\forall \mathbf{v} \in \mathbf{V}$ e $\forall \lambda \in \mathbb{R}$ (ou $\forall \lambda \in \mathbb{C}$)

$$f(\lambda \mathbf{v}) = \lambda f(\mathbf{v})$$

Uma função real $\xi : \mathbf{V} \rightarrow \mathbb{R}$ (ou complexa $\xi : \mathbf{V} \rightarrow \mathbb{C}$) aditiva e homogénea, ou seja, tal que

$$\xi(\lambda \mathbf{v} + \mu \mathbf{w}) = \lambda \xi(\mathbf{v}) + \mu \xi(\mathbf{w})$$

é dita *forma linear*, ou *co-vetor* (ou *funcional linear* quando $\mathbf{V}$ é um espaço de funções). Uma notação simétrica para o valor da forma linear ξ sobre o vetor $\mathbf{v}$ é

$$\langle \xi, \mathbf{v} \rangle := \xi(\mathbf{v}).$$

Espaço dual. O espaço $\mathbf{V}^* := \text{Hom}_{\mathbb{R}}(\mathbf{V}, \mathbb{R})$ (ou $\text{Hom}_{\mathbb{C}}(\mathbf{V}, \mathbb{C})$) das formas lineares, dito *espaço dual (algébrico)* de $\mathbf{V}$, é um espaço linear real (ou complexo) se a adição e o produto por um escalar são definidos por

$$\langle \xi + \eta, \mathbf{v} \rangle := \langle \xi, \mathbf{v} \rangle + \langle \eta, \mathbf{v} \rangle \quad \text{e} \quad \langle \lambda \xi, \mathbf{v} \rangle := \langle \xi, \lambda \mathbf{v} \rangle = \lambda \langle \xi, \mathbf{v} \rangle$$

respetivamente, e a forma nula $\mathbf{0}^* \in \mathbf{V}^*$ é definida por $\langle \mathbf{0}^*, \mathbf{v} \rangle = 0$ para todo $\mathbf{v} \in \mathbf{V}$.

Espaço dual em dimensão finita. Se o espaço vetorial $\mathbf{V}$ tem dimensão finita, e se $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ é uma base (por exemplo, a base canónica de $\mathbf{V} \approx \mathbb{R}^n$), então uma forma linear ξ é determinada pelos seus valores $\xi_i := \langle \xi, \mathbf{e}_i \rangle$ nos vetores da base, pois se $\mathbf{v} = v_1 \mathbf{e}_1 + v_2 \mathbf{e}_2 + \dots + v_n \mathbf{e}_n$ então

$$\begin{aligned} \langle \xi, \mathbf{v} \rangle &= \langle \xi, v_1 \mathbf{e}_1 + v_2 \mathbf{e}_2 + \dots + v_n \mathbf{e}_n \rangle \\ &= v_1 \langle \xi, \mathbf{e}_1 \rangle + v_2 \langle \xi, \mathbf{e}_2 \rangle + \dots + v_n \langle \xi, \mathbf{e}_n \rangle = \xi_1 v_1 + \xi_2 v_2 + \dots + \xi_n v_n. \end{aligned}$$

Portanto, também o espaço dual $\mathbf{V}^*$ tem dimensão finita $\dim(\mathbf{V}^*) = \dim(\mathbf{V})$, e uma sua base, dita *base dual*, é o conjunto ordenado dos co-vetores $\mathbf{e}_1^*, \mathbf{e}_2^*, \dots, \mathbf{e}_n^*$ definidos por ²¹

$$\langle \mathbf{e}_i^*, \mathbf{e}_j \rangle = \delta_{ij}$$

²⁰D’Arcy Wentworth Thompson, *On growth and form*, 1917 and 1942.

²¹O *símbolo de Kronecker* é definido por

$$\delta_{ij} = \begin{cases} 1 & \text{se } i = j \\ 0 & \text{se } i \neq j \end{cases}$$

As coordenadas da forma linear $\xi = \xi_1 \mathbf{e}_1^* + \xi_2 \mathbf{e}_2^* + \cdots + \xi_n \mathbf{e}_n^*$ relativamente à base dual são os números $\xi_i = \langle \xi, \mathbf{e}_i \rangle$, assim que $\langle \xi, \mathbf{v} \rangle = \sum_i \xi_i v_i$. Observe que a k -ésima coordenada do vetor $\mathbf{v}$ relativamente à base $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ é precisamente $v_k = \langle \mathbf{e}_k^*, \mathbf{v} \rangle$.

Uma maneira conveniente de representar o valor $\langle \xi, \mathbf{x} \rangle$ da forma linear ξ sobre o vetor $\mathbf{x}$ é usando o “produto linhas por colunas”

$$\langle \xi, \mathbf{x} \rangle = \Xi X = \begin{pmatrix} \xi_1 & \xi_2 & \cdots & \xi_n \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$$

do vetor linha

$$\Xi = \begin{pmatrix} \xi_1 & \xi_2 & \cdots & \xi_n \end{pmatrix}$$

(que representa a forma linear ξ) pelo vetor coluna

$$X = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$$

que representa o vetor $\mathbf{x}$.

Dual do espaço dual. Cada vetor $\mathbf{v} \in \mathbf{V}$ define uma forma linear $\xi \mapsto \langle \xi, \mathbf{v} \rangle$ em $\mathbf{V}^*$, e portanto existe uma inclusão $\mathbf{V} \subset (\mathbf{V}^*)^*$. Se $\mathbf{V}$ tem dimensão finita, então todas as formas lineares $g \in (\mathbf{V}^*)^*$ são deste género, ou seja, podem ser representadas como $g(\xi) = \langle \xi, \mathbf{v} \rangle$ para algum $\mathbf{v} \in \mathbf{V}$. De fato, fixada uma base $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ de $\mathbf{V}$, e portanto a sua base dual $\mathbf{e}_1^*, \mathbf{e}_2^*, \dots, \mathbf{e}_n^*$ de $\mathbf{V}^*$, basta escolher $\mathbf{v} = v_1 \mathbf{e}_1 + v_2 \mathbf{e}_2 + \cdots + v_n \mathbf{e}_n$ com coeficientes $v_i = g(\mathbf{e}_i^*)$. Portanto, o espaço dual do espaço dual é isomorfo ao próprio espaço, ou seja, $(\mathbf{V}^*)^* \approx \mathbf{V}$. É útil também observar que a correspondência entre bases duais é simétrica: a cada base de $\mathbf{V}^*$ corresponde uma base dual de $\mathbf{V}$, e a primeira é também a base dual da segunda.

e.g. Projeções coordenadas. As projeções $\pi_k : \mathbb{R}^n \rightarrow \mathbb{R}$, com $k = 1, 2, \dots, n$, definidas por

$$\pi_k(x_1, x_2, \dots, x_n) = x_k$$

são formas lineares. São precisamente as formas $\mathbf{e}_k^*$ da base dual da base canónica.

ex: Mostre que uma função homogénea $f : \mathbb{R} \rightarrow \mathbb{R}$ é necessariamente uma homotetia $f(x) = \lambda x$, com $\lambda = f(1) \in \mathbb{R}$. Deduza que uma função homogénea $f : \mathbb{R} \rightarrow \mathbb{R}$ é também aditiva, e portanto linear.

ex: Diga se as seguintes funções $f : \mathbb{R}^n \rightarrow \mathbb{R}$ são ou não lineares:

$$\begin{aligned} x &\mapsto 3x & x &\mapsto 2x - 1 & x &\mapsto \sin(2\pi x) & (x, y) &\mapsto 3x - 5y & (x, y) &\mapsto x^2 - xy \\ (x, y, z) &\mapsto 2x - y + 3z & (x, y, z) &\mapsto 2x - y + 3z + 8 & (x, y, z) &\mapsto 0 & (x, y, z) &\mapsto \sqrt{3} \end{aligned}$$

ex: Seja $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ uma forma linear tal que $f(\mathbf{i}) = 5$ e $f(\mathbf{j}) = -2$. Determine $f(x, y)$.

ex: Seja $f : \mathbb{R}^3 \rightarrow \mathbb{R}$ uma forma linear tal que $f(\mathbf{i}) = -3$, $f(\mathbf{j}) = 1$ e $f(\mathbf{k}) = 7$. Determine $f(x, y, z)$.

Formas lineares e produto escalar no espaço euclidiano $\mathbb{R}^n$. No espaço euclidiano $\mathbb{R}^n$, munido do produto escalar (2.1), há uma relação simples entre formas e vetores. Todo vetor $\mathbf{y} \in \mathbb{R}^n$ define uma forma linear $\mathbf{y}^\# \in (\mathbb{R}^n)^*$ de acordo com

$$\mathbf{x} \mapsto \langle \mathbf{y}^\#, \mathbf{x} \rangle := \mathbf{y} \cdot \mathbf{x} \quad (8.1)$$

(a linearidade é consequência da linearidade do produto escalar).

Vice-versa, toda forma linear em $\mathbb{R}^n$ é deste género. Para provar isto, convém escolher uma base ortonormada, por exemplo a base canónica $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$. Então é claro que uma forma ξ , com coordenadas $\xi_k = \langle \xi, \mathbf{e}_k \rangle$ relativamente à base dual, é igual a $\xi = \mathbf{y}^\#$ se definimos o vetor $\mathbf{y} = \xi_1 \mathbf{e}_1 + \xi_2 \mathbf{e}_2 + \dots + \xi_n \mathbf{e}_n$. Assim, a correspondência $\mathbf{y} \mapsto \mathbf{y}^\#$ é um isomorfismo $\mathbb{R}^n \approx (\mathbb{R}^n)^*$ entre o espaço euclidiano $\mathbb{R}^n$ e o seu dual $(\mathbb{R}^n)^*$ (que depende da estrutura euclidiana, ou seja, do produto escalar euclidiano). O isomorfismo inverso $(\mathbb{R}^n)^* \approx \mathbb{R}^n$ pode ser denotado por $\xi \mapsto \xi^\flat$.

ex: Determine o vetor $\mathbf{y} = (\xi_1, \xi_2, \dots, \xi_n) \in \mathbb{R}^n$ que define as seguintes formas lineares $\mathbf{y}^\#$ no espaço euclidiano $\mathbb{R}^n$:

$$x \mapsto 3x \quad (x, y) \mapsto 0 \quad (x, y) \mapsto 5x + 9y$$

$$(x, y, z) \mapsto -3x + 7y - z \quad (x_1, x_2, \dots, x_n) \mapsto 3x_k \quad \text{com } 0 \leq k \leq n$$

8.2 Núcleo e hiperplanos

Núcleo e hiperplanos. Seja $\mathbf{V}$ um espaço linear. O *núcleo/espço nulo* (em inglês *kernel*) da forma linear $\xi \in \mathbf{V}^*$ é

$$\text{Ker}(\xi) := \{\mathbf{x} \in \mathbf{V} \text{ t.q. } \langle \xi, \mathbf{x} \rangle = 0\}$$

É claro que o núcleo é um subespaço vetorial de $\mathbf{V}$. O núcleo da forma nula é, naturalmente, o próprio espaço $\mathbf{V}$.

Se ξ não é a forma identicamente nula, então o seu núcleo é um subespaço próprio de $\mathbf{V}$, pois existe pelo menos um vetor $\mathbf{v}$ onde $\langle \xi, \mathbf{v} \rangle = \alpha \neq 0$. Consequentemente, existe um vetor $\mathbf{v}_1$ tal que $\langle \xi, \mathbf{v}_1 \rangle = 1$ (basta escolher $\mathbf{v}_1 = \alpha^{-1} \mathbf{v}$). Vice-versa, dado um vetor não nulo $\mathbf{v}_1$ num espaço linear de dimensão finita $\mathbf{V}$, existe uma forma $\xi \in \mathbf{V}^*$ tal que $\langle \xi, \mathbf{v}_1 \rangle = 1$ (basta escolher uma base de $\mathbf{V}$ contendo o vetor $\mathbf{v}_1$, e depois o vetor $\mathbf{v}_1^*$ da base dual).

Se a forma ξ não é nula, e se $\mathbf{v}_1$ é um vetor tal que $\langle \xi, \mathbf{v}_1 \rangle = 1$, então cada vetor $\mathbf{x} \in \mathbf{V}$ pode ser representado de uma única maneira como soma

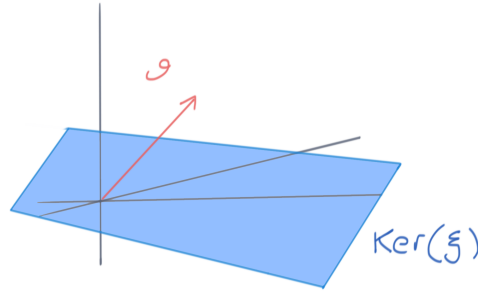
$$\mathbf{x} = \lambda \mathbf{v}_1 + \mathbf{w}$$

de um vetor proporcional a $\mathbf{v}$ e de um vetor $\mathbf{w} \in \text{Ker}(\xi)$. De fato, a condição $\mathbf{x} - \lambda \mathbf{v}_1 \in \text{Ker}(\xi)$, ou seja, $\langle \xi, \mathbf{x} - \lambda \mathbf{v}_1 \rangle = 0$, obriga a escolher o valor $\lambda = \langle \xi, \mathbf{x} \rangle$. Portanto,

Teorema 8.1. *Se ξ é uma forma linear não nula no espaço vetorial $\mathbf{V}$, e se $\mathbf{v}_1$ é um vetor tal que $\langle \xi, \mathbf{v}_1 \rangle = 1$, então o espaço total é uma soma direta*

$$\mathbf{V} = \text{Ker}(\xi) \oplus \mathbb{R}\mathbf{v}_1$$

do núcleo de ξ e da reta gerada pelo vetor $\mathbf{v}_1$.



O núcleo $\text{Ker}(\xi)$ de uma forma não nula ξ é chamado *hiperplano* do espaço linear $\mathbf{V}$, ou seja, subespaço de “co-dimensão” 1 (falta apenas uma reta para reconstruir o espaço total). Em particular, se $\mathbf{V}$ tem dimensão finita, então também $\text{Ker}(\xi)$ tem dimensão finita e

$$\dim(\text{Ker}(\xi)) + 1 = \dim \mathbf{V}$$

O *hiperplano afim* que passa pelo ponto $\mathbf{a} \in \mathbf{V}$ e é paralelo ao hiperplano $\text{Ker}(\xi)$ é

$$\mathbf{a} + \text{Ker}(\xi) = \{\mathbf{v} \in \mathbf{V} \text{ t.q. } \langle \xi, \mathbf{v} \rangle = \lambda\} \quad \text{onde } \lambda = \langle \xi, \mathbf{a} \rangle.$$

No espaço $\mathbb{R}^n$ euclidiano, onde o produto escalar define um isomorfismo (8.1) entre vetores e formas, o núcleo de uma forma ξ é simplesmente o hiperplano $\text{Ker}(\xi) = \mathbf{y}^\perp$ ortogonal ao vetor $\mathbf{y} = \xi^b$.

e.g. Hiperplanos coordenados. O núcleo da projeção $\pi_k : \mathbb{R}^n \rightarrow \mathbb{R}$, definida por $\pi_k(x_1, x_2, \dots, x_n) = x_k$, é o hiperplano $x_k = 0$.

ex: Mostre que o núcleo de uma forma linear $\xi \in \mathbf{V}^*$ é um subespaço linear de $\mathbf{V}$.

Equações lineares homogêneas. Uma “equação linear homogênea”

$$a_1x_1 + a_2x_2 + \dots + a_nx_n = 0$$

não trivial, ou seja, com pelo menos um coeficiente $a_k \neq 0$, define um hiperplano de $\mathbb{R}^n$, um subespaço vetorial de dimensão $n - 1$. É o núcleo da forma $\xi = a_1\mathbf{e}_1^* + a_2\mathbf{e}_2^* + \dots + a_n\mathbf{e}_n^*$, ou seja, no isomorfismo entre $(\mathbb{R}^n)^*$ e $\mathbb{R}^n$ definido pela estrutura euclidiana, o hiperplano $\mathbf{a}^\perp$ ortogonal ao vetor $\mathbf{a} = (a_1, a_2, \dots, a_n)$. Uma “equação linear”

$$a_1x_1 + a_2x_2 + \dots + a_nx_n = b$$

define um hiperplano afim (que não passa pela origem se $b \neq 0$).

Integral. Em dimensão infinita, o espaço dual deixa de ser isomorfo ao próprio espaço. No entanto, formas lineares típicas são objetos básicos da análise matemática. Por exemplo, o integral

$$I(f) := \int_a^b f(t) dt$$

é uma forma linear no espaço $\mathcal{C}^0([a, b])$ das funções contínuas no intervalo $[a, b] \subset \mathbb{R}$. O seu núcleo é o hiperplano das funções com média nula no intervalo. Toda função contínua pode ser representada de forma única como soma $f(x) = c + g(x)$ onde $c = I(f)$ é uma constante (a média de f vezes o comprimento do intervalo) e $g(x) = f(x) - I(f)$ é uma função com média nula.

Dada uma função $h(t)$, por exemplo contínua, também é possível definir a forma linear

$$I_h(f) := \int_a^b f(t)h(t) dt$$

uma média pesada dos valores de $f(t)$.

Delta de Dirac. Outra forma linear importante na análise das EDPs da física-matemática foi “inventada” pelo físico Paul Dirac, e formalizada pelo matemático Laurent Schwartz. A *delta de Dirac* (no ponto 0) é definida por

$$\delta(f) := “ \int_{-\infty}^{\infty} \delta(t)f(t) dt ” = f(0)$$

É uma forma linear no espaço $\mathcal{C}^0(\mathbb{R})$ das funções contínuas $f: \mathbb{R} \rightarrow \mathbb{R}$. Se pensarmos, informalmente, que os valores $f(t)$ são as “coordenadas” do vetor f , uma para cada t , então a delta de Dirac é o vetor da base dual que corresponde a coordenada $t = 0$. O núcleo de δ é o conjunto das funções (contínuas) tais que $f(0) = 0$, que é um hiperplano do espaço $\mathcal{C}^0(\mathbb{R})$.

Interseções de hiperplanos e sistemas homogêneos. Sejam $\xi_1, \xi_2, \dots, \xi_m$ formas lineares em $\mathbb{R}^n$ (ou, em geral, num espaço linear de dimensão finita). Então a interseção dos núcleos

$$W = \bigcap_{k=1}^m \text{Ker}(\xi_k)$$

é um subespaço vetorial de $\mathbb{R}^n$. Se $\xi_1 = (a_{11}, a_{12}, \dots, a_{1n})$, $\xi_2 = (a_{21}, a_{22}, \dots, a_{2n})$, $\dots$, $\xi_m = (a_{m1}, a_{m2}, \dots, a_{mn})$ denotam as coordenadas das ξ_k 's na base dual da base canônica de $\mathbb{R}^n$, então os vetores de W são as soluções $\mathbf{x} = (x_1, x_2, \dots, x_n)$ do “sistema linear homogêneo”

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n = 0 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n = 0 \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n = 0 \end{cases} \quad (8.2)$$

Teorema 8.2. Se as formas lineares $\xi_1, \xi_2, \dots, \xi_m \in (\mathbb{R}^n)^*$ são linearmente independentes (e portanto $m \leq n$), então a interseção $W = \bigcap_{k=1}^m \text{Ker}(\xi_k)$ é um subespaço vetorial de co-dimensão m , e portanto de dimensão $n - m$.

Demonstração. De fato, se completamos o sistema de formas independentes até obter uma base $\xi_1, \xi_2, \dots, \xi_m, \xi_{m+1}, \dots, \xi_n$ de $(\mathbb{R}^n)^*$, e se $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ denota a base dual de $\mathbb{R}^n$ (definida pelas condições $\langle \xi_i, \mathbf{b}_j \rangle = \delta_{ij}$), então nas coordenadas relativas a esta base o sistema homogêneo é

$$x_1 = 0, \quad x_2 = 0, \quad \dots \quad x_m = 0,$$

e as soluções são todos os vetores do gênero $\mathbf{x} = x_{m+1}\mathbf{b}_{m+1} + x_{m+2}\mathbf{b}_{m+2} + \dots + x_n\mathbf{b}_n$. Portanto, $W \approx \mathbb{R}^{n-m}$, e uma sua base é formada pelos vetores $\mathbf{b}_{m+1}, \mathbf{b}_{m+2}, \dots, \mathbf{b}_n$. $\square$

Naturalmente, se uma das formas lineares que definem o sistema homogêneo, por exemplo ξ_m , é uma combinação linear das outras, então a equação que define, neste caso m -ésima, pode ser retirada do sistema homogêneo (8.2) sem alterar o espaço das soluções. De fato, se $\xi_m = b_1\xi_1 + b_2\xi_2 + \dots + b_{m-1}\xi_{m-1}$, então o núcleo de ξ_m contém a interseção dos núcleos das outras formas lineares $\xi_1, \xi_2, \dots, \xi_{m-1}$.

ex: Determine o espaço das soluções dos seguintes sistemas lineares homogêneos:

$$\begin{cases} 3x + 2y = 0 \\ 5x - y = 0 \end{cases} \quad \begin{cases} x - 2y = 0 \\ -3x + 6y = 0 \end{cases}$$

$$\begin{cases} 3x + 2y + z = 0 \\ 5x + 3y + 3z = 0 \\ -x + y + z = 0 \end{cases} \quad \begin{cases} x + 2y + z = 0 \\ 2x - y - 6z = 0 \\ -x + 3y + 7z = 0 \end{cases}$$

Aniquilador. Sejam $\mathbf{V}$ um espaço vetorial e $\mathbf{V}^*$ o seu dual algébrico. O *aniquilador* (em inglês *annihilator*, do latim NIHIL=nada) do subconjunto $S \subset \mathbf{V}$ é

$$\text{Annih}(S) := \{\xi \in \mathbf{V}^* \text{ t.w. } \langle \xi, \mathbf{v} \rangle = 0 \text{ para todo } \mathbf{v} \in S\}$$

ou seja, o conjunto das formas que se anulam em todos os vetores de S . Outra notação utilizada para o aniquilador é S^0 .

É imediato verificar que o aniquilador de um conjunto arbitrário $S \subset \mathbf{V}$ (finito ou não) é um subespaço linear de $\mathbf{V}^*$, e é claro que $\text{Annih}(S) = \text{Annih}(\text{Span}(S))$. O aniquilador do conjunto vazio é $\text{Annih}(\emptyset) = \mathbf{V}^*$, e o aniquilador do espaço total $\mathbf{V}$ é o subespaço trivial formado por apenas a forma nula. Se $S \subset S' \subset \mathbf{V}$, então é evidente que $\text{Annih}(S') \subset \text{Annih}(S)$.

No espaço euclidiano $\mathbb{R}^n$, onde o produto escalar define um isomorfismo natural $(\mathbb{R}^n)^* \approx \mathbb{R}^n$, o aniquilador de um subconjunto S é o subespaço ortogonal $S^\perp$.

Também é possível definir o aniquilador de um subconjunto $T \subset \mathbf{V}^*$ como

$$\text{Annih}(T) := \{\mathbf{v} \in \mathbf{V} \text{ t.w. } \langle \xi, \mathbf{v} \rangle = 0 \text{ para todo } \xi \in T\}$$

(e não como um subespaço do dual do espaço dual!) Desta forma, o aniquilador do conjunto $\{\xi\} \subset \mathbf{V}^*$ formado por apenas um vetor do espaço dual $\mathbf{V}^*$ coincide com o núcleo $\text{Ker}(\xi)$ da forma ξ .

Nesta linguagem, o teorema 8.2 é traduzido no seguinte.

Teorema 8.3. Sejam $\mathbf{V}$ um espaço linear de dimensão finita e $W \subset \mathbf{V}$ um subespaço linear. Então

$$\dim(W) + \dim(\text{Annih}(W)) = \dim(\mathbf{V})$$

Demonstração. Basta observar que se $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ é uma base de $\mathbf{V}$ tal que $W = \text{Span}(\mathbf{b}_1, \dots, \mathbf{b}_k)$, então $\text{Annih}(W)$ é o subespaço de $\mathbf{V}^*$ gerado pelos vetores $\mathbf{b}_{k+1}^*, \dots, \mathbf{b}_n^*$ da base dual. $\square$

ex: Mostre que $\text{Annih}(S) + \text{Annih}(T) \subset \text{Annih}(S \cap T)$ e $\text{Annih}(S \cup T) = \text{Annih}(S) \cap \text{Annih}(T)$.

ex: Deduza que se $V, W \subset \mathbf{V}$ são subespaços então $\text{Annih}(V + W) = \text{Annih}(V) \cap \text{Annih}(W)$.

ex: Mostre que se X é um subespaço do espaço linear de dimensão finita $\mathbf{V}$ então $\text{Annih}(\text{Annih}(X)) = X$.

Wild additive functions in the real line. For real valued functions of a real variable $f : \mathbb{R} \rightarrow \mathbb{R}$, homogeneity implies additivity, since $x + y = (1 + y/x)x$ (if $x \neq 0$, of course), and therefore an homogeneous function satisfies

$$f(x + y) = f((1 + y/x)x) = (1 + y/x)f(x) = f(x) + (y/x)f(x) = f(x) + f(y).$$

Surprisingly, there exist additive functions which are not homogeneous (hence not linear), at least if we accept the axiom of choice. Indeed, additivity only implies linearity on “rational lines” $\mathbb{Q}x \subset \mathbb{R}$, i.e.

$$f(rx) = rf(x) \quad \forall r = p/q \in \mathbb{Q} \text{ and } \forall x \in \mathbb{R}.$$

Therefore, if we could choose a different slope $\lambda_{\mathbb{Q}x}$, hence a different homogeneous function $rx \mapsto \lambda_{\mathbb{Q}x}rx$, for any orbit $\mathbb{Q}x$ of the quotient space $\mathbb{Q} \backslash \mathbb{R}$, the resulting function $f : \mathbb{R} \rightarrow \mathbb{R}$ would be additive but not homogeneous. Any such wild additive but not homogeneous function $f : \mathbb{R} \rightarrow \mathbb{R}$ cannot be continuous, and indeed has a dense graph.

Let $f : \mathbb{R} \rightarrow \mathbb{R}$ be an homogeneous function, and $x \in \mathbb{R}$. For $r = p/q \in \mathbb{Q}$ with $p, q \in \mathbb{Z}$ and $q \neq 0$, show that $f(rx) = f(px/q) = pf(x/q)$ and also $f(x) = f(qx/q) = qf(x/q)$. Deduce that $f(rx) = rf(x)$ for all $r \in \mathbb{Q}$. Let $f : \mathbb{R} \rightarrow \mathbb{R}$ be an additive but not homogeneous function,

so that there exists two points $a, b \in \mathbb{R}$ where $f(a)/a \neq f(b)/b$. This implies that $\mathbf{v} = (a, f(a))$ and $\mathbf{w} = (b, f(b))$ are linearly independent vectors, hence a basis, of $\mathbb{R}^2$. From $\mathbb{R}\mathbf{v} + \mathbb{R}\mathbf{w} = \mathbb{R}^2$, deduce that the “rational plane” $\mathbb{Q}\mathbf{v} + \mathbb{Q}\mathbf{w}$ is dense in $\mathbb{R}^2$, i.e. any point $\mathbf{r} = (x, y) \in \mathbb{R}^2$ may be approximated with arbitrary precision by a point in $\mathbb{Q}\mathbf{v} + \mathbb{Q}\mathbf{w}$, i.e. for any $\mathbf{r} = (x, y) \in \mathbb{R}^2$ and any $\varepsilon > 0$ there exist rationals $\lambda, \lambda' \in \mathbb{Q}$ such that $\|\mathbf{r} - (\lambda\mathbf{v} + \lambda'\mathbf{w})\| < \varepsilon$. Use again the additivity of f to show that this implies the existence of a point $c \in \mathbb{R}$ such that $\|\mathbf{r} - (c, f(c))\| < \varepsilon$.

8.3 Transformações lineares

Homomorfismos/transformações lineares. Uma *transformação/aplicação linear*, ou *homomorfismo*, entre os espaços vetoriais $\mathbf{V}$ e $\mathbf{W}$ (reais ou complexos, os dois definidos sobre o mesmo corpo) é uma função $L : \mathbf{V} \rightarrow \mathbf{W}$ *aditiva* e *homogênea*, ou seja, tal que

$$L(\mathbf{v} + \mathbf{v}') = L(\mathbf{v}) + L(\mathbf{v}') \quad \text{e} \quad L(\lambda\mathbf{v}) = \lambda L(\mathbf{v})$$

respetivamente, para todos $\mathbf{v}, \mathbf{v}' \in \mathbf{V}$ e todo escalar λ . Combinando as duas propriedades, uma transformação linear é também caracterizada pela propriedade

$$L(\lambda\mathbf{v} + \lambda'\mathbf{v}') = \lambda L(\mathbf{v}) + \lambda' L(\mathbf{v}')$$

Por indução, a linearidade implica que a imagem de uma combinação linear (finita) dos vetores $\mathbf{v}_k$'s com coeficientes λ_k 's é uma combinação linear

$$L(\lambda_1\mathbf{v}_1 + \lambda_2\mathbf{v}_2 + \dots + \lambda_m\mathbf{v}_m) = \lambda_1 L(\mathbf{v}_1) + \lambda_2 L(\mathbf{v}_2) + \dots + \lambda_m L(\mathbf{v}_m)$$

das imagens $L(\mathbf{v}_k)$'s com coeficientes λ_k 's. A homogeneidade também implica que a imagem do vetor nulo é o vetor nulo. É usual omitir as parêntesis, e denotar a imagem $L(\mathbf{v})$ do vetor $\mathbf{v}$ simplesmente por $L\mathbf{v}$.

Um exemplo trivial é a *transformação nula* $0 : \mathbf{V} \rightarrow \mathbf{W}$, que envia todo vetor no vetor nulo, ou seja, $0(\mathbf{v}) = \mathbf{0}$.

O espaço linear das transformações lineares. O espaço $\text{Lin}(\mathbf{V}, \mathbf{W})$ das transformações lineares de $\mathbf{V}$ em $\mathbf{W}$ é um espaço linear (real ou complexo) se a adição e a multiplicação por um escalar são definidas por

$$(L + M)(\mathbf{v}) := L(\mathbf{v}) + M(\mathbf{v}) \quad (\lambda L)(\mathbf{v}) := \lambda L(\mathbf{v})$$

O elemento neutro é a transformação nula, que envia todo vetor no vetor nulo. O oposto da transformação linear L é a transformação $-L$, que envia $(-L)(\mathbf{v}) = -L(\mathbf{v})$.

Endomorfismos/operadores. Em particular, é um espaço linear o espaço $\text{End}(\mathbf{V}) := \text{Lin}(\mathbf{V}, \mathbf{V})$ dos *endomorfismos* de $\mathbf{V}$, as transformações lineares $L : \mathbf{V} \rightarrow \mathbf{V}$ de um espaço linear em si próprio. Os endomorfismos de um espaço vetorial são também chamados *operadores*, em particular em dimensão infinita, quando o espaço é um espaço de funções.

Um exemplo elementar mas básico é a *transformação identidade* $I_{\mathbf{V}} : \mathbf{V} \rightarrow \mathbf{V}$ (denotada simplesmente por I ou também 1 quando o espaço $\mathbf{V}$ é claro pelo contexto), que envia todo vetor em si próprio, ou seja, $I(\mathbf{v}) = \mathbf{v}$. Outro exemplo importante é, se λ é um escalar, a *homotetia* $\lambda I_{\mathbf{V}} : \mathbf{V} \rightarrow \mathbf{V}$, também denotada simplesmente por λ , que envia cada vetor $\mathbf{v}$ no seu múltiplo $\lambda\mathbf{v}$.

Se o espaço linear é uma soma direta $\mathbf{V} = X \oplus Y$ de dois subespaços, e $L \in \text{End}(X)$ e $M \in \text{End}(Y)$ são dois endomorfismos, é definido o endomorfismo “soma direta” $L \oplus M$ de acordo com

$$(L \oplus M)(\mathbf{x} + \mathbf{y}) := L(\mathbf{x}) + M(\mathbf{y})$$

ex: Se $L : \mathbf{V} \rightarrow \mathbf{W}$ é uma transformação linear, então $L(\mathbf{0}) = \mathbf{0}$ e $L(-\mathbf{v}) = -L(\mathbf{v})$.

ex: Uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ envia retas afins $\mathbf{a} + \mathbb{R}\mathbf{v} \subset \mathbf{V}$ em retas afins $\mathbf{b} + \mathbb{R}\mathbf{w} \subset \mathbf{W}$, se $\mathbf{b} = L(\mathbf{a})$ e $\mathbf{w} = L(\mathbf{v}) \neq \mathbf{0}$, ou em pontos $\mathbf{b}$, se $L(\mathbf{v}) = \mathbf{0}$. Em particular, envia retas $\mathbb{R}\mathbf{v} \subset \mathbf{V}$ passando pela origem em retas $\mathbb{R}\mathbf{w} \subset \mathbf{W}$ passando pela origem ou na própria origem.

ex: Uma transformação $T : \mathbb{R}^n \rightarrow \mathbb{R}^m$, definida por

$$(x_1, x_2, \dots, x_n) \mapsto (T_1(x_1, x_2, \dots, x_n), T_2(x_1, x_2, \dots, x_n), \dots, T_m(x_1, x_2, \dots, x_n))$$

é linear sse todas as suas “coordenadas” $T_i : \mathbb{R}^n \rightarrow \mathbb{R}$, com $i = 1, 2, \dots, m$, são lineares.

ex: Diga se as seguintes aplicações de $\mathbb{R}^n$ em $\mathbb{R}^m$ são lineares.

$$\begin{aligned} (x, y) &\mapsto (3x - 5y, x - y) & (x, y) &\mapsto (x^2, xy) \\ (x, y, z) &\mapsto (x, y + z, 2) & (x, y, z) &\mapsto (x, y + z, 0) \\ (x, y, z) &\mapsto (x - y, z + y, 3x + 2y - z) & (x, y, z) &\mapsto (1, 2, 3) \\ (x_1, x_2, \dots, x_n) &\mapsto (0, 0, 1) & (x_1, x_2, \dots, x_n) &\mapsto (0, 0, \dots, 0) \in \mathbb{R}^m \end{aligned}$$

ex: Diga se as seguintes aplicações de $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ são lineares.

T transforma cada ponto no seu simétrico em relação à reta $y = 0$

T transforma cada ponto no seu simétrico em relação à reta $y = x$

T transforma o ponto de coordenadas polares (r, θ) no ponto de coordenadas polares $(2r, \theta)$

T transforma o ponto de coordenadas polares (r, θ) no ponto de coordenadas polares $(r, \theta + \pi/2)$

ex: Seja $L : \mathbf{V} \rightarrow \mathbf{W}$ uma transformação linear. Se $V \subset \mathbf{V}$ é um subespaço linear de $\mathbf{V}$ e então a imagem $L(V)$ é um subespaço linear de $\mathbf{W}$. Se $W \subset \mathbf{W}$ é um subespaço linear de $\mathbf{W}$ e então a imagem inversa $L^{-1}(W)$ é um subespaço linear de $\mathbf{V}$.

ex: As translações de $\mathbb{R}^n$, as transformações $\mathbf{v} \mapsto \mathbf{v} + \mathbf{a}$ com $\mathbf{a} \in \mathbb{R}^n$, são transformações lineares?

ex: As homotetias de $\mathbb{R}^n$, as transformações $\mathbf{v} \mapsto \lambda \mathbf{v}$ com $\lambda \in \mathbb{R}$, são transformações lineares?

ex: [Ap69] Vol 2 2.4.

Transformações lineares determinadas pelos valores numa base. Se $\mathbf{V}$ tem dimensão finita, e $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ é uma sua base, então uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ é determinada pelos seus valores $\mathbf{w}_k = L(\mathbf{b}_k)$ sobre os vetores da base. De fato, um vetor arbitrário $\mathbf{v} = v_1 \mathbf{b}_1 + v_2 \mathbf{b}_2 + \dots + v_n \mathbf{b}_n$ tem imagem

$$\begin{aligned} L(\mathbf{v}) &= L(v_1 \mathbf{b}_1 + v_2 \mathbf{b}_2 + \dots + v_n \mathbf{b}_n) \\ &= v_1 L(\mathbf{b}_1) + v_2 L(\mathbf{b}_2) + \dots + v_n L(\mathbf{b}_n) \\ &= v_1 \mathbf{w}_1 + v_2 \mathbf{w}_2 + \dots + v_n \mathbf{w}_n \end{aligned}$$

Por exemplo, uma transformação linear $T : \mathbb{R}^n \rightarrow \mathbb{R}^m$ é determinada pelos n vetores $\mathbf{w}_k = T(\mathbf{e}_k) \in \mathbb{R}^m$, as imagens dos vetores $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ da base canônica de $\mathbb{R}^n$. O seu valor sobre um vetor genérico $\mathbf{x} = (x_1, x_2, \dots, x_n)$ é

$$T(x_1, x_2, \dots, x_n) = x_1 \mathbf{w}_1 + x_2 \mathbf{w}_2 + \dots + x_n \mathbf{w}_n$$

e.g. Por exemplo, se a transformação linear $T : \mathbb{R}^2 \rightarrow \mathbb{R}^3$ envia $T(1, 0) = (2, 3, 4)$ e $T(0, 1) = (5, 6, 7)$, então a imagem do vetor genérico $(x, y) = x(1, 0) + y(0, 1)$ é

$$T(x, y) = xT(1, 0) + yT(0, 1) = x(2, 3, 4) + y(5, 6, 7) = (2x + 5y, 3x + 6y, 4x + 7y)$$

ex: Seja $L : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ uma transformação linear tal que $L(1, 1) = (1, 4)$ e $L(2, -1) = (-2, 3)$. Determine $L(5, -1)$ (observe que $1 + 2 \cdot 2 = 5$ e $1 + 2 \cdot (-1) = -1$).

ex: Determine a transformação linear $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ (ou seja, a imagem $T(x, y)$ de um vetor genérico) tal que

$$T(\mathbf{i}) = 2\mathbf{i} + \mathbf{j} \quad \text{e} \quad T(\mathbf{j}) = \mathbf{i} - 3\mathbf{j}$$

ex: Determine a transformação linear $T : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ tal que

$$T(\mathbf{i} + \mathbf{j}) = 3\mathbf{i} \quad \text{e} \quad T(\mathbf{j} + \mathbf{k}) = 2\mathbf{j} - \mathbf{k} \quad \text{e} \quad T(\mathbf{k} + \mathbf{i}) = 5\mathbf{i} + \mathbf{j}$$

8.4 Núcleo e imagem

Núcleo e imagem. Seja $L : \mathbf{V} \rightarrow \mathbf{W}$ uma transformação linear entre os espaços lineares $\mathbf{V}$ e $\mathbf{W}$. O *núcleo/espaco nulo* (em inglês, *kernel*) de L é

$$\text{Ker}(L) := \{\mathbf{v} \in \mathbf{V} \text{ t.q. } L(\mathbf{v}) = \mathbf{0}\} \subset \mathbf{V}$$

É imediato verificar que $\text{Ker}(L)$ é um subespaço vetorial de $\mathbf{V}$. De fato, se $L(\mathbf{v}) = L(\mathbf{v}') = \mathbf{0}$ e λ é um escalar, então também $L(\lambda\mathbf{v}) = \lambda L(\mathbf{v}) = \mathbf{0}$ e $L(\mathbf{v} + \mathbf{v}') = L(\mathbf{v}) + L(\mathbf{v}') = \mathbf{0}$. O núcleo é uma medida da falta de injetividade, de acordo com o seguinte

Teorema 8.4. *Uma transformação linear é injetiva sse o seu núcleo é trivial.*

Demonstração. Se $L(\mathbf{v}) = L(\mathbf{v}')$ para dois vetores distintos $\mathbf{v} \neq \mathbf{v}'$, então a diferença $\mathbf{v} - \mathbf{v}'$ é um vetor não trivial de $\text{Ker}(L)$, pois $L(\mathbf{v} - \mathbf{v}') = L(\mathbf{v}) - L(\mathbf{v}') = \mathbf{0}$. Vice-versa, se $\mathbf{v}$ é um vetor não nulo de $\text{Ker}(L)$ e $\mathbf{v}'$ é um vetor arbitrário, então $L(\mathbf{v}' + \mathbf{v}) = L(\mathbf{v}') + L(\mathbf{v}) = L(\mathbf{v}')$, e portanto $\mathbf{v}'$ e $\mathbf{v}' + \mathbf{v}$ são dois vetores distintos com a mesma imagem. $\square$

Em particular,

Teorema 8.5. *Uma transformação linear injetiva preserva a independência, ou seja, envia um conjunto independente de vetores em um conjunto independente de vetores.*

Demonstração. Sejam $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m \in \mathbf{V}$, e sejam $\mathbf{w}_1 = L(\mathbf{v}_1), \mathbf{w}_2 = L(\mathbf{v}_2), \dots, \mathbf{w}_m = L(\mathbf{v}_m)$ as suas imagens pela transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$. Se os $\mathbf{w}_k$'s são dependentes, então existem coeficientes t_k 's não todos nulos tais que $t_1\mathbf{w}_1 + t_2\mathbf{w}_2 + \dots + t_m\mathbf{w}_m = \mathbf{0}$. Pela linearidade

$$L(t_1\mathbf{v}_1 + t_2\mathbf{v}_2 + \dots + t_m\mathbf{v}_m) = t_1\mathbf{w}_1 + t_2\mathbf{w}_2 + \dots + t_m\mathbf{w}_m = \mathbf{0}$$

Se L é injetiva, então $t_1\mathbf{v}_1 + t_2\mathbf{v}_2 + \dots + t_m\mathbf{v}_m = \mathbf{0}$ e portanto os $\mathbf{v}_k$'s também são dependentes. $\square$

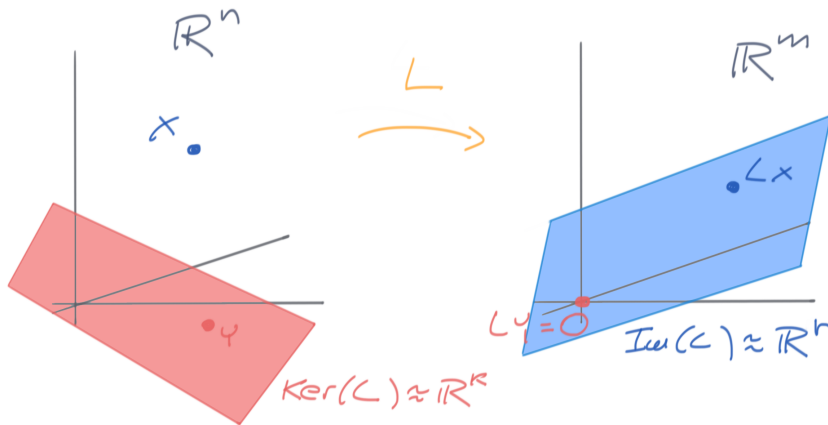
A *imagem* (em inglês, *range*) de L é

$$\text{Im}(L) := L(\mathbf{V}) = \{L(\mathbf{v}) \text{ com } \mathbf{v} \in \mathbf{V}\} \subset \mathbf{W}$$

É imediato verificar que $\text{Im}(L)$ é um subespaço vetorial de $\mathbf{W}$. De fato, se $\mathbf{w} = L(\mathbf{v})$ e $\mathbf{w}' = L(\mathbf{v}')$ com $\mathbf{v}, \mathbf{v}' \in \mathbf{V}$, e λ é um escalar, então $\mathbf{w} + \mathbf{w}' = L(\mathbf{v}) + L(\mathbf{v}') = L(\mathbf{v} + \mathbf{v}')$ e também $\lambda \mathbf{w} = \lambda L(\mathbf{v}) = L(\lambda \mathbf{v})$. É uma tautologia que

Teorema 8.6. *Uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ é sobrejetiva sse a sua imagem $\text{Im}(L)$ é o próprio $\mathbf{W}$.*

A dimensão $k = \dim(\text{Ker}(L))$ do núcleo é dita *nulidade* (em inglês, *nullity*) de L , e denotada por $\text{Null}(L)$. A dimensão $r = \dim(\text{Im}(L))$ da imagem é dita *ordem* (em inglês, *rank*) ou *caraterística* de L , e denotada por $\text{Rank}(L)$. Acontece que, quando $\mathbf{V}$ tem dimensão finita, estes dois números satisfazem um “princípio de conservação” que diz que $k + r = n$, se n é a dimensão do domínio da transformação.



Teorema 8.7 (teorema da ordem-nulidade). *Seja $L : \mathbf{V} \rightarrow \mathbf{W}$ uma transformação linear. Se $\mathbf{V}$ tem dimensão finita, então também a imagem $L(\mathbf{V})$ tem dimensão finita e*

$$\dim(\text{Ker}(L)) + \dim(\text{Im}(L)) = \dim(\mathbf{V})$$

ou seja, $\text{Rank}(L) + \text{Null}(L) = \dim(\mathbf{V})$.

Demonstração. Sejam $n = \dim(\mathbf{V})$, $m = \dim(\text{Ker}(L))$, e seja $\mathbf{v}_1, \dots, \mathbf{v}_m$ uma base de $\text{Ker}(L)$. Se $k = n$, então a imagem $\text{Im}(L)$ é trivial, e o teorema é verdadeiro. Se $m < n$, de acordo com o teorema 4.5 existem $\mathbf{v}_{m+1}, \dots, \mathbf{v}_n$ tais que os $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m, \mathbf{v}_{m+1}, \dots, \mathbf{v}_n$ formam uma base de $\mathbf{V}$. Queremos então mostrar que os vetores $\mathbf{w}_1 = L(\mathbf{v}_{m+1})$, $\mathbf{w}_2 = L(\mathbf{v}_{m+2})$, $\dots$, $\mathbf{w}_{n-m} = L(\mathbf{v}_n)$ geram $\text{Im}(L)$ e são independentes, assim que $\dim(\text{Im}(L)) = n - m$. Se $\mathbf{w} \in \text{Im}(L)$, existe um vetor $\mathbf{v} = a_1 \mathbf{v}_1 + a_2 \mathbf{v}_2 + \dots + a_n \mathbf{v}_n$ tal que $\mathbf{w} = L(\mathbf{v})$. Mas

$$\begin{aligned} \mathbf{w} &= L(a_1 \mathbf{v}_1 + a_2 \mathbf{v}_2 + \dots + a_n \mathbf{v}_n) \\ &= a_{m+1} L(\mathbf{v}_{m+1}) + a_{m+2} L(\mathbf{v}_{m+2}) + \dots + a_n L(\mathbf{v}_n) = a_{m+1} \mathbf{w}_1 + a_{m+2} \mathbf{w}_2 + \dots + a_n \mathbf{w}_{n-m} \end{aligned}$$

porque os $\mathbf{v}_1, \dots, \mathbf{v}_m$ estão no núcleo de L . Isto prova que os $\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_{n-m}$ geram $\text{Im}(L)$. Por outro lado, se os $\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_{n-m}$ são dependentes, então existem coeficientes b_k não todos nulos tais que

$$b_1 \mathbf{w}_1 + b_2 \mathbf{w}_2 + \dots + b_{n-m} \mathbf{w}_{n-m} = \mathbf{0}$$

Pela própria definição dos $\mathbf{w}_k$'s e a linearidade de L ,

$$L(b_1 \mathbf{v}_{m+1} + b_2 \mathbf{v}_{m+2} + \cdots + b_{n-m} \mathbf{v}_n) = \mathbf{0}$$

Isto quer dizer que $\mathbf{v}' = b_1 \mathbf{v}_{m+1} + b_2 \mathbf{v}_{m+2} + \cdots + b_{n-m} \mathbf{v}_n \in \text{Ker}(L)$. Então, sendo $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ uma base do núcleo, temos também

$$\mathbf{v}' = c_1 \mathbf{v}_1 + c_2 \mathbf{v}_2 + \cdots + c_m \mathbf{v}_m$$

para alguns coeficientes c_k 's. Consequentemente,

$$\mathbf{0} = \mathbf{v}' - \mathbf{v}' = c_1 \mathbf{v}_1 + c_2 \mathbf{v}_2 + \cdots + c_m \mathbf{v}_m - (b_1 \mathbf{v}_{m+1} + b_2 \mathbf{v}_{m+2} + \cdots + b_{n-m} \mathbf{v}_n)$$

Sendo pelo menos um dos $b_k \neq 0$, isto contradiz a independência dos $\mathbf{v}_k$'s. Consequentemente, os $\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_{n-m}$ são independentes, logo formam uma base de $\text{Im}(L)$. $\square$

Em particular, uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ entre dois espaços lineares de dimensão finita não pode ser injetiva se $\dim(\mathbf{W}) < \dim(\mathbf{V})$, e não pode ser sobrejetiva se $\dim(\mathbf{W}) > \dim(\mathbf{V})$. Por outro lado, quando $\dim(\mathbf{V}) = \dim(\mathbf{W})$, então uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ é injetiva sse é sobrejetiva sse é bijetiva.

ex: Determine o núcleo, a imagem, a nulidade e a ordem das seguintes transformações lineares

$$L(x, y) = (x + y, x - y) \quad L(x, y) = (y, -x) \quad L(x, y) = (x + y, 3x - 2y)$$

$$L(x, y, z) = (2x, 3y, 0) \quad L(x, y, z) = (x + y, y + z, z + x)$$

$$L(x, y, z) = (0, 0, 0) \quad L(x, y, z) = (x, y)$$

ex: Determine o núcleo, a imagem, a nulidade e a ordem da transformação $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ que transforma cada ponto no seu simétrico em relação à reta $x = 0$

ex: Determine o núcleo, a imagem, a nulidade e a ordem da transformação $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ que transforma cada ponto no seu simétrico em relação à origem.

ex: [Ap69] Vol 2 2.4.

8.5 Álgebra dos operadores

Composição e álgebra dos endomorfismos. A composição de duas transformações lineares $L : \mathbf{V} \rightarrow \mathbf{W}$ e $M : \mathbf{W} \rightarrow \mathbf{Z}$, definida por $(ML)(\mathbf{v}) := M(L(\mathbf{v}))$, é uma transformação linear. De fato, se $\mathbf{v}, \mathbf{v}' \in \mathbf{V}$ e λ é um escalar, então

$$\begin{aligned} (ML)(\mathbf{v} + \mathbf{v}') &= M(L(\mathbf{v} + \mathbf{v}')) = M(L(\mathbf{v}) + L(\mathbf{v}')) \\ &= M(L(\mathbf{v})) + M(L(\mathbf{v}')) = (ML)(\mathbf{v}) + (ML)(\mathbf{v}') \end{aligned}$$

e

$$(ML)(\lambda \mathbf{v}) = M(L(\lambda \mathbf{v})) = M(\lambda L(\mathbf{v})) = \lambda M(L(\mathbf{v})) = \lambda (ML)(\mathbf{v})$$

pela linearidade de M e L .

Em particular, a composição de dois endomorfismos de um espaço vetorial $\mathbf{V}$ é um endomorfismo de $\mathbf{V}$. A n -ésima iterada do endomorfismo $L \in \text{End}(\mathbf{V})$ é o endomorfismo L^n definido indutivamente por

$$L^0 = I, \quad L^{n+1} = LL^n \quad \text{se } n \geq 1.$$

A composição de transformações lineares é associativa, ou seja,

$$L(MN) = (LM)N$$

(como a composição de funções não necessariamente lineares) e satisfaz as propriedades distributivas

$$(L + M)N = LN + MN \quad L(M + N) = LM + LN$$

(que justificam a notação “multiplicativa” LM em vez de $L \circ M$). Se I denota o endomorfismo identidade, então $IL = LI = L$ para todo $L \in \text{End}(\mathbf{V})$. Assim, os endomorfismos $\text{End}(\mathbf{V})$ de um espaço vetorial formam uma álgebra associativa.

A composição não é comutativa! Ou seja, em geral LM é diferente ML . Os endomorfismos $L, M \in \text{End}(\mathbf{V})$ *comutam* (ou *são permutáveis*) se $LM = ML$. A obstrução é o *comutador*, definido por

$$[L, M] := LM - ML,$$

que é igual a transformação nula sse L e M comutam. É claro que todo endomorfismo comuta com si próprio. Também, todo endomorfismo comuta com a identidade I e com os seus múltiplos λI . De fato, é possível provar que os únicos endomorfismos que comutam com todos os endomorfismos são precisamente os múltiplos da identidade, ou seja, as homotetias.

ex: Calcule a composição ML quando

$$\begin{aligned} L(x, y) &= (x + y, x - y) & M(x, y) &= 2x - 3y \\ L(x, y, z) &= (x - y + z, z - y) & M(x, y) &= (x, y, x + y) \end{aligned}$$

ex: Calcule o comutador entre os endomorfismos do plano $E_+(x, y) = (y, 0)$ e $E(x, y) = (x, -y)$.

ex: Se $[L, M] = 0$ e $[M, N] = 0$, é verdade que $[L, N] = 0$?

ex: Determine todos os endomorfismos do plano que comutam com a inversão $J(x, y) = (x, y)$.

ex: Prove que os únicos endomorfismo de um espaço vetorial $\mathbf{V}$ que comutam com todos os $L \in \text{End}(\mathbf{V})$ são os múltiplos da identidade, ou seja, da forma $M = \lambda I$ para algum escalar λ .

ex: Verifique que se L e M são dois endomorfismos, então

$$(L + M)^2 = L^2 + LM + ML + M^2$$

e, se $[L, M] = 0$, também $(L + M)^2 = L^2 + 2LM + M^2$. Deduza fórmulas análogas para as potências superiores ...

ex: Seja $L : \text{End}(\mathbf{V})$ um operador. Observe que $L^{k+j} = L^j L^k$ e portanto se $L^k \mathbf{v} = \mathbf{0}$ então também $L^{k+j} \mathbf{v} = \mathbf{0}$. Deduza as desigualdades

$$\{\mathbf{0}\} = \text{Ker}(L^0) \subset \text{Ker}(L) \subset \text{Ker}(L^2) \subset \dots$$

Operadores translação, derivação, multiplicação e primitivação. Em análise, e no estudo das equações diferenciais da física-matemática, são importantes certos operadores diferenciais e integrais definidos em espaços de funções.

Consideramos o espaço linear $\mathbb{R}^{\mathbb{R}}$ das funções reais de uma variável real, e os seus subespaços naturais $\mathcal{C}^k(\mathbb{R})$ das funções com k derivadas contínuas, com $0 \leq k \leq \infty$. Fixado $a \in \mathbb{R}$, o operador *translação* $T_a : \mathbb{R}^{\mathbb{R}} \rightarrow \mathbb{R}^{\mathbb{R}}$ envia uma função $f(x)$ na função

$$(T_a f)(x) := f(x + a)$$

O operador *derivação* envia uma função derivável $f(x)$ na função

$$(Df)(x) := f'(x).$$

Pode ser pensado como um operador $D : \mathcal{C}^{k+1}(\mathbb{R}) \rightarrow \mathcal{C}^k(\mathbb{R})$, ou também como um endomorfismo de $\mathcal{C}^\infty(\mathbb{R})$. Naturalmente, D é igual ao limite $\lim_{\varepsilon \rightarrow 0} (T_\varepsilon - I) / \varepsilon$, que existe desde que pensado no subespaço das funções deriváveis. O operador *multiplicação* envia uma função $f(x)$ na função

$$(Xf)(x) := x \cdot f(x).$$

O operador *primitivação* envia uma função integrável (na reta real ou num intervalo da reta) $f(x)$ na função

$$(Rf)(x) := \int_c^x f(t) dt$$

Pode ser pensado como um operador $R : \mathcal{C}^k(\mathbb{R}) \rightarrow \mathcal{C}^{k+1}(\mathbb{R})$, ou como um endomorfismo de $\mathcal{C}^0(\mathbb{R})$ ou de $\mathcal{C}^\infty(\mathbb{R})$.

A álgebra do grupo de Weyl-Heisenberg²² (em dimensão um) é gerada pelos operadores multiplicação e derivação. Os operadores $Q = X$ e $P := -i\hbar D$ (onde $\hbar$ é a constante de Planck reduzida) representam a *posição* e o *momento*, respetivamente, de uma partícula quântica livre na “representação de Schrödinger”. Os operadores P e Q não comutam, pois

$$[Q, P] = i\hbar$$

e está é a causa do “princípio de incerteza de Heisenberg”.

ex: Determine o núcleo e a imagem dos operadores D e R .

ex: Mostre que $DR = I$ mas $RD \neq I$. Descreva o núcleo e a imagem de RD .

ex: Calcule o comutador $[D, X]$ entre os operadores derivação e multiplicação definidos em $\mathcal{C}^\infty(\mathbb{R})$.

Derivadas e primitivas discretas. Consider the spaces $\mathbb{R}^{\mathbb{N}_0}$ of real sequences (i.e. discrete time signals) $x = (x_n) = (x_0, x_1, x_2, \dots)$ (observe that time starts with “zero”). We can “integrate” a sequence and get the sequence Sx , defined by $(Sx)_0 := 0$ and

$$(Sx)_n := \sum_{k=0}^{n-1} x_k = x_0 + x_1 + \dots + x_{n-1} \quad \text{for } n \geq 1.$$

The operator S should be called *sum operator*, or also *discrete primitive*. Also, we may define its (*forward*) *discrete derivative* taking differences, as the sequence Dx defined by

$$(Dx)_n := x_{n+1} - x_n.$$

It is clear that S and D are discrete versions of the integration and derivation operators, respectively. Indeed, one easily checks that Newton’s fundamental theorem of calculus and Leibniz rule look like

$$(DSx)_n = x_n \quad \text{and} \quad (SDx)_n = x_n - x_0,$$

respectively. Observe that D and S do not commute, and that their commutator $[D, S]$ is the operator sending a sequence x to the constant sequence equal to x_0 .

Projeções. Se o espaço vetorial é uma soma direta $\mathbf{V} = X \oplus Y$, então é possível definir uma “projeção” $P : \mathbf{V} \rightarrow \mathbf{V}$ sobre o subespaço X da seguinte maneira: como todo vetor é uma soma única $\mathbf{v} = \mathbf{x} + \mathbf{y}$ com $\mathbf{x} \in X$ e $\mathbf{y} \in Y$, então

$$P(\mathbf{x} + \mathbf{y}) := \mathbf{x}$$

É evidente que $\text{Ker}(P) = Y$ e $\text{Im}(P) = X$, e que o operador P satisfaz $P^2 = P$.

²²Os físicos dizem “grupo de Weyl”, que era um matemático (colega de Einstein em Zurich e depois em Princeton), e os matemáticos dizem “grupo de Heisenberg”, que era um físico, um dos pais da mecânica quântica.

Em geral, um endomorfismo $P : \mathbf{V} \rightarrow \mathbf{V}$ de um espaço vetorial que satisfaz a relação algébrica

$$P^2 = P$$

é chamado *projeção*. Exemplos são as projeções $\pi_{\leq k} : (x_1, x_2, \dots, x_n) \mapsto (x_1, \dots, x_k, 0, \dots, 0)$ sobre os subespaços $\mathbb{R}^k \subset \mathbb{R}^n$ definidos pelas primeiras $k \leq n$ coordenadas de $\mathbb{R}^n$. Outros exemplos são as projeções ortogonais sobre subespaços do espaço euclidiano $\mathbb{R}^n$.

Se P é uma projeção, então também o endomorfismo $Q = I - P$ é uma projeção, pois satisfaz

$$Q^2 = (I - P)^2 = I - 2P + P^2 = I - P = Q$$

Também é evidente que $QP = PQ = 0$, ou seja, $\text{Im}(P) \subset \text{Ker}(Q)$ e $\text{Im}(Q) \subset \text{Ker}(P)$. De fato, é imediato verificar que $\text{Im}(P) = \text{Ker}(Q)$ e $\text{Ker}(P) = \text{Im}(Q)$. Mas a soma é o operador identidade $P + Q = I$. Isto claramente implica que o espaço total é uma soma direta

$$\mathbf{V} = \text{Im}(P) \oplus \text{Ker}(P)$$

Relativamente a esta decomposição de $\mathbf{V}$, a própria projeção é uma soma direta $P = 1 \oplus 0$ do operador identidade em $X = \text{Im}(P)$ e do operador nulo em $Y = \text{Ker}(P)$. Finalmente, as projeções estão biunivocamente associadas a representações de um espaço como soma direta de dois subespaços.

Em geral, uma família de projeções $P_1, P_2, P_3, \dots$ é *ortogonal* se $P_i P_j = P_j P_i = 0$ (ou seja, se $\text{Im}(P_j) \subset \text{Ker}(P_i)$ e $\text{Im}(P_i) \subset \text{Ker}(P_j)$) cada vez que $i \neq j$ (esta noção é algébrica, e não depende de um produto escalar!). Se o espaço é uma soma direta $\mathbf{V} = X_1 \oplus X_2 \oplus \dots \oplus X_m$ de uma família finita de subespaços X_i , então é possível definir umas projeções $P_i : \mathbf{V} \rightarrow \mathbf{V}$, com $i = 1, 2, \dots, m$, da seguinte maneira: como todo vetor é uma soma única $\mathbf{v} = \mathbf{x}_1 + \mathbf{x}_2 + \dots + \mathbf{x}_m$ de $\mathbf{x}_i \in X_i$, então

$$P_i(\mathbf{x}_1 + \mathbf{x}_2 + \dots + \mathbf{x}_m) := \mathbf{x}_i$$

É imediato então verificar que as P_i 's são ortogonais, e que o operador identidade é uma soma

$$I = P_1 + P_2 + \dots + P_m$$

destas m projeções.

ex: Sejam P uma projeção e $Q = I - P$. Verifique que $\text{Im}(P) = \text{Ker}(Q)$ e $\text{Ker}(P) = \text{Im}(Q)$.

ex: Verifique que se P é uma projeção então todo vetor é uma soma única $\mathbf{v} = \mathbf{x} + \mathbf{y}$ de um vetor $\mathbf{x} \in \text{Im}(P)$ e um vetor $\mathbf{y} \in \text{Ker}(P)$.

ex: Dê exemplos de duas projeções diferentes $P, P' : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ com a mesma imagem, por exemplo $\text{Im}(P) = \text{Im}(P') = \mathbb{R}\mathbf{i}$.

Operadores nilpotentes. Particularmente importante, para a compreensão dos operadores genéricos, é a seguinte classe de endomorfismos. Um operador $N : \mathbf{V} \rightarrow \mathbf{V}$ é dito *nilpotente* se alguma sua potência é nula. Isto significa que existe um expoente minimal $m \geq 1$ tal que $N^m = 0$, ou seja, $\text{Ker}(N^m) = \mathbf{V}$. Por outro lado, os núcleos de um operador (não necessariamente nilpotente) satisfazem as inclusões

$$\{0\} = \text{Ker}(N^0) \subset \text{Ker}(N) \subset \text{Ker}(N^2) \subset \dots$$

pois $N^{k+j} = N^j N^k$ e portanto se $N^k \mathbf{v} = 0$ então também $N^{k+j} \mathbf{v} = 0$. Se o espaço tem dimensão finita $\dim \mathbf{V} = n$ e N é nilpotente, então a potência minimal tal que $N^m = 0$ satisfaz

$$0 = \dim \text{Ker}(N^0) < \dim \text{Ker}(N) < \dim \text{Ker}(N^2) < \dots < \dim \text{Ker}(N^{m-1}) < \dim \text{Ker}(N^m) = n$$

e portanto necessariamente $m \leq n$. Consequentemente, também $\text{Ker}(N^n) = \mathbf{V}$. Os operadores nilpotentes são moralmente “raízes de zero”, e como tais devem ser considerados “desprezáveis” ou pelo menos “pequenos” em algum sentido.

e.g. Consideramos o endomorfismo $L(x, y, z) = (y, z, 0)$ de $\mathbb{R}^3 \approx \mathbb{R}\mathbf{i} \oplus \mathbb{R}\mathbf{j} \oplus \mathbb{R}\mathbf{k}$. O seu núcleo é $\mathbb{R}\mathbf{i}$ e a sua imagem é $\mathbb{R}\mathbf{i} \oplus \mathbb{R}\mathbf{j}$. O quadrado $L^2(x, y, z) = (z, 0, 0)$ tem núcleo $\mathbb{R}\mathbf{i} \oplus \mathbb{R}\mathbf{j}$ e imagem $\mathbb{R}\mathbf{i}$. O cubo L^3 é o operador nulo, cujo núcleo é o próprio $\mathbb{R}^3$.

e.g. Deslocamentos em dimensão finita. Mais em geral, os deslocamentos esquerdo $L : \mathbb{R}^n \rightarrow \mathbb{R}^n$ e direito $R : \mathbb{R}^n \rightarrow \mathbb{R}^n$, definido por

$$L(x_1, x_2, \dots, x_n) = (x_2, x_3, \dots, x_{n-1}, 0) \quad \text{e} \quad R(x_1, x_2, \dots, x_n) = (0, x_1, x_2, \dots, x_{n-1})$$

respetivamente, são nilpotentes, pois $L^n = R^n = 0$.

e.g. Derivação de polinómios de grau limitado. O operador derivação D , pensado definido no espaço $\text{Pol}_n(\mathbb{R})$ dos polinómios de grau $\leq n$, é nilpotente, pois claramente $D^{n+1} = 0$ (ou seja, a derivada $(n+1)$ -ésima de um polinómio de grau $\leq n$ é nula).

8.6 Transformações lineares invertíveis

Transformações lineares invertíveis. Sejam $\mathbf{V}$ e $\mathbf{W}$ dois espaços vetoriais, reais ou complexos. A transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ é *invertível* se admite uma *transformação inversa*, uma transformação $L^{-1} : \mathbf{W} \rightarrow \mathbf{V}$ tal que

$$L^{-1}L = I_{\mathbf{V}} \quad \text{e} \quad LL^{-1} = I_{\mathbf{W}},$$

ou seja, $L^{-1}(L(\mathbf{v})) = \mathbf{v}$ para todo $\mathbf{v} \in \mathbf{V}$ e $L(L^{-1}(\mathbf{w})) = \mathbf{w}$ para todo $\mathbf{w} \in \mathbf{W}$. É um fato geral, válido para funções não necessariamente lineares, que uma função é invertível sse é injetiva (não existem dois elementos do domínio com a mesma imagem, em inglês, *one-to-one*) e sobrejetiva (todo elemento do contradomínio é imagem de pelos menos um elemento do domínio, em inglês, *onto*). A transformação inversa é então definida por $L^{-1}(\mathbf{w}) := \mathbf{v}$ se $\mathbf{v} \in \mathbf{V}$ é o único vetor tal que $L(\mathbf{v}) = \mathbf{w} \in \mathbf{W}$.

Esta definição não é consensual. É usada, por exemplo, por Serge Lang in [La87] e por Sheldon Axler in [Ax15]. Por outro lado, Tom Apostol em [Ap69] chama “invertíveis” todas as transformações injetivas, que portanto admitem uma inversa esquerda (e consequentemente direita) definida apenas na imagem $\text{Im}(L)$, e não necessariamente em todo o contradomínio $\mathbf{W}$. Naturalmente, uma transformação injetiva é invertível no nosso sentido se pensada como uma transformação $L : \mathbf{V} \rightarrow \text{Im}(L)$.

Um argumento standard mostra que a inversa de uma função bijetiva é única. No caso de transformações lineares, temos também

Teorema 8.8. *A inversa de uma transformação linear invertível é também uma transformação linear.*

Demonstração. Se $L(\mathbf{v}) = \mathbf{w}$ e $L(\mathbf{v}') = \mathbf{w}'$, a linearidade de L implica que $L(\mathbf{v} + \mathbf{v}') = \mathbf{w} + \mathbf{w}'$. Consequentemente,

$$L^{-1}(\mathbf{w} + \mathbf{w}') = \mathbf{v} + \mathbf{v}' = L^{-1}(\mathbf{w}) + L^{-1}(\mathbf{w}')$$

Se λ é um escalar, a linearidade de L também implica que $L(\lambda\mathbf{v}) = \lambda L(\mathbf{v}) = \lambda\mathbf{w}$. Consequentemente,

$$L^{-1}(\lambda\mathbf{w}) = \lambda\mathbf{v} = \lambda L^{-1}(\mathbf{w})$$

□

Em dimensão finita, os teoremas 8.5 e 8.7 implicam que

Teorema 8.9. *Se os espaços lineares $\mathbf{V}$ e $\mathbf{W}$ têm dimensão finita e igual, então uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ é invertível sse $\text{Ker}(L) = \{\mathbf{0}\}$ sse $\text{Im}(L) = \mathbf{W}$ sse envia bases de $\mathbf{V}$ em bases de $\mathbf{W}$.*

Em particular, um endomorfismo de espaço vetorial de dimensão finita é invertível sse é injetivo sse é sobrejetivo.

Uma transformação linear invertível $L : \mathbf{V} \rightarrow \mathbf{W}$ é também dita *isomorfismo (linear)* entre os espaços lineares $\mathbf{V}$ e $\mathbf{W}$, e os dois espaços são ditos *isomorfos*. É imediato verificar que a composição LM de duas transformações invertíveis é também invertível, e a sua inversa é

$$(LM)^{-1} = M^{-1}L^{-1}$$

A transformação identidade é claramente um isomorfismo de um espaço com si próprio. Consequentemente, a existência de um isomorfismo entre dois espaços lineares é uma relação de equivalência.

É claro que um isomorfismo entre espaços lineares de dimensão finita apenas pode existir quando as dimensões são iguais. De fato,

Teorema 8.10. *Dois espaços vetoriais de dimensão finita (sobre o mesmo corpo) são isomorfos sse têm a mesma dimensão.*

Demonstração. O teorema 8.7 implica que dois espaços isomorfos têm a mesma dimensão, pois o isomorfismo tem nulidade igual a zero. Vice-versa, sejam $\mathbf{V}$ e $\mathbf{W}$ espaços lineares de dimensão n , e sejam $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ uma base de $\mathbf{V}$ e $\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_n$ uma base de $\mathbf{W}$. Então é imediato verificar que a transformação

$$x_1\mathbf{b}_1 + x_2\mathbf{b}_2 + \dots + x_n\mathbf{b}_n \mapsto x_1\mathbf{c}_1 + x_2\mathbf{c}_2 + \dots + x_n\mathbf{c}_n$$

é um isomorfismo entre os dois. □

Involuções. Um endomorfismo $R : \mathbf{V} \rightarrow \mathbf{V}$ que satisfaz a relação

$$R^2 = I$$

ou seja, igual ao seu próprio inverso $R^{-1} = R$, é chamado *involução*. Um exemplo trivial é o operador identidade. Exemplos não triviais são as reflexões, como a transformação $R(x, y) = (y, x)$ do plano.

ex: Se P é uma projeção, então $R = 2P - I$ é uma involução. Vice-versa, se R é uma involução então $P = (I + R)/2$ é uma projeção.

Grupo dos automorfismos. Os endomorfismos invertíveis de um espaço linear são chamados *automorfismos*. Exemplos são a transformação identidade I , ou as homotetias λI com $\lambda \neq 0$.

Seja $\mathbf{V}$ um espaço linear de dimensão finita. Se $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ é uma base ordenada de $\mathbf{V}$, e se $T : \mathbf{V} \rightarrow \mathbf{V}$ é um automorfismo, então as imagens $\mathbf{c}_k = T\mathbf{e}_k$ também formam uma base de $\mathbf{V}$. Vice-versa, dadas duas bases ordenadas $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ e $\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_n$ de $\mathbf{V}$, existe um único automorfismo $T : \mathbf{V} \rightarrow \mathbf{V}$ que envia a primeira base na segunda, ou seja, $\mathbf{c}_k = T\mathbf{b}_k$.

O conjunto $\text{Aut}(\mathbf{V})$ dos automorfismos de um espaço linear $\mathbf{V}$ forma um “grupo”, no seguinte sentido. A cada dois automorfismos L e M pode ser associado um automorfismo “produto” LM , a composição. Existe um automorfismo I , a identidade, tal que

$$LI = IL = L$$

para todo automorfismo L . Todo automorfismo L admite um automorfismo inverso L^{-1} , tal que

$$LL^{-1} = L^{-1}L = I$$

Em geral, o produto não é comutativo, ou seja, LM pode ser diferente de ML .

ex: Diga se L é invertível e, caso afirmativo, determine a imagem $\text{Im}(L)$ e a transformação inversa

$$L(x, y) = (x, x) \quad L(x, y) = (y, x) \quad L(x, y) = (x - y, x + y) \quad L(x, y) = (0, y)$$

$$L(x, y, z) = (x + y, y + z, z + x) \quad L(x, y, z) = (3x, 2y, z) \quad L(x, y, z) = (y, z, 0)$$

$$L(x, y, z) = (x + y + z, y, z) \quad L(x, y) = (x, 0, y) \quad L(x, y) = (x - x, x + y, 0)$$

ex: Determine a transformação inversa de uma homotetia $\mathbf{x} \mapsto \lambda \mathbf{x}$ de $\mathbb{R}^n$, com $\lambda \neq 0$.

ex: Mostre que se $L : \mathbf{V} \rightarrow \mathbf{V}$ é invertível então também L^n é invertível e $(L^n)^{-1} = (L^{-1})^n$.

ex: Mostre que se $L : \mathbf{V} \rightarrow \mathbf{V}$ e $M : \mathbf{V} \rightarrow \mathbf{V}$ comutam, então também as inversas L^{-1} e M^{-1} comutam, e $(LM)^n = L^n M^n$.

ex: Mostre que operador $S : \mathbb{R}^n \rightarrow \mathbb{R}^n$ definido por

$$S(x_1, x_2, \dots, x_n) = (x_2, x_3, \dots, x_n, x_1)$$

(que representa uma translação no espaço das funções no grupo finito $\mathbb{Z}/n\mathbb{Z}$) é invertível, e calcule a sua inversa (observe que $S^n = I$) ...)

ex: O operador *multiplicação*, definido por $(Xf)(x) := x f(x)$, em $C^\infty(\mathbb{R})$, é invertível?

ex: O operador *deslocamento* $\sigma : \mathbb{R}^{\mathbb{N}} \rightarrow \mathbb{R}^{\mathbb{N}}$, definido por

$$(x_1, x_2, x_3, \dots) \mapsto (x_2, x_3, x_4, \dots),$$

é invertível?

ex: [Ap69] Vol 2 2.8.

9 Transformações lineares e matrizes

ref: [Ap69] Vol 2, 2.9-16 ; [La87] Ch. IV

9.1 Matrizes

17 nov 2022

Matrizes. Uma *matriz* real (ou complexa) $m \times n$ é uma função real (ou complexa) A definida no produto cartesiano $\{1, 2, \dots, m\} \times \{1, 2, \dots, n\}$, que associa portanto um número $A(i, j) = a_{ij}$ a cada par ordenado de inteiros (i, j) com $1 \leq i \leq m$ e $1 \leq j \leq n$.

Mais conveniente é “visualizar” uma matriz como uma tabela retangular

$$A = (a_{ij}) = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}$$

de $m \cdot n$ números reais (ou complexos) dispostos em m linhas e n colunas. O número real (ou complexo) a_{ij} é dito *elemento/componente/entrada* ij da matriz A . Os vetores

$$A_{i*} = (a_{i1}, a_{i2}, \dots, a_{in}) \in \mathbb{R}^n \quad \text{e} \quad A_{*j} = \begin{pmatrix} a_{1j} \\ a_{2j} \\ \vdots \\ a_{mj} \end{pmatrix} \in \mathbb{R}^m$$

são ditos *i-ésima linha* e *j-ésima coluna* da matriz A , respetivamente. Em particular, uma matriz com m linhas e apenas uma coluna é um vetor de $\mathbb{R}^m$ representado como um vetor coluna, e uma matriz com n colunas e apenas uma linha é um vetor de $\mathbb{R}^n$ representado como um vetor linha. Se o número de linhas é igual ao número de colunas, ou seja, se $n = m$, a matriz é dita *quadrada*.

Listas de listas. As matrizes ocorrem principalmente, no contexto da álgebra, enquanto “coordenadas” de transformações lineares (ou de certos “tensores”), como veremos à seguir. No entanto, podem simplesmente ser pensadas como listas formadas por listas. Numa linguagem de programação como **Python**, a matriz

$$A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix}$$

pode ser definida como

```
8
9  A = [[1, 2, 3], [4, 5, 6]]
10
```

uma lista formada pelas duas listas $A[0] = [1, 2, 3]$ e $A[1] = [4, 5, 6]$.

Por exemplo, o sensor de uma boa câmara digital contemporânea captura uma imagem dividindo a luz em 5568×3712 pixel (ou seja, aproximadamente 20.7 Megapixel). A imagem é assim codificada numa matriz $A = (a_{ij})$ de dimensão 5568×3712 . Numa imagem em preto e branco, cada pixel pode assumir valores inteiros entre $0 \leq a_{ij} < 2^8 = 256$ (ou seja, um número de 8 bits, entre 00000000 e 11111111), que correspondem a diferentes gradações de cinzento (uma imagem em cores parametriza cada pixel com 3 número de 8 bits, associados a três cores básicas, e câmaras reais também registam outras informações ...). É claro que uma imagem típica do mundo real (uma chita que corre atrás de uma gazela na savana) não é que uma matriz de bits, uma lista de listas, sem mais estrutura matematicamente interessante ...

Espaço linear das matrizes. Fixados os inteiros m e n , o espaço $\text{Mat}_{m \times n}(\mathbb{R})$ (ou $\text{Mat}_{m \times n}(\mathbb{C})$) das matrizes reais (ou complexas) $m \times n$ é um espaço linear real (ou complexo) se a adição e a multiplicação por um escalar são definidas por

$$A + B := (a_{ij} + b_{ij}) \quad \text{e} \quad \lambda A := (\lambda a_{ij})$$

se $A = (a_{ij})$ e $B = (b_{ij}) \in \text{Mat}_{m \times n}(\mathbb{R})$, e $\lambda \in \mathbb{R}$ (ou $\mathbb{C}$). Por exemplo,

$$\begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} + \begin{pmatrix} 5 & 6 \\ 7 & 8 \end{pmatrix} = \begin{pmatrix} 1+5 & 2+6 \\ 3+7 & 4+8 \end{pmatrix} = \begin{pmatrix} 6 & 8 \\ 10 & 12 \end{pmatrix}$$

e

$$3 \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} = \begin{pmatrix} 3 \cdot 1 & 3 \cdot 2 \\ 3 \cdot 3 & 3 \cdot 4 \end{pmatrix} = \begin{pmatrix} 3 & 6 \\ 9 & 12 \end{pmatrix}$$

O elemento neutro, ou vetor nulo, é a “matriz nula” $0 = (0)$, cujas entradas são todas nulas, e que satisfaz $A + 0 = A$ para toda a matriz A . A matriz “oposta” da matriz A é a matriz $-A := (-1)A$, tal que $A + (-A) = 0$. Então, podemos simplificar a notação e escrever $A - B := A + (-B)$.

A dimensão do espaço $\text{Mat}_{m \times n}(\mathbb{R})$ é igual ao produto $m \cdot n$, o número de elementos das matrizes. De fato, uma base natural é o conjunto das matrizes I_{ij} ’s que têm um única entrada não nula, e igual a um, na posição ij , ou seja, as matrizes do gênero

$$I_{ij} = \begin{pmatrix} 0 & \dots & 0 & \dots & 0 \\ \vdots & & \vdots & & \vdots \\ 0 & \dots & 1 & \dots & 0 \\ \vdots & & \vdots & & \vdots \\ 0 & \dots & 0 & \dots & 0 \end{pmatrix} \quad \begin{matrix} \leftarrow i \\ \\ \uparrow j \end{matrix} \quad (9.1)$$

ao variar $1 \leq i \leq m$ e $1 \leq j \leq n$. É tautológico que toda matriz $m \times n$ é uma combinação linear única

$$A = \sum_{1 \leq i \leq m} \sum_{1 \leq j \leq n} a_{ij} I_{ij}$$

com coordenadas reais ou complexas a_{ij} ’s, e esta representação define um isomorfismo natural $\text{Mat}_{m \times n}(\mathbb{R}) \approx \mathbb{R}^{m \cdot n}$ ou $\text{Mat}_{m \times n}(\mathbb{C}) \approx \mathbb{C}^{m \cdot n}$

e.g. Espaço linear das matrizes 2×2 . Por exemplo, as matrizes

$$I_{11} = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \quad I_{12} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \quad I_{21} = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} \quad I_{22} = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}$$

formam uma base do espaço linear $\text{Mat}_{2 \times 2}(\mathbb{R})$ ou $\text{Mat}_{2 \times 2}(\mathbb{C})$. Uma matriz genérica é uma combinação linear

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix} = aI_{11} + bI_{12} + cI_{21} + dI_{22}$$

Álgebra das matrizes. Se $A = (a_{ij}) \in \text{Mat}_{m \times n}(\mathbb{R})$ e $B = (b_{ij}) \in \text{Mat}_{n \times p}(\mathbb{R})$, então o *produto (linhas por colunas)* das matrizes A e B (nesta ordem) é a matriz $AB = (c_{ij}) \in \text{Mat}_{m \times p}(\mathbb{R})$ definida por

$$c_{ij} := \sum_{k=1}^n a_{ik} b_{kj} \quad (9.2)$$

Ou seja, o elemento c_{ij} do produto $C = AB$ é igual ao produto escalar $c_{ij} = A_{i*} \cdot B_{*j}$ entre a i -ésima linha de A e a j -ésima coluna de B (ou, melhor, é o valor $\langle A_{i*}, B_{*j} \rangle$ da forma linear A_{i*} sobre o vetor B_{*j}). Por exemplo, se

$$A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} \quad \text{e} \quad B = \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix}$$

então

$$\begin{aligned}
 AB &= \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix} \\
 &= \begin{pmatrix} (1, 2, 3) \cdot (1, 3, 5) & (1, 2, 3) \cdot (2, 4, 6) \\ (4, 5, 6) \cdot (1, 3, 5) & (4, 5, 6) \cdot (2, 4, 6) \end{pmatrix} \\
 &= \begin{pmatrix} \boxed{1 \cdot 1 + 2 \cdot 3 + 3 \cdot 5} & \boxed{1 \cdot 2 + 2 \cdot 4 + 3 \cdot 6} \\ \boxed{4 \cdot 1 + 5 \cdot 3 + 6 \cdot 5} & \boxed{4 \cdot 2 + 5 \cdot 4 + 6 \cdot 6} \end{pmatrix} = \begin{pmatrix} 22 & 28 \\ 49 & 64 \end{pmatrix}
 \end{aligned}$$

Observe que o produto AB apenas pode ser definido quando o número de colunas de A (a dimensão do espaço onde vivem as linhas de A) é igual ao número de linhas de B (a dimensão do espaço onde vivem as colunas de B).

Em particular, o produto linha por colunas de uma matriz linha

$$A = (a_1, a_2, \dots, a_n),$$

por uma matriz coluna

$$B = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{pmatrix}$$

é um escalar

$$AB = a_1 b_1 + a_2 b_2 + \dots + a_n b_n$$

igual ao produto escalar $\mathbf{a} \cdot \mathbf{b}$ dos vetores $\mathbf{a} = (a_1, a_2, \dots, a_n)$ e $\mathbf{b} = (b_1, b_2, \dots, b_n)$.

A *matriz identidade* de dimensão n é a matriz quadrada $I \in \text{Mat}_{n \times n}(\mathbb{R})$ definida por

$$I = (\delta_{ij}) = \begin{pmatrix} 1 & 0 & 0 & \dots \\ 0 & 1 & 0 & \dots \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & 1 \end{pmatrix}.$$

(também denotada por I_n quando é necessário lembrar a dimensão), onde o *símbolo de Kronecker* é definido por

$$\delta_{ii} := \begin{cases} 1 & \text{se } i = j \\ 0 & \text{se } i \neq j \end{cases}$$

É imediato verificar que

$$IA = A \quad \text{e} \quad BI = B$$

para todas as matrizes A e B (quando o produto faz sentido).

O produto é “associativo”,

$$\boxed{A(BC) = (AB)C}$$

e satisfaz as “propriedades distributivas” à esquerda e à direita,

$$\boxed{A(B + C) = AB + AC \quad \text{e} \quad (A + B)C = AC + BC}$$

Estas propriedades são consequências imediatas da definição (9.2) e das leis associativas e distributivas da aritmética elementar. São também consequências naturais da interpretação do produto em termos de composição de transformações lineares, que é a verdadeira motivação do produto linhas por colunas. É também evidente que se λ é um escalar então

$$(\lambda A)B = A(\lambda B) = \lambda(AB)$$

Juntamente com a associatividade, isto implica que o produto é bilinear, ou seja, satisfaz

$$A(\lambda B + \mu C) = \lambda AB + \mu AC \quad \text{e} \quad (\lambda A + \mu B)C = \lambda AC + \mu BC$$

para todos escalares λ e μ , e para todas matrizes A, B, C com as dimensões corretas para os produtos serem definidos.

Em particular, o produto de duas matrizes quadradas $n \times n$ é ainda uma matriz quadrada $n \times n$. Assim, os espaços $\text{Mat}_{n \times n}(\mathbb{R})$ e $\text{Mat}_{n \times n}(\mathbb{C})$ formam umas “álgebras associativas”, ou seja, uns espaços lineares onde está definido um produto interno bilinear e associativo munido de uma identidade multiplicativa.

e.g. Álgebra das matrizes 2×2 . Por exemplo, o conjunto

$$\mathbf{1} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad \mathbf{i} = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \quad \mathbf{j} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \mathbf{k} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

(apesar da notação, nada a ver com os vetores da base canónica do espaço de dimensão 3!) é também uma base do espaço linear $\text{Mat}_{2 \times 2}(\mathbb{R})$, ou seja, toda matriz 2×2 é uma combinação linear única $M = x\mathbf{1} + y\mathbf{i} + z\mathbf{j} + w\mathbf{k}$. Os quadrados satisfazem

$$\mathbf{i}^2 = -\mathbf{1} \quad \mathbf{j}^2 = \mathbf{k}^2 = \mathbf{1},$$

os produtos mistos satisfazem

$$\mathbf{ij} = -\mathbf{ji} = \mathbf{k} \quad \mathbf{jk} = -\mathbf{kj} = -\mathbf{i} \quad \mathbf{ki} = -\mathbf{ik} = \mathbf{j},$$

e portanto $\mathbf{ijk} = \mathbf{1}$. Esta é também conhecida como álgebra dos *co-quaterniões*.

Números complexos e matrizes reais 2×2 . Mais interessantes é que a álgebra das matrizes reais 2×2 “contém” a álgebra dos números complexos. A cada número complexo $z = x + iy$ é possível associar uma matriz real 2×2 definida por

$$M(z) := xI + yJ = \begin{pmatrix} x & -y \\ y & x \end{pmatrix},$$

onde I denota a matriz identidade e

$$J = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$$

é uma matriz que satisfaz $J^2 = -I$ (que lembra a identidade $i^2 = -1$). Observem que a imagem desta aplicação $M : \mathbb{C} \rightarrow \text{Mat}_{2 \times 2}(\mathbb{R})$ é um subconjunto próprio do espaço linear de todas as matrizes reais 2×2 , o plano das matrizes $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ definido pelas equações $a - d = 0$ e $b + c = 0$.

É imediato verificar que este é um isomorfismo entre as duas estruturas algébricas. Ou seja, a soma $z + w$ entre os números complexos z e w corresponde à soma $M(z + w) = M(z) + M(w)$ entre as matrizes, e o produto zw entre os números complexos z e w corresponde ao produto $M(zw) = M(z)M(w)$ entre as matrizes (que é portanto comutativo para esta classe de matrizes). O determinante da matriz $M(z)$ é

$$\text{Det } M(z) = x^2 + y^2 = |z|^2,$$

o quadrado do valor absoluto de z . Em particular, se $z \neq 0$ então a matriz $M(z)^{-1} := M(1/z)$ satisfaz $M(z)^{-1}M(z) = I$. Também interessante é observar que $M(\bar{z}) = M(z)^\top$ (a transposta da matriz A , definida mais a frente, é uma matriz $A^\top$ cujas linhas são as colunas de A).

ex: Considere as matrizes

$$A = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \quad B = \begin{pmatrix} 5 & 6 & 7 \\ 7 & 6 & 5 \end{pmatrix} \quad C = \begin{pmatrix} 1 & -1 \\ -1 & 1 \\ 1 & -1 \end{pmatrix}$$

Calcule $3A$, AB , BC , CA , $A - BC$, ...

ex: Existem matrizes $A \neq 0$ e $B \neq 0$ tais que $AB = 0$?

9.2 Matrizes e transformações lineares

Matriz de uma transformação linear. Uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$, definida num espaço de linear dimensão finita $\mathbf{V}$, é determinada pelos seus valores nos vetores de uma base. De fato, se $\mathcal{B} = (\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n)$ é uma base ordenada de $\mathbf{V} \approx \mathbb{R}^n$, e os vetores

$$\mathbf{w}_j := L(\mathbf{b}_j)$$

com $j = 1, 2, \dots, n$, são as imagens dos $\mathbf{b}_k$'s pela transformação linear L , então a imagem do vetor genérico $\mathbf{x} = x_1\mathbf{b}_1 + x_2\mathbf{b}_2 + \dots + x_n\mathbf{b}_n$ é

$$\begin{aligned} L(\mathbf{x}) &= L(x_1\mathbf{b}_1 + x_2\mathbf{b}_2 + \dots + x_n\mathbf{b}_n) \\ &= x_1L(\mathbf{b}_1) + x_2L(\mathbf{b}_2) + \dots + x_nL(\mathbf{b}_n) \\ &= x_1\mathbf{w}_1 + x_2\mathbf{w}_2 + \dots + x_n\mathbf{w}_n. \end{aligned}$$

Em particular, os vetores $\mathbf{w}_1, \dots, \mathbf{w}_n$ geram $\text{Im}(L) \subset \mathbf{W}$.

Se também $\mathbf{W}$ tem dimensão finita, e $\mathcal{C} = (\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_m)$ é uma base ordenada de $\mathbf{W} \approx \mathbb{R}^m$, então podemos definir os a_{ij} como sendo as coordenadas dos $\mathbf{w}_j$'s relativamente à base $\mathcal{C}$, ou seja,

$$\mathbf{w}_j = a_{1j}\mathbf{c}_1 + a_{2j}\mathbf{c}_2 + \dots + a_{mj}\mathbf{c}_m$$

com $j = 1, 2, \dots, n$. Em outras palavras, se $\mathbf{c}_1^*, \mathbf{c}_2^*, \dots, \mathbf{c}_m^*$ denota a base dual, então

$$a_{ij} = \langle \mathbf{c}_i, L\mathbf{b}_j \rangle \quad (9.3)$$

Então a imagem do vetor genérico $\mathbf{x} = x_1\mathbf{b}_1 + x_2\mathbf{b}_2 + \dots + x_n\mathbf{b}_n$ é

$$\begin{aligned} L(\mathbf{x}) &= L\left(\sum_j x_j \mathbf{b}_j\right) = \sum_{j=1}^n x_j L(\mathbf{b}_j) \\ &= \sum_{j=1}^n x_j \left(\sum_{i=1}^m a_{ij} \mathbf{c}_i\right) = \sum_{i=1}^m \left(\sum_{j=1}^n a_{ij} x_j\right) \mathbf{c}_i \end{aligned}$$

Portanto, as coordenadas do vetor $\mathbf{y} = L(\mathbf{x})$ na base escolhida $\mathcal{C}$ são

$$y_i = \sum_{j=1}^n a_{ij} x_j \quad i = 1, 2, \dots, m.$$

Os números a_{ij} são os elementos de uma matriz $A = (a_{ij}) \in \text{Mat}_{m \times n}(\mathbb{R})$, a matriz que representa a transformação L nas bases escolhidas, ou também “a matriz de L relativamente às bases $\mathcal{B}$ e $\mathcal{C}$ ” (uma notação pedante deve lembrar que a matriz depende da transformação e das duas bases, e portanto deve ser algo do género $A_L^{\mathcal{B}, \mathcal{C}}$...). As colunas de A são os vetores $L(\mathbf{b}_j)$'s na base dos $\mathbf{c}_i$'s.

Transformação linear definida por uma matriz. Vice-versa, fixadas as bases canónicas dos espaços $\mathbb{R}^n$ e $\mathbb{R}^m$ (ou duas bases de dois espaços lineares de dimensão finita n e m , respetivamente), uma matriz $A \in \text{Mat}_{m \times n}(\mathbb{R})$ define uma transformação linear $L_A : \mathbb{R}^n \rightarrow \mathbb{R}^m$ da seguinte maneira. Se X e Y denotam os “vetores coluna” (matrizes com apenas uma coluna)

$$X := \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{R}^n \quad \text{e} \quad Y := \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{pmatrix} \in \mathbb{R}^m$$

e

$$A := \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \in \text{Mat}_{m \times n}(\mathbb{R}),$$

então a transformação linear $L_A : \mathbb{R}^n \rightarrow \mathbb{R}^m$ é definida pela equação matricial

$$X \mapsto Y = L_A(X) := AX$$

ou seja, explicitamente,

$$\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{pmatrix} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$$

De acordo com a definição de produto linha por colunas, isto quer dizer precisamente que as coordenadas de $Y = L_A(X)$, a imagem do vetor X de coordenadas $x_1, x_2, \dots, x_n$, são $y_i = \sum_{j=1}^n a_{ij}x_j$, com $i = 1, 2, \dots, m$.

É claro que as matrizes $A + B$ e λA definem as transformações $L_A + L_B$ e λL_A , respetivamente. Assim, a correspondência $A \mapsto L_A$ é um isomorfismo entre os espaços vetoriais $\text{Mat}_{m \times n}(\mathbb{R})$ e $\text{Lin}(\mathbb{R}^n, \mathbb{R}^m)$. O isomorfismo inverso será denotado por $L \mapsto A_L$.

Finalmente, as matrizes podem ser pensadas como “coordenadas” de transformações lineares entre espaços vetoriais de dimensão finita, como tais dependentes da escolha de umas bases, que podem não ser canónicas.

É sem dúvida útil aprender a ler e manipular matrizes, e é um dos objetivos desta UC. No entanto, não deixam de ter sentido as palavras de Emil Artin:²³ ‘It is my experience that proofs involving matrices can be shortened by 50% if one throws the matrices out.’

ex: Fixados $1 \leq i \leq m$ e $1 \leq j \leq n$, seja $I_{ij} \in \text{Mat}_{m \times n}(\mathbb{R})$ a matriz cuja única entrada não nula é a unidade na posição ij , definida em (9.1). Verifique que o operador $L : \mathbb{R}^n \rightarrow \mathbb{R}^m$, definido pela matriz I_{ij} relativamente às bases $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ de $\mathbb{R}^n$ e $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_m$ de $\mathbb{R}^m$, envia

$$L(x_1\mathbf{e}_1 + x_2\mathbf{e}_2 + \dots + x_n\mathbf{e}_n) = x_j\mathbf{b}_i$$

Em particular, a sua imagem é a reta $\text{Im}(L) = \mathbb{R}\mathbf{b}_i$ gerada pelo i -ésimo vetor da base de $\mathbb{R}^m$, e o seu núcleo é o hiperplano $\text{Ker}(L) = \text{Ker}(\mathbf{e}_i^*) = \{(x_1, x_2, \dots, x_n) \text{ t.q. } x_j = 0\}$, o espaço nulo da forma $\mathbf{e}_j^* \in (\mathbb{R}^n)^*$.

Formas lineares e matrizes linha. Se, fixada uma base, representamos os vetores de $\mathbb{R}^n$ como vetores coluna, então as formas lineares, na base dual, são representadas por vetores linha, e a dualidade é dada pelo produto linhas por colunas. De fato, o produto linhas por colunas

$$\Xi X = \begin{pmatrix} \xi_1 & \xi_2 & \dots & \xi_n \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \xi_1 x_1 + \xi_2 x_2 + \dots + \xi_n x_n$$

do vetor linha

$$\Xi = \begin{pmatrix} \xi_1 & \xi_2 & \dots & \xi_n \end{pmatrix}$$

pelo vetor coluna

$$X = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$$

representa o valor $\langle \xi, \mathbf{x} \rangle$ da forma linear $\xi = \xi_1 \mathbf{e}_1^* + \xi_2 \mathbf{e}_2^* + \dots + \xi_n \mathbf{e}_n^* \in (\mathbb{R}^n)^*$ sobre o vetor $\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_n \mathbf{e}_n \in \mathbb{R}^n$.

²³E. Artin, *Geometric Algebra*, Interscience, 1957.

Caraterística. Seja $L_A : \mathbb{R}^n \rightarrow \mathbb{R}^m$ a transformação linear $X \mapsto Y = AX$, definida pela matriz $A \in \text{Mat}_{m \times n}(\mathbb{R})$. Se $E_1, E_2, \dots, E_n$ denotam os vetores coluna da base canônica de $\mathbb{R}^n$, ou seja,

$$E_1 := \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad E_2 := \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix}, \quad \dots \quad E_n := \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix},$$

então o produto $AE_k = A_{*k}$ é a k -ésima coluna da matriz A , que representa portanto a imagem $L_A(E_k)$ do vetor E_k .

A imagem do vetor (coluna) genérico $X = x_1 E_1 + x_2 E_2 + \dots + x_n E_n \in \mathbb{R}^n$ é uma combinação linear

$$AX = x_1 A_{*1} + x_2 A_{*2} + \dots + x_n A_{*n}$$

das colunas da matriz A . A ordem da transformação linear L_A , ou seja, a dimensão de $\text{Im}(L_A) \subset \mathbb{R}^m$, é portanto igual ao número de colunas linearmente independentes de A . É também chamada *caraterística* da matriz A , e denotada por

$$\text{Rank}(A) := \dim \text{Im}(L_A)$$

Por outro lado, o vetor $X \in \mathbb{R}^n$ pertence ao espaço nulo $\text{Ker}(L_A)$ da transformação linear L_A se $AX = 0$, ou seja, se

$$A_{1*} \cdot X = 0 \quad A_{2*} \cdot X = 0 \quad \dots \quad A_{m*} \cdot X = 0,$$

onde $A_{i*} = (a_{i1}, a_{i2}, \dots, a_{in}) \in \mathbb{R}^n$ é a i -ésima linha da matriz A . O espaço nulo é portanto o espaço ortogonal ao subespaço vetorial $\text{Span}(A_{1*}, A_{2*}, \dots, A_{m*}) \subset \mathbb{R}^n$ gerado pelas linhas de A , ou seja,

$$\text{Ker}(L_A) = \text{Span}(A_{1*}, A_{2*}, \dots, A_{m*})^\perp.$$

Pela (4.2), a sua dimensão é igual a $n - k$, se k é o número de linhas linearmente independentes de A . Em particular, sendo $\dim \text{Ker}(L_A) + \dim \text{Im}(L_A) = n$ pelo teorema 8.7, a caraterística da matriz $\text{Rank}(A)$ é também igual ao número de linhas linearmente independentes de A . Resumindo,

Teorema 9.1. A caraterística $\text{Rank}(A)$ de uma matriz $A \in \text{Mat}_{m \times n}(\mathbb{R})$, ou seja, a ordem da transformação linear $L_A \in \text{Lin}(\mathbb{R}^n, \mathbb{R}^m)$, é igual ao número de colunas linearmente independentes assim como ao número de linhas linearmente independentes da matriz.

9.3 Composição e produto de matrizes

Produto e composição. Se uma primeira matriz $A \in \text{Mat}_{m \times n}(\mathbb{R})$ define a transformação linear $L_A : \mathbb{R}^n \rightarrow \mathbb{R}^m$, e se uma segunda matriz $B = (b_{ij}) \in \text{Mat}_{p \times m}(\mathbb{R})$ define a transformação linear $L_B : \mathbb{R}^m \rightarrow \mathbb{R}^p$, então a matriz produto $BA \in \text{Mat}_{p \times n}(\mathbb{R})$ define a composição $L_B L_A : \mathbb{R}^n \rightarrow \mathbb{R}^p$. Ou seja,

$$L_{BA} = L_B L_A$$

De fato, se as coordenadas de $Y = L_A(X)$ são $y_k = \sum_{j=1}^n a_{kj} x_j$ e as coordenadas de $Z = L_B(Y)$ são $z_i = \sum_{k=1}^m b_{ik} y_k$, então

$$z_i = \sum_{k=1}^m b_{ik} y_k = \sum_{k=1}^m b_{ik} \left(\sum_{j=1}^n a_{kj} x_j \right) = \sum_{j=1}^n \left(\sum_{k=1}^m b_{ik} a_{kj} \right) x_j$$

Em notação matricial, as coisas são ainda mais simples:

$$\text{se } Y = AX \quad \text{e} \quad Z = BY \quad \text{então} \quad Z = BAX$$

Assim, o produto linha por colunas corresponde à composição entre transformações lineares. A associatividade e a bilinearidade do produto são então também consequências das análogas propriedades da composição entre transformações lineares.

ex: Determine a matriz que define a transformação

$$L(x, y) = (x - y, 2x - 3y) \quad L(x, y, z) = (3x + y - z, -x + 2y + z)$$

$$L(x, y, z) = (3x, 3y, 3z) \quad L(x, y) = (x + y, x - y, 2x - 7y)$$

$$L(x, y, z) = (x, y) \quad L(x, y, z) = (x, z)$$

ex: Determine a transformação linear $L : \mathbb{R}^n \rightarrow \mathbb{R}^m$ definida pela matriz

$$\begin{pmatrix} 1 & 3 \\ -2 & 0 \end{pmatrix} \quad \begin{pmatrix} -5 & 0 & -1 \\ 3 & 1 & 2 \end{pmatrix} \quad \begin{pmatrix} 0 & -1 \\ 0 & 8 \\ -1 & -3 \end{pmatrix} \quad \begin{pmatrix} 1 & 3 \\ -2 & 0 \end{pmatrix}$$

ex: Determine a matriz 2×2 que define a transformação $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ que

transforma cada ponto no seu simétrico em relação à reta $x = 0$

transforma cada ponto no seu simétrico em relação à reta $y = -x$

transforma o ponto de coordenadas polares (r, θ) no ponto de coordenadas polares $(r/2, \theta)$

transforma o ponto de coordenadas polares (r, θ) no ponto de coordenadas polares $(r, \theta - \pi/2)$

ex: [Ap69] Vol 2 2.12.

ex: [Ap69] Vol 2 2.16.

9.4 Transposta e dualidade

Matrizes transpostas. Seja $A = (a_{ij}) \in \text{Mat}_{m \times n}(\mathbb{R})$ uma matriz. A *matriz transposta* é a matriz $A^\top = (a'_{ij}) \in \text{Mat}_{n \times m}(\mathbb{R})$, definida por $a'_{ij} = a_{ji}$ (ou seja, as linhas de $A^\top$ são as colunas de A e vice-versa). Por exemplo,

$$\text{se } A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} \quad \text{então } A^\top = \begin{pmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{pmatrix}$$

É imediato verificar que $(A^\top)^\top = A$ e que $(AB)^\top = B^\top A^\top$ (quando o primeiro produto faz sentido).

Por exemplo, se X e Y são dois vetores/matrizes coluna de $\mathbb{R}^n$, então $Y^\top$ é um vetor/matriz linha, e o produto linha por coluna $Y^\top X$ é igual ao produto escalar

$$Y^\top X = Y \cdot X.$$

Em geral, se X é um vetor coluna de $\mathbb{R}^n$ e Y é um vetor coluna de $\mathbb{R}^m$, e se a matriz $A \in \text{Mat}_{m \times n}(\mathbb{R})$ define a transformação linear $X \mapsto AX$, então a matriz transposta define a transformação linear $Y \mapsto A^\top Y$ tal que

$$Y \cdot AX = (A^\top Y) \cdot X,$$

pois $Y \cdot AX = Y^\top AX = (A^\top Y)^\top X = (A^\top Y) \cdot X$.

Uma matriz quadrada A é dita *simétrica* se $A = A^\top$ (ou seja, se $a_{ij} = a_{ji}$) e *anti-simétrica* se $A^\top = -A$ (ou seja, se $a_{ij} = -a_{ji}$).

ex: Verifique que $(A^\top)^\top = A$ e que $(AB)^\top = B^\top A^\top$.

ex: Mostre que matrizes simétricas e matrizes anti-simétricas formam dois subespaços do espaço $\text{Mat}_{n \times n}(\mathbb{R})$ das matrizes quadradas, e calcule as dimensões.

ex: Mostre que, se A é uma matriz quadrada, então $A + A^\top$ é simétrica, e $A - A^\top$ é anti-simétrica. Deduza que cada matriz quadrada pode ser decomposta de maneira única como soma $A = A_+ + A_-$ de uma matriz simétrica $A_+ = (A + A^\top)/2$ e uma matriz anti-simétrica $A_- = (A - A^\top)/2$.

Transformação dual. O verdadeiro significado da matriz transposta fica escondido pelos isomorfismos entre $\mathbb{R}^n$ e o seu dual definido pelo produto escalar. Seja $L : \mathbf{V} \rightarrow \mathbf{W}$ uma transformação linear entre os espaços vetoriais $\mathbf{V}$ e $\mathbf{W}$. A *transformação transposta*, ou *dual*, é a transformação linear $L^\top : \mathbf{W}^* \rightarrow \mathbf{V}^*$ entre os espaços duais definida por $L^\top \xi := \xi \circ L$, ou seja,

$$\langle L^\top \xi, \mathbf{v} \rangle = \langle \xi, L\mathbf{v} \rangle \quad (9.4)$$

se $\xi \in \mathbf{W}^*$ e $\mathbf{v} \in \mathbf{V}$ um vetor arbitrário (uma notação mais natural é L^* , que tem o defeito de ser também a notação preferida para o operador adjunto ...). A linearidade é evidente. É também claro que esta definição não depende das bases, e de fato faz sentido também em dimensão infinita.

É fácil verificar que a transformação transposta de uma soma é a soma das transpostas, ou seja,

$$(L + T)^\top = L^\top + T^\top$$

e que

$$(\lambda L)^\top = \lambda L^\top$$

se λ é um escalar. Ou seja, a transposição $L \mapsto L^\top$ é uma transformação linear de $\text{Lin}(\mathbf{V}, \mathbf{W})$ em $\text{Lin}(\mathbf{W}^*, \mathbf{V}^*)$.

Também é imediato verificar que a transformação transposta de uma composição é

$$(LT)^\top = T^\top L^\top$$

Observe a ordem invertida, que é a única que faz sentido se as dimensões dos três espaços envolvidos são diferentes. Os matemáticos dizem então que a transposição é uma operação “contra-variante”, pois inverte as direções das setas..

Se os espaços têm dimensão finita, então, fixadas uma base ordenada $(\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n)$ de $\mathbf{V} \approx \mathbb{R}^n$ e uma base ordenada $(\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_m)$ de $\mathbf{W} \approx \mathbb{R}^m$, a transformação linear L é definida por uma matriz $A = (a_{ij}) \in \text{Mat}_{m \times n}(\mathbb{R})$, e envia o vetor $\mathbf{x}$ de coordenadas x_j 's no vetor $\mathbf{y} = L\mathbf{x}$ de coordenadas

$$y_i = \sum_{j=1}^n a_{ij} x_j.$$

Em termos de vetores coluna, a transformação é $X \mapsto Y = AX$. Uma forma linear $\xi \in \mathbf{W}^*$ é representada na base dual de $\mathbf{W}$ por um vetor linha $\Xi = (\xi_1, \xi_2, \dots, \xi_m)$, e o seu valor sobre o vetor Y é simplesmente $\langle \xi, \mathbf{y} \rangle = \sum_{i=1}^m \xi_i y_i$, o produto linhas por coluna ΞY . Então a forma $L^* \xi$ é representada, na base dual de $\mathbf{V}$, por um vetor linha $\Xi' = (\xi'_1, \xi'_2, \dots, \xi'_n)$ tal que, de acordo com a (9.4),

$$\sum_{j=1}^n \xi'_j x_j = \sum_{i=1}^m \xi_i \left(\sum_{j=1}^n a_{ij} x_j \right) = \sum_{j=1}^n \sum_{i=1}^m \xi_i a_{ij} x_j$$

ou seja, na notação matricial

$$\Xi' X = \Xi A X$$

Pela arbitrariedade de X , e a associatividade do produto entre matrizes, isto significa que a forma $L^\top \xi$ é representada pelo vetor linha $\Xi' = \Xi A$. As coordenadas da forma $\xi' = L^* \xi$ são dadas por

$$\xi'_j = \sum_{i=1}^m \xi_i a_{ij}.$$

A transformação $L^\top : \mathbf{W}^* \rightarrow \mathbf{V}^*$ é portanto definida, nas bases duais, pela matriz transposta $A^\top \in \text{Mat}_{n \times m}(\mathbb{R})$.

ex: Calcule o operador transposto de $L(x, y) = (x + y, x - y)$, de $M(x, y) = (x + y, y)$ e de $N(x, y) = (y, x)$.

ex: Considere o operador derivação $D : \text{Pol}_n(\mathbb{R}) \rightarrow \text{Pol}_n(\mathbb{R})$, e as forma lineares $I, \delta \in \text{Pol}_n(\mathbb{R})^*$ definidas por $\langle I, f \rangle = \int_0^1 f(t) dt$ e $\langle \delta, f \rangle = f(0)$, respetivamente. Calcule $D^\top I$ e $D^\top \delta$.

Dualidade. A relação de dualidade entre uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ e a sua transposta $L^\top : \mathbf{W}^* \rightarrow \mathbf{V}^*$ é refletida em certas relações entre os espaços nulos e as imagens de L e $L^\top$. A primeira observação é quase tautológica.

Teorema 9.2. O núcleo de $L^\top$ é o aniquilador da imagem de L , ou seja,

$$\text{Ker}(L^\top) = \text{Annih}(\text{Im}(L))$$

Demonstração. Se $\xi \in \text{Ker} L^\top$ então $\langle L^\top \xi, \mathbf{v} \rangle = \langle \xi, L\mathbf{v} \rangle = 0$ para todo $\mathbf{v} \in \mathbf{V}$. Isto significa que ξ anula todos os vetores $L\mathbf{v}$, ou seja, que ξ está no aniquilador da imagem de L .

Vice-versa, seja ξ uma forma no aniquilador de $\text{Im}(L)$. Então $\langle \xi, L\mathbf{v} \rangle = 0$ para todos $\mathbf{v} \in \mathbf{V}$. Mas então $\langle L^\top \xi, \mathbf{v} \rangle = \langle \xi, L\mathbf{v} \rangle = 0$ para todos $\mathbf{v} \in \mathbf{V}$. Isto significa que $L^\top \xi$ é a forma nula. $\square$

Se $\mathbf{V}$ e $\mathbf{W}$ têm dimensão finita, é então possível relacionar as dimensões dos núcleos de $L^\top$ e L . De fato, usando os teoremas 8.3 e 8.7, a dimensão do núcleo de $L^\top$ é

$$\begin{aligned} \dim(\text{Ker}(L^\top)) &= \dim(\text{Annih}(\text{Im}(L))) \\ &= \dim(\mathbf{W}) - \dim(\text{Im}(L)) \\ &= \dim(\mathbf{W}) - \dim(\mathbf{V}) + \dim(\text{Ker}(L)) \end{aligned}$$

e portanto

$$\dim(\mathbf{V}) - \dim(\text{Ker}(L)) = \dim(\mathbf{W}) - \dim(\text{Ker}(L^\top))$$

O seguinte teorema é o dual do 9.2 e a versão abstrata do teorema 9.1 sobre uma matriz e a sua transposta.

Teorema 9.3. Seja $L : \mathbf{V} \rightarrow \mathbf{W}$ uma transformação linear entre espaços de dimensão finita. Então

$$\text{Im}(L^\top) = \text{Annih}(\text{Ker}(L))$$

e

$$\dim(\text{Im}(L^\top)) = \dim(\text{Im}(L))$$

Demonstração. Se $\xi = L^* \eta$ é a imagem de uma forma $\eta \in \mathbf{W}^*$, e se $\mathbf{v} \in \text{Ker}(L)$, então

$$\langle \xi, \mathbf{v} \rangle = \langle L^* \eta, \mathbf{v} \rangle = \langle \eta, L\mathbf{v} \rangle = 0$$

Isto prova a inclusão $\text{Im}(L^\top) \subset \text{Annih}(\text{Ker}(L))$ (que portanto vale também em dimensão infinita). Para provar a outra inclusão basta provar que os dois espaços têm a mesma dimensão. Esta é uma consequência da segunda parte do teorema e dos teoremas 8.3 e 8.7. Usando os teoremas 8.3 e 9.2

$$\begin{aligned} \dim(\text{Im}(L^\top)) &= \dim(\mathbf{W}^\top) - \dim(\text{Ker}(L^\top)) \\ &= \dim(\mathbf{W}) - \dim(\text{Annih}(\text{Im}(L))) \\ &= \dim(\text{Im}(L)) \end{aligned}$$

$\square$

ex: Mostre que uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ entre espaços de dimensão finita é sobrejetiva sse a transformação dual $L^\top : \mathbf{W}^* \rightarrow \mathbf{V}^*$ é injetiva.

ex: Mostre que uma transformação linear $L : \mathbf{V} \rightarrow \mathbf{W}$ entre espaços de dimensão finita é injetiva sse a transformação dual $L^\top : \mathbf{W}^* \rightarrow \mathbf{V}^*$ é sobrejetiva.

Einstein's sum convention. It is often important to take care of the distinction between vectors and co-vectors, and then different types of tensors, as well as to shorten formulas and computations. One possibility is the convention introduced by Einstein. We denote the coordinates of vectors using upper indices as $\mathbf{x} = (x^i) \in \mathbb{R}^n$, and denote the coordinates of co-vectors using lower indices as $\boldsymbol{\xi} = (\xi_j) \in (\mathbb{R}^n)^*$. The pairing between a co-vector $\boldsymbol{\xi}$ and a vector $\mathbf{x}$ reads $\langle \boldsymbol{\xi}, \mathbf{x} \rangle = \xi_1 x^1 + \xi_2 x^2 + \dots + \xi_n x^n$, and is shortened as $\langle \boldsymbol{\xi}, \mathbf{x} \rangle = \xi_i x^i$, Einstein's convention being that a repeated index which appears once as an upper index and once as a lower index implies summation. Using Einstein's sum convention, the coordinates of the image of a linear transformation represented by the matrix $T = (t^i_j)$ are given by $y^i = t^i_j x^j$. The composition of $T = (t^i_j)$ followed by $S = (s^i_j)$ is then represented by the matrix $ST = (s^i_k t^k_j)$.

This notation is particularly useful in differential and Riemannian geometry, the language of Einstein's theory of gravitation, the general theory of relativity²⁴, where important objects are "tensors", as the metric itself $g_{\mu\nu}$ or Riemann curvature $R^\alpha_{\beta\gamma\delta}$. For example, the squared norm of a four-vector $\mathbf{x}$ reads $g_{\mu\nu} x^\mu x^\nu$, the Ricci curvature is the contraction $R_{\mu\kappa} = R^\nu_{\mu\nu\kappa}$ of the Riemann curvature, ...

²⁴C.W. Misner, K.S. Thorne and J.A. Wheeler, *Gravitation*, Freeman, 1973.

10 Operadores e matrizes quadradas

ref: [Ap69] Vol 2, 2.1-8 ; [La87] Ch. IV-V

10.1 Álgebra das matrizes quadradas

Endomorfismos e matrizes quadradas. Seja $\mathbf{V}$ um espaço vetorial de dimensão n , por exemplo o próprio $\mathbb{R}^n$. Fixada uma base (por exemplo, a base canônica de $\mathbb{R}^n$), o espaço linear $\text{End}(\mathbf{V}) := \text{Lin}(\mathbf{V}, \mathbf{V})$ dos *endomorfismos* de $\mathbf{V}$ é isomorfo ao espaço linear $\text{Mat}_{n \times n}(\mathbb{R})$ das matrizes “quadradas” $n \times n$ reais, os seus elementos sendo as transformações lineares

$$X \mapsto Y = AX$$

ou seja, $x_i \mapsto y_i = \sum_j a_{ij} x_j$, com $X \in \mathbb{R}^n$ e $A \in \text{Mat}_{n \times n}(\mathbb{R})$. O caso complexo é análogo.

Naturalmente, a transformação “identidade” $\mathbf{x} \mapsto \mathbf{x}$ é definida pela matriz identidade I , e a transformação “nula” $\mathbf{x} \mapsto \mathbf{0}$ é definida pela matriz nula 0 . Uma homotetia $\mathbf{x} \mapsto \lambda \mathbf{x}$ é definida pela matriz λI .

A composição de dois endomorfismos corresponde ao produto de matrizes, ou seja,

$$L_A \circ L_B = L_{AB}$$

Em particular, a k -ésima iterada $L^k = L \circ L \circ \dots \circ L$ (k vezes) do endomorfismo $L : X \mapsto AX$ é representada pela k -ésima *potência* de A , definida recursivamente por

$$A^0 = I \quad \text{e} \quad A^{k+1} = AA^k \quad \text{se } k \geq 0.$$

Um espaço linear munido de um produto bilinear é chamado *álgebra*. É o caso do espaço $\text{End}(\mathbf{V})$ munido da lei “composição”, e do espaço $\text{Mat}_{n \times n}(\mathbb{R})$ munido do produto linhas por colunas. Assim, a correspondência $A \mapsto L_A$ é de fato um isomorfismo entre álgebras.

Matrizes diagonais. A *diagonal* da matriz quadrada

$$A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{pmatrix}$$

é o conjunto ordenado dos elementos “diagonais” $a_{11}, a_{22}, \dots, a_{nn}$ (em vermelho na imagem).

Uma matriz quadrada é uma *matriz diagonal* se os elementos que não pertencem à diagonal são nulos, ou seja, se é da forma

$$\Lambda = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n) := \sum_{i=1}^n \lambda_i I_{ii} = \begin{pmatrix} \lambda_1 & 0 & 0 & \dots \\ 0 & \lambda_2 & 0 & \dots \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & \lambda_n \end{pmatrix}$$

O conjunto $\text{Diag}_n(\mathbb{R})$ das matrizes diagonais $n \times n$ reais é claramente um subespaço linear de $\text{Mat}_{n \times n}(\mathbb{R})$. Um cálculo elementar mostra que o produto de duas matrizes diagonais é também uma matriz diagonal, cujas entradas são os produtos das entradas dos fatores:

$$\begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix} \begin{pmatrix} \mu_1 & 0 & \dots & 0 \\ 0 & \mu_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \mu_n \end{pmatrix} = \begin{pmatrix} \lambda_1 \mu_1 & 0 & \dots & 0 \\ 0 & \lambda_2 \mu_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \mu_n \end{pmatrix}$$

Assim, $\text{Diag}_n(\mathbb{R})$ é uma subálgebra da álgebra das matrizes quadradas $n \times n$.

Em particular, calcular potências de uma matriz diagonal Λ é fácil: Λ^k é ainda uma matriz diagonal, com entradas iguais às potências λ_i^k das entradas de Λ . Isto permite calcular polinômios, logo séries de potências, e finalmente funções analíticas de matrizes diagonais ...

Matrizes nilpotentes e unipotentes. A matriz quadrada N é dita *nilpotente* se existe um inteiro $k \geq 1$ tal que $N^k = 0$. A matriz quadrada U é dita *unipotente* se a diferença $U - I$ é nilpotente, e portanto existe um inteiro k tal que $(U - I)^k = 0$. Típicas matrizes nilpotentes e unipotentes (mas não as únicas, naturalmente!) são matrizes do género

$$N = \begin{pmatrix} 0 & * & \dots & * \\ 0 & 0 & \dots & * \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix} \quad \text{e} \quad U = I + N = \begin{pmatrix} 1 & * & \dots & * \\ 0 & 1 & \dots & * \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$$

respetivamente, com entradas $n_{ij} = 0$ se $i \geq j$ as primeiras, e com entradas $u_{ii} = 1$ e $u_{ij} = 0$ se $i > j$ as segundas (as outras entradas sendo arbitrárias, logo representadas por $*$'s).

ex: Calcule as potências $A^0, A^1, A^2, A^3, \dots, A^k, \dots$ quando

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \quad A = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \quad A = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \quad A = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

$$A = \begin{pmatrix} 0 & 1 & 1 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix} \quad A = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{pmatrix} \quad A = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 2 & 1 \\ 0 & 0 & 3 \end{pmatrix}$$

ex: Determine as matrizes $A \in \text{Mat}_{2 \times 2}(\mathbb{R})$ tais que $A^2 = 0$.

ex: Determine as matrizes $A \in \text{Mat}_{2 \times 2}(\mathbb{R})$ tais que $A^2 = I$.

ex: Mostre que (a matriz que representa) uma projecção ortogonal $P : \mathbb{R}^n \rightarrow \mathbb{R}^n$ sobre um subespaço não trivial $\mathbb{R}^m \subset \mathbb{R}^n$ (por exemplo, a projecção $(x_1, \dots, x_n) \mapsto (x_1, x_2, \dots, x_m, 0, \dots, 0)$ com $m < n$) satisfaz $P^2 = P$.

ex: Determine a matriz que representa o operador derivação, definido por $(Df)(t) := f'(t)$, no espaço $\text{Pol}_n(\mathbb{R}) \approx \mathbb{R}^{n+1}$ dos polinómios de grau $\leq n$, por exemplo na base formada pelos monómios $1, t, t^2/2, t^3/6, \dots, t^n/n!$ (comece pelos casos $n = 2$ ou 3). Mostre que é nilpotente.

Comutador. A composição de transformações lineares, e portanto o produto de matrizes, não são comutativos! Ou seja, em geral, $AB \neq BA$. As matrizes quadradas $A, B \in \text{Mat}_{n \times n}(\mathbb{R})$ (e portanto os endomorfismos que representam), *comutam entre si/são permutáveis* se

$$AB = BA.$$

A obstrução é o *comutador*, definido por

$$[A, B] := AB - BA$$

O comutador é linear em cada variável, ou seja,

$$[A + B, C] = [A, C] + [B, C] \quad [\lambda A, B] = \lambda[A, B] \quad \dots$$

anti-simétrico, ou seja,

$$[A, B] = -[B, A]$$

e satisfaz a *identidade de Jacobi*

$$[[A, B], C] + [[B, C], A] + [[C, A], B] = 0$$

como é possível verificar com um cálculo elementar.

Em particular, cada matriz quadrada A define um operador $\mathcal{L}_A : \text{Mat}_{n \times n}(\mathbb{R}) \rightarrow \text{Mat}_{n \times n}(\mathbb{R})$ no espaço linear das matrizes quadradas, dado por

$$\mathcal{L}_A B := AB - BA \quad (10.1)$$

É claro que $\mathcal{L}_I$ é o operador nulo. A identidade de Jacobi diz então que

$$\mathcal{L}_A[B, C] = [\mathcal{L}_A B, C] + [B, \mathcal{L}_A C]$$

Esta fórmula lembra a regra de Leibniz, e portanto diz que $\mathcal{L}_A$ é uma “derivação” na álgebra das matrizes quadradas, se o produto é o comutador.

ex: Mostre que cada matriz quadrada A comuta com si própria, i.e. $[A, A] = 0$.

ex: Mostre que duas matrizes diagonais comutam.

ex: Mostre que um múltiplo λI da matriz identidade comuta com toda matriz.

ex: Considere as matrizes 2×2

$$E = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad E_+ = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \quad E_- = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}$$

Calcule os comutadores $[E, E_+]$, $[E, E_-]$ e $[E_+, E_-]$.

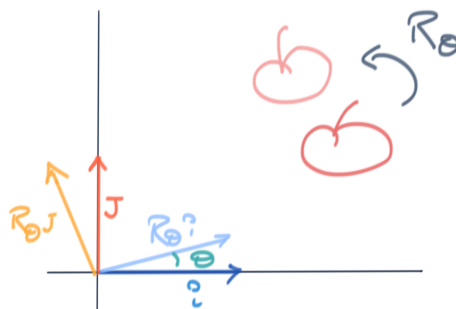
10.2 Exemplos geométricos

e.g. Rotações do plano. Uma rotação anti-horária de um ângulo θ envia o ponto do plano de coordenadas polares (ρ, φ) no ponto de coordenadas polares $(\rho, \varphi + \theta)$, e fixa a origem. Em particular, envia o vetor unitário $\mathbf{i}$ em $(\cos \theta, \sin \theta)$ e o vetor unitário $\mathbf{j}$ em $(-\sin \theta, \cos \theta)$. Envia o ponto genérico de coordenadas cartesianas $(r \cos \varphi, r \sin \varphi)$ no ponto de coordenadas cartesianas $(r \cos(\varphi + \theta), r \sin(\varphi + \theta))$. As fórmulas de adição das funções trigonométricas implicam que

$$\begin{pmatrix} r \cos(\varphi + \theta) \\ r \sin(\varphi + \theta) \end{pmatrix} = \begin{pmatrix} r (\cos \varphi \cos \theta - \sin \varphi \sin \theta) \\ r (\sin \varphi \cos \theta + \cos \varphi \sin \theta) \end{pmatrix} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} r \cos \varphi \\ r \sin \varphi \end{pmatrix}$$

e portanto que a rotação é uma transformação linear $R_\theta : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ definida, na base canónica, pela matriz

$$R_\theta = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}.$$



Em particular, uma rotação de um ângulo nulo, ou múltiplo inteiro de 2π , é a transformação identidade, definida pela matriz $R_0 = I$. As fórmulas de adição também mostram que as potências de uma rotação são rotações, e de fato

$$R_\theta^2 = R_{2\theta} \quad R_\theta^3 = R_{3\theta} \quad \dots \quad R_\theta^n = R_{n\theta}.$$

Também é imediato verificar que

$$R_\theta R_{-\theta} = R_{-\theta} R_\theta = I,$$

o seja, as rotações são invertíveis, e a inversa de uma rotação anti-horária de um ângulo θ é uma rotação anti-horária de um ângulo $-\theta$. Em geral, a composição de duas rotações de ângulos θ e ϕ é uma rotação de um ângulo $\theta + \phi$, ou seja,

$$R_\theta R_\phi = R_{\theta+\phi}.$$

Esta fórmula também mostra que as rotações do plano comutam, ou seja, $R_\theta R_\phi = R_\phi R_\theta$.

Um caso particular é a matriz

$$J = R_{\pi/2} = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$$

que define uma rotação de $\pi/2$. O seu quadrado é $J^2 = -I$, assim que J é uma “raiz quadrada” de $-I$. Se identificamos os pontos (x, y) do plano com os números complexos $z = x + iy$, a matriz J define a transformação $z \mapsto iz$, ou seja, uma multiplicação por i (um número tal que $i^2 = -1$).

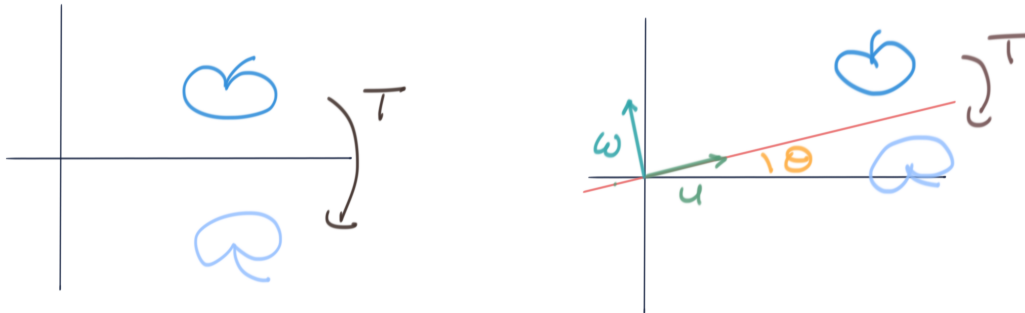
e.g. Reflexões no plano. Uma reflexão $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ ao longo de uma reta passando pela origem satisfaz a identidade $T^2 = I$, e portanto é uma involução. Por exemplo, as matrizes $\pm E$, com

$$E = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix},$$

definem reflexões nos eixos dos x e dos y , respetivamente. A reflexão numa reta genérica ℓ de equação cartesiana $y \cos \theta = x \sin \theta$ (ou seja, com declive $\tan \theta$) pode ser obtida como a composição

$$R_\theta E R_{-\theta} = \begin{pmatrix} \cos 2\theta & \sin 2\theta \\ \sin 2\theta & -\cos 2\theta \end{pmatrix},$$

ou seja, transformando a reta ℓ no eixo dos x , aplicando a reflexão no eixo dos x , e depois voltando a transformar o eixo dos x na reta ℓ .



Outro ponto de vista é considerar um vetor, por exemplo unitário, $\mathbf{u} = (\cos \theta, \sin \theta)$ que gera a reta, e observar que, se $\mathbf{v} = \lambda \mathbf{u} + \mathbf{w}$ é a decomposição de vetor genérico como soma de um vetor $\lambda \mathbf{u}$ proporcional a $\mathbf{u}$ e um vetor $\mathbf{w}$ ortogonal a $\mathbf{u}$, então a reflexão deve enviar

$$\mathbf{v} \mapsto T(\mathbf{v}) = \lambda \mathbf{u} - \mathbf{w}$$

Sendo $\lambda = \mathbf{u} \cdot \mathbf{v}$ (pois $\mathbf{u}$ é unitário) e $\mathbf{w} = \mathbf{v} - \lambda \mathbf{u}$, temos finalmente

$$T(\mathbf{v}) = 2(\mathbf{u} \cdot \mathbf{v})\mathbf{u} - \mathbf{v}$$

Em coordenadas cartesianas (associadas a uma base ortonormada), isto significa

$$\begin{aligned} T(x, y) &= 2(x \cos \theta + y \sin \theta)(\cos \theta, \sin \theta) - (x, y) \\ &= ((2 \cos^2 \theta - 1)x + (2 \sin \theta \cos \theta)y, (2 \cos \theta \sin \theta)x + (2 \sin^2 \theta - 1)y) \\ &= ((\cos 2\theta)x + (\sin 2\theta)y, (\sin 2\theta)x - (\cos 2\theta)y) \end{aligned}$$

e.g. Projeções ortogonais no plano. Uma projeção ortogonal $P : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ sobre um subespaço $V \subset \mathbb{R}^2$ satisfaz $P^2 = P$. Por exemplo, a projeção sobre a reta $y = 0$ é definida pela matriz

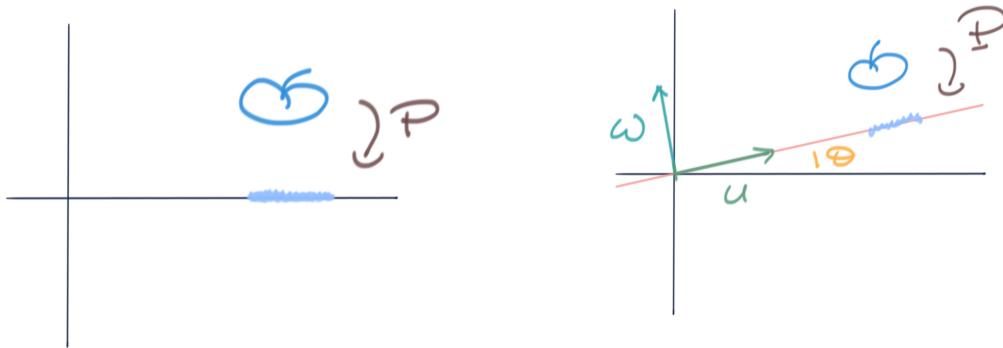
$$\begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}.$$

Em geral, a projeção sobre a reta passando pelo vetor unitário $\mathbf{u} = (\cos \theta, \sin \theta)$ envia $P\mathbf{u} = \mathbf{u}$, e é nula, ou seja, $P\mathbf{w} = 0$, sobre o vetor ortogonal $\mathbf{w} = (-\sin \theta, \cos \theta)$. Consequentemente, o seu valor sobre um vetor genérico $\mathbf{v} = (x, y)$ é dado por

$$P\mathbf{v} = (\mathbf{v} \cdot \mathbf{u})\mathbf{u} = (x \cos^2 \theta + y \cos \theta \sin \theta, x \cos \theta \sin \theta + y \sin^2 \theta)$$

e portanto a sua matriz na base canónica é

$$\begin{pmatrix} \cos^2 \theta & \cos \theta \sin \theta \\ \cos \theta \sin \theta & \sin^2 \theta \end{pmatrix}.$$



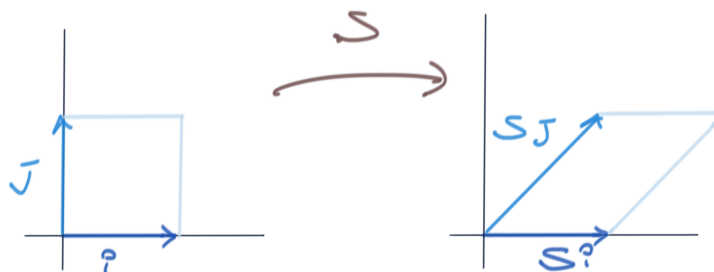
e.g. Deslizamentos. Um *deslizamento*, ou *cisalhamento* (em inglês, *shear*) horizontal é uma transformação do plano $(x, y) \mapsto (x + \alpha y, y)$, definida pela matriz

$$S_\alpha = \begin{pmatrix} 1 & \alpha \\ 0 & 1 \end{pmatrix},$$

com $\alpha \neq 0$. Em particular, o quadrado unitário de lados $\mathbf{i}$ e $\mathbf{j}$ é enviado no paralelogramo de lados $\mathbf{i}$ e $\alpha\mathbf{i} + \mathbf{j}$ (que tem a mesma área). É evidente que as potências de um cisalhamento são ainda cisalhamentos, e que de fato

$$S_\alpha^2 = \begin{pmatrix} 1 & 2\alpha \\ 0 & 1 \end{pmatrix} \quad S_\alpha^3 = \begin{pmatrix} 1 & 3\alpha \\ 0 & 1 \end{pmatrix} \quad \dots \quad S_\alpha^n = \begin{pmatrix} 1 & n\alpha \\ 0 & 1 \end{pmatrix}$$

Da mesma forma é possível definir cisalhamentos verticais.



e.g. Deslocamentos. Os *deslocamentos esquerdo* e *direito* são os operadores $L : \mathbb{R}^n \rightarrow \mathbb{R}^n$ e $R : \mathbb{R}^n \rightarrow \mathbb{R}^n$ definido por

$$L(x_1, x_2, \dots, x_n) = (x_2, x_3, \dots, x_n, 0) \quad \text{e} \quad R(x_1, x_2, \dots, x_n) = (0, x_1, x_2, \dots, x_{n-1})$$

respetivamente. Na base canónica, são definidos pelas matrizes nilpotentes

$$N = \begin{pmatrix} 0 & 1 & & & \\ & 0 & 1 & & \\ & & \ddots & \ddots & \\ & & & 0 & 1 \\ & & & & 0 \end{pmatrix} \quad \text{e} \quad N^T = \begin{pmatrix} 0 & & & & \\ 1 & 0 & & & \\ & \ddots & \ddots & & \\ & & 1 & 0 & \\ & & & 1 & 0 \end{pmatrix}$$

respetivamente. Observe que N é também a matriz que representa o operador derivação D no espaço $\text{Pol}_n(\mathbb{R})$ dos polinómios de grau $\leq n$ na base $1, t, t^2/2, \dots, t^n/n!$.

e.g. Translações na circunferência discreta. Também importante é o operador $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ definido por

$$T(x_1, x_2, \dots, x_n) = (x_2, x_3, \dots, x_n, x_1)$$

que pode ser pensado como uma translação no espaço das funções no grupo finito $\mathbb{Z}/n\mathbb{Z}$. Observe que $T^n = I$. Em particular, T é invertível e o operador inverso é $S = T^{n-1}$, definido por

$$S(x_1, x_2, \dots, x_n) = (x_n, x_1, x_2, \dots, x_{n-1})$$

A matriz que representa o operador T na base canónica é

$$C = \begin{pmatrix} 0 & 1 & & & \\ & 0 & 1 & & \\ & & \ddots & \ddots & \\ 0 & \dots & & 0 & 1 \\ 1 & 0 & \dots & & 0 \end{pmatrix}$$

Esta matriz, e as suas potências, são conhecidas como “matrizes circulantes”. Observe que C satisfaz $C^n = I$, assim que é uma raiz n -ésimas da matriz identidade.

10.3 Matrizes invertíveis

Operadores invertíveis em dimensão finita são representados por matrizes invertíveis.

Inversão de transformações do plano. A matriz 2×2

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

representa o endomorfismo genérico do plano $L_A(x, y) = (ax + by, cx + dy)$. A transformação L_A é invertível se para cada vetor $(\alpha, \beta) \in \mathbb{R}^2$ é possível encontrar um vetor $(x, y) \in \mathbb{R}^2$ tal que $L_A(x, y) = (\alpha, \beta)$, ou seja, resolver o sistema linear

$$\begin{cases} ax + by = \alpha \\ cx + dy = \beta \end{cases} \Rightarrow \begin{cases} (ad - bc)x = d\alpha - c\beta \\ (ad - bc)y = a\beta - c\alpha \end{cases}$$

(o segundo sistema é obtido ao retirar b vezes a segunda equação de d vezes a primeira equação, e depois ao retirar c vezes a primeira equação de a vezes a segunda equação). Portanto, a transformação L_A é invertível sse $\text{Det} A := ad - bc \neq 0$, e a sua inversa é a transformação linear

$$L_A^{-1}(\alpha, \beta) = \frac{1}{ad - bc}(d\alpha - c\beta, a\beta - c\alpha),$$

representada pela matriz

$$A^{-1} := \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}.$$

ex: Calcule a inversa das transformações lineares do plano definidas pelas matrizes

$$\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad \begin{pmatrix} 2 & 0 \\ 0 & 3 \end{pmatrix} \quad \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \quad \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$$

Automorfismos e matrizes invertíveis. A matriz quadrada $A \in \text{Mat}_{n \times n}(\mathbb{R})$ é *invertível* (ou *não-singular*, ou *regular*) se existe uma matriz quadrada $A^{-1} \in \text{Mat}_{n \times n}(\mathbb{R})$, dita *inversa* de A , tal que

$$A^{-1}A = AA^{-1} = I$$

A transformação linear $L_A : \mathbb{R}^n \rightarrow \mathbb{R}^n$, representada pela equação matricial $X \mapsto Y = AX$, é invertível sse a matriz A é invertível, e a sua inversa é a transformação linear $Y \mapsto X = A^{-1}Y$, definida pela matriz A^{-1} .

Se $A = (a_{ij})$ é invertível, então as entradas da inversa $A^{-1} = (b_{ij})$ satisfazem as n^2 equações lineares

$$\sum_{k=1}^n b_{ik} a_{kj} = \delta_{ij}$$

Se A e B são invertíveis (e têm a mesma dimensão), então também a matriz produto AB é invertível, e a sua inversa é

$$(AB)^{-1} = B^{-1}A^{-1}$$

De fato, $B^{-1}A^{-1}AB = B^{-1}B = I$, e $ABB^{-1}A^{-1} = AA^{-1} = I$. Por indução, a inversa de um produto finito $ABC \dots$ é o produto $(ABC \dots)^{-1} = \dots C^{-1}B^{-1}A^{-1}$. Se A é invertível então também a transposta $A^\top$ é invertível e

$$(A^\top)^{-1} = (A^{-1})^\top$$

De fato, $(A^{-1})^\top A^\top = (AA^{-1})^\top = I^\top = I$, e $A^\top (A^{-1})^\top = (AA^{-1})^\top = I$, pois $I^\top = I$.

Se $AX = B$ é um sistema de n equações com n incógnitas, e se $A \in \text{Mat}_{n \times n}(\mathbb{R})$ é uma matriz invertível, então o sistema admite uma solução única dada por $X = A^{-1}B$. Isto acontece quando o núcleo da transformação linear $X \mapsto AX$ é trivial, e portanto a característica da matriz (o número de linhas ou de colunas linearmente independentes) é n .

e.g. Matrizes diagonais invertíveis. Por exemplo, uma matriz diagonal

$$A = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix}$$

é invertível sse todos os λ_k são diferentes de zero, e a sua inversa é

$$\Lambda^{-1} = \begin{pmatrix} \lambda_1^{-1} & 0 & \dots & 0 \\ 0 & \lambda_2^{-1} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n^{-1} \end{pmatrix}.$$

e.g. Matrizes diagonais superiores invertíveis. Uma outra classe de matrizes invertíveis é importante na resolução algorítmica de sistemas lineares. Uma matriz quadrada $S = (s_{ij})$ é dita *diagonal superior* se é da forma

$$S = \begin{pmatrix} \lambda_1 & * & \dots & * \\ 0 & \lambda_2 & \dots & * \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix}$$

ou seja, se todos os elementos abaixo da diagonal forem nulos (ou seja, se $s_{ij} = 0$ quando $i > j$). Os $*$'s denotam escalares arbitrários. Se os termos diagonais $s_{kk} = \lambda_k$ forem todos diferentes de zero, esta matriz é claramente invertível. De fato, a equação $AX = B$ para o vetor coluna X (ou seja, o sistema linear nas incógnitas $x_1, x_2, \dots, x_n$) admite uma única solução para qualquer vetor coluna B . As coordenadas desta solução podem ser calculadas recursivamente a partir da última. De fato, a última equação, $\lambda_n x_n = b_n$ determina x_n . A penúltima equação, $\lambda_{n-1} x_{n-1} + s_{n-1,n} x_n = b_{n-1}$, determina x_{n-1} e assim a seguir.

Naturalmente, o mesmo é possível dizer para matrizes diagonais inferiores ...

ex: Mostre que, se $A, B \in \text{Mat}_{n \times n}(\mathbb{R})$, então $BA = I \Rightarrow AB = I$ (ou seja, uma inversa esquerda é também uma inversa direita, logo uma inversa).

ex: Diga se as seguintes matrizes são invertíveis e, caso afirmativo, calcule a inversa.

$$\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix}$$

$$\begin{pmatrix} 2 & 0 \\ 0 & 3 \end{pmatrix} \quad \begin{pmatrix} 2 & 1 \\ 0 & 3 \end{pmatrix} \quad \begin{pmatrix} 1 & -3 \\ -2 & 6 \end{pmatrix} \quad \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix}$$

$$\begin{pmatrix} 1 & 2 & 3 \\ 0 & 1 & 2 \\ 0 & 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 & 0 \\ 3 & 2 & 0 \\ 5 & 4 & 3 \end{pmatrix} \quad \begin{pmatrix} a & 0 & 0 \\ 0 & b & 0 \\ 0 & 0 & c \end{pmatrix} \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ a & 0 & 1 \end{pmatrix}$$

ex: Verifique que as rotações do plano são invertíveis, e determine a inversa da matriz R_θ que representa uma rotação anti-horária de um ângulo θ .

ex: Verifique que as reflexões do plano (relativamente a uma reta passando pela origem) são invertíveis. Calcule a inversa de uma matriz que representa uma reflexão.

ex: Existem infinitas matrizes quadradas 2×2 A tais que $A^2 = I$?

ex: As projeções são invertíveis?

ex: Seja N uma matriz (quadrada) nilpotente, ou seja, tal que $N^k = 0$ para algum inteiro minimal $k \geq 1$. Mostre que N não é invertível. Mostre que $I - N$ é invertível, e que a sua inversa é

$$(I - N)^{-1} = I + N + N^2 + \dots + N^{k-1}$$

(sugestão: multiplique por $I - N$...)

ex: [Ap69] Vol 2 2.20.

10.4 Mudança de bases e matrizes semelhantes

A escolha de um referencial conveniente pode ajudar na compreensão de um fenómeno. Um exemplo famoso é a diferença entre as órbitas dos planetas no sistema de Ptolomeu (descritas pelos gregos por meio de um “deferente” e um certo número de “epiciclos”) e as órbitas elípticas no sistema de Copérnico, dependendo se o referencial é solidário com a Terra e com o Sol. Da mesma forma, o estudo de uma transformação linear, por exemplo um operador, pode ser muito simplificado ao escolher umas bases oportunas.

5 jan 2023

Mudança de bases/coordenadas. Sejam $\mathcal{B} = (\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n)$ e $\mathcal{B}' = (\mathbf{b}'_1, \mathbf{b}'_2, \dots, \mathbf{b}'_n)$ duas bases ordenadas do espaço vetorial $\mathbf{V}$. Então existem duas matrizes quadradas $n \times n$ (reais ou complexas, dependendo se os espaços são reais ou complexos), $U = (u_{ij})$ e $U' = (u'_{ij})$, tais que

$$\mathbf{b}'_j = \sum_{i=1}^n u_{ij} \mathbf{b}_i \quad \text{e} \quad \mathbf{b}_j = \sum_{i=1}^n u'_{ij} \mathbf{b}'_i. \quad (10.2)$$

para todo $j = 1, 2, \dots, n$. Mas

$$\mathbf{b}'_j = \sum_{i=1}^n u_{ij} \mathbf{b}_i = \sum_{i=1}^n u_{ij} \left(\sum_{k=1}^n u'_{ki} \mathbf{b}'_k \right) = \sum_{k=1}^n \left(\sum_{i=1}^n u'_{ki} u_{ij} \right) \mathbf{b}'_k$$

e portanto, pela (10.2) e pela unicidade das coordenadas, $\sum_{i=1}^n u'_{ki} u_{ij} = \delta_{kj}$, ou seja, $U'U = I$. Isto significa que a matriz U é invertível, e que a sua inversa é

$$U' = U^{-1}$$

Uma caracterização mais abstrata dos coeficientes u_{ij} 's e u'_{ij} 's é a seguinte. Se $(\mathbf{b}_1^*, \mathbf{b}_2^*, \dots, \mathbf{b}_n^*)$ e $(\mathbf{b}'_1, \mathbf{b}'_2, \dots, \mathbf{b}'_n)$ denotam as bases ordenadas duais de $\mathcal{B}$ e $\mathcal{B}'$, respetivamente, então

$$u_{ij} = \langle \mathbf{b}_i^*, \mathbf{b}'_j \rangle \quad \text{e} \quad u'_{ij} = \langle \mathbf{b}'_i, \mathbf{b}_j \rangle$$

Então um vetor genérico $\mathbf{v} \in \mathbf{V}$ tem uma representação $\mathbf{v} = \sum_{i=1}^n x_i \mathbf{b}_i$ na base $\mathcal{B}$ e uma representação $\mathbf{v} = \sum_{i=1}^n x'_i \mathbf{b}'_i$ na base $\mathcal{B}'$. Para compreender a relação entre as coordenadas x_i 's e as coordenadas x'_i 's, basta observar que, pelas (10.2),

$$\mathbf{v} = \sum_{j=1}^n x_j \mathbf{b}_j = \sum_{j=1}^n \sum_{i=1}^n u'_{ij} x_j \mathbf{b}'_i$$

Consequentemente,

$$\boxed{x_i = \sum_{j=1}^n u_{ij} x'_j \quad \text{ou também} \quad x'_i = \sum_{j=1}^n u'_{ij} x_j} \quad (10.3)$$

pois U' é a matriz inversa de U .

A notação matricial é mais eficaz. A matriz

$$U = \begin{pmatrix} u_{11} & u_{12} & \dots & u_{1n} \\ u_{21} & u_{22} & & \\ \vdots & \vdots & \ddots & \vdots \\ u_{n1} & u_{n2} & \dots & u_{nn} \end{pmatrix}$$

cujas colunas são as coordenadas dos vetores da base $\mathcal{B}'$ relativamente à base $\mathcal{B}$, é a matriz que realiza a mudança de coordenadas. A matriz inversa $U' = U^{-1}$ tem como colunas as coordenadas

dos vetores da base $\mathcal{B}$ relativamente à base $\mathcal{B}'$. Se X e X' são os vetores coluna de coordenadas x_i 's e x'_i 's respetivamente, então a mudança de coordenadas (10.3) assume a forma

$$\boxed{X = U X' \quad \text{ou também} \quad X' = U^{-1} X} \quad (10.4)$$

As matrizes U e U^{-1} podem ser pensadas como matrizes das derivadas parciais, pois é claro, pelas (10.3), que $u_{ij} = \partial x_i / \partial x'_j$ e $v_{ij} = \partial x'_i / \partial x_j$. Nesta notação, as fórmulas (10.3) para a mudança de coordenadas parecem tautológicas (o que explica o valor da notação de Leibniz para as derivadas):

$$x'_i = \sum_{j=1}^n \frac{\partial x'_i}{\partial x_j} x_j \quad \text{e} \quad x_i = \sum_{j=1}^n \frac{\partial x_i}{\partial x'_j} x'_j.$$

Efeito sobre as matrizes que representam transformações lineares. Seja $L : \mathbf{V} \rightarrow \mathbf{W}$ uma transformação linear entre os espaços vetoriais de dimensão finita $\mathbf{V}$ e $\mathbf{W}$, reais ou complexos. Fixadas uma base ordenada $\mathcal{B} = (\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n)$ de $\mathbf{V}$ e uma base ordenada $\mathcal{C} = (\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_m)$ de $\mathbf{W}$, os vetores destes dois espaços são representados por vetores coluna $X = (x_1, x_2, \dots, x_n)^\top$ e $Y = (y_1, y_2, \dots, y_m)^\top$, respetivamente. A transformação L é então definida por uma matriz $m \times n$ $A = (a_{ij})$, e é dada por

$$X \mapsto Y = AX$$

ou seja, em coordenadas, por $x_j \mapsto y_i = \sum_{j=1}^n a_{ij} x_j$.

Sejam $\mathcal{B}' = (\mathbf{b}'_1, \mathbf{b}'_2, \dots, \mathbf{b}'_n)$ e $\mathcal{C}' = (\mathbf{c}'_1, \mathbf{c}'_2, \dots, \mathbf{c}'_m)$ umas outras bases de $\mathbf{V}$ e $\mathbf{W}$, respetivamente. De acordo com as (10.4), as coordenadas $X' = (x'_1, x'_2, \dots, x'_n)^\top$ e $Y' = (y'_1, y'_2, \dots, y'_m)^\top$ dos vetores dos dois espaços nas novas bases são dadas por $X' = U^{-1} X$ e $Y' = V^{-1} Y$, onde $U = (u_{ij})$ é a matriz invertível $n \times n$ tal que $\mathbf{b}'_j = \sum_{i=1}^n u_{ij} \mathbf{b}_i$ e $V = (v_{ij})$ é a matriz invertível $m \times m$ tal que $\mathbf{c}'_j = \sum_{i=1}^m v_{ij} \mathbf{c}_i$. Então

$$X \mapsto Y = AX \quad \Rightarrow \quad X' = U^{-1} X \mapsto Y' = V^{-1} Y = V^{-1} AX = V^{-1} A U X'.$$

Consequentemente, a matriz da transformação L relativamente às bases ordenadas $\mathcal{B}'$ e $\mathcal{C}'$ é

$$\boxed{A' = V^{-1} A U}$$

ou seja, a transformação linear L é definida, nas novas coordenadas x'_i 's e y'_i 's, por

$$X' \mapsto Y' = A' X'$$

O caso mais importante é o caso de um endomorfismo, ou seja, um operador $L : \mathbf{V} \rightarrow \mathbf{V}$. Se A é a matriz $n \times n$ do endomorfismo $L \in \text{End}(\mathbf{V})$ relativamente à base $\mathcal{B}$, então a matriz de L relativamente à base $\mathcal{B}'$ é

$$\boxed{A' = U^{-1} A U} \quad (10.5)$$

Obsem que então $A = U A' U^{-1}$. Matrizes quadradas A e A' (da mesma dimensão) relacionadas pela identidade (10.5), para alguma matriz invertível U , são ditas *semelhantes/similares* ou também *conjugadas*. Representam o mesmo endomorfismo em bases possivelmente diferentes. É um exercício verificar que a semelhança é uma relação de equivalência.

Conjugação. Cada matriz $n \times n$ invertível U define um operador $\mathcal{M}_U : \text{Mat}_{n \times n}(\mathbb{R}) \rightarrow \text{Mat}_{n \times n}(\mathbb{R})$, chamado *conjugação*, de acordo com

$$\mathcal{M}_U A := U A U^{-1} \quad (10.6)$$

É imediato verificar que $\mathcal{M}_I$ é o operador identidade, e que

$$\mathcal{M}_{U^{-1}} = (\mathcal{M}_U)^{-1}$$

Também, se U e V são invertíveis, então

$$\mathcal{M}_{UV} = \mathcal{M}_U \circ \mathcal{M}_V$$

e.g. Por exemplo, consideramos, no plano euclidiano, o problema de determinar a matriz A que define a reflexão $R : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ na reta $y = 2x$ relativamente à base canónica. Na base $\mathbf{e}'_1 = (1, 2)$ e $\mathbf{e}'_2 = (-2, 1)$, os vetores que geram a reta $y = 2x$ e a reta ortogonal, esta reflexão é muito mais simples, pois claramente envia $R\mathbf{e}'_1 = \mathbf{e}'_1$ e $R\mathbf{e}'_2 = -\mathbf{e}'_2$. Isto significa que a matriz A' que define esta reflexão nas coordenadas x', y' relativa a esta base é

$$A' = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

A matriz mudança de coordenadas, cujas colunas são os vetores da nova base relativamente à base canónica, é

$$U = \begin{pmatrix} 1 & -2 \\ 2 & 1 \end{pmatrix}$$

Portanto, pela (10.5),

$$A = UA'U^{-1} = \begin{pmatrix} 1 & -2 \\ 2 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} 1/5 & 2/5 \\ -2/5 & 1/5 \end{pmatrix} = \begin{pmatrix} -3/5 & 4/5 \\ 4/5 & 3/5 \end{pmatrix}$$

Finalmente, nas coordenadas x, y relativas à base canónica, esta reflexão é

$$R(x, y) = ((-3x + 4y)/5, (4x + 3y)/5)$$

ex: Verifique que se $A' = U^{-1}AU$ então $A = UA'U^{-1}$ e vice-versa.

ex: Verifique que a semelhança é uma relação de equivalência.

ex: Verifique que se A e A' são semelhantes e A é invertível então também A' é invertível.

ex: Verifique que se $B = U^{-1}AU$ então $B^k = U^{-1}A^kU$ para toda potência $k \geq 0$, e, se A é invertível, também para toda potência $k \in \mathbb{Z}$.

ex: Determine a matriz de $L(x, y) = (3x, 2y)$ relativamente à base $\mathbf{b}_1 = (1, 1)$ e $\mathbf{b}_2 = (1, -1)$.

ex: Seja $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ a reflexão na reta $y = x$. Determine a matriz de T relativamente à base canónica e relativamente à base $(1, 1)$ e $(1, -1)$.

ex: Seja $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ a reflexão na reta $y = \sqrt{3}x$. Determine a matriz de T relativamente à base canónica.

ex: Seja $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ a projecção ortogonal sobre a reta $x + y = 0$. Determine a matriz de T relativamente à base canónica.

ex: Determine a matriz que representa o operador derivação $D : \text{Pol}_3(\mathbb{R}) \rightarrow \text{Pol}_2(\mathbb{R})$, definido por $(Df)(t) := f'(t)$, relativamente às bases ordenadas $(1, t, t^2, t^3)$ e $(1, t, t^2)$. Determine umas bases de $\text{Pol}_3(\mathbb{R})$ e $\text{Pol}_2(\mathbb{R})$ tais que o operador D seja diagonal.

Polinómios de operadores e matrizes. Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador definido no espaço vetorial $\mathbf{V}$, real ou complexo. A cada polinómio $f(z) = a_0 + a_1z + a_2z^2 + \cdots + a_mz^m$ com coeficientes escalares (ou seja, reais ou complexos dependendo se o espaço é real ou complexo) numa incógnita z é possível associar um operador $f(L) : \mathbf{V} \rightarrow \mathbf{V}$, definido por

$$f(L) := a_0I + a_1L + a_2L^2 + \cdots + a_mL^m$$

onde I denota o operador identidade. Esta definição é bem posta porque todas as potências de L e os seus múltiplos comutam entre si. De fato, todos estes operadores $f(L)$, combinações lineares de potências de L , comutam.

É claro que se o polinómio $f(z)$ é uma soma $f(z) = p(z) + q(z)$ de dois polinómios, então $f(L) = p(L) + q(L)$. Também evidente é que se o polinómio $f(z)$ fatoriza num produto $f(z) = p(z)q(z)$ de dois polinómios, então também o operador $f(L)$ fatoriza no produto $f(L) = p(L)q(L)$, ou seja, na composição dos dois operadores $p(L)$ e $q(L)$.

Se $\mathbf{V}$ tem dimensão finita, e se a matriz quadrada A representa o operador L numa base fixada, então as mesmas considerações se aplicam aos polinómios $f(A)$, que são então matrizes quadradas que representam os operadores $f(L)$. Se mudamos base, e representamos o mesmo operador com a matriz semelhante $B = U^{-1}AU$, então também

$$f(B) = U^{-1}f(A)U$$

para todo polinómio $f(z)$, pois isto acontece com as potências de B .

Operadores diferenciais. Por exemplo, se $p(z) = az^2 + bz + c$ é um polinómio de grau dois e D é o operador derivação, definido por $(Df)(t) = f'(t)$ no espaço linear $\mathcal{C}^\infty(\mathbb{R})$, então o polinómio $p(D)$ é o operador diferencial $L = aD^2 + bD + cI$, que envia a função $f(t)$ em $(Lf)(t) = af''(t) + bf'(t) + cf(t)$. O núcleo deste operador é então o espaço das soluções da equação diferencial ordinária linear de segunda ordem

$$af''(t) + bf'(t) + cf(t) = 0$$

Da mesma forma, um polinómio $p(z)$ de grau m em D define um operador diferencial $p(D)$ de ordem m , e o seu núcleo é o espaço das soluções de uma equação diferencial ordinária linear de ordem m .

ex: Verifique que o núcleo de $D - \lambda$ é a reta dos exponenciais $f(t) = ce^{\lambda t}$.

10.5 Traço

Traço de uma matriz. Seja $A = (a_{ij})$ uma matriz quadrada $n \times n$, real ou complexa. O *traço* (em inglês, *trace*) de A é a soma dos elementos da diagonal, ou seja,

$$\text{Tr}(A) := a_{11} + a_{22} + \cdots + a_{nn}$$

Por exemplo,

$$\text{Tr} \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{pmatrix} = 1 + 5 + 9 = 15$$

É imediato verificar que o traço é uma função linear da matriz (ou seja, uma forma linear no espaço linear das matrizes quadradas), pois satisfaz

$$\text{Tr}(\lambda A) = \lambda \text{Tr}(A) \quad \text{Tr}(A + B) = \text{Tr}(A) + \text{Tr}(B)$$

se A e B são matrizes $n \times n$ e λ é um escalar. Outra propriedade elementar é que

$$\text{Tr}(A^\top) = \text{Tr}(A) \tag{10.7}$$

que pode ser verificada diretamente.

Menos óbvio é que se A e B são duas matrizes quadradas $n \times n$, então

$$\text{Tr}(AB) = \text{Tr}(BA) \tag{10.8}$$

embora os produtos AB e BA sejam, em geral, diferentes. De fato,

$$\text{Tr}(AB) = \sum_i \sum_k a_{ik} b_{ki} \quad \text{e} \quad \text{Tr}(BA) = \sum_i \sum_k b_{ik} a_{ki}$$

e estas duas expressões são claramente iguais. Observe que isto implica que traços de comutadores são nulos, ou seja, $\text{Tr}([A, B]) = 0$.

Traço de um operador. Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador definido no espaço vetorial $\mathbf{V}$, real ou complexo, de dimensão finita. Fixada uma base de $\mathbf{V}$, o operador é definido por uma matriz quadrada A . De acordo com a (10.5), uma mudança de base transforma a matriz A na matriz semelhante $A' = U^{-1}AU$, onde U é uma matriz invertível. Pela (10.8),

$$\text{Tr}(A') = \text{Tr}(U^{-1}AU) = \text{Tr}(AUU^{-1}) = \text{Tr}(A)$$

ou seja, o traço de uma matriz depende apenas da sua classe de semelhança. É portanto possível definir o *traço do operador* L como $\text{Tr}(L) := \text{Tr}(A)$, onde A é a matriz que representa L numa base arbitrária. Uma definição intrínseca (e que generaliza em dimensão infinita), de acordo com a (9.3), é então a seguinte. Seja $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ uma base arbitrária de $\mathbf{V}$, e seja $\mathbf{e}_1^*, \mathbf{e}_2^*, \dots, \mathbf{e}_n^*$ a sua base dual de $\mathbf{V}^*$. O traço do operador $L : \mathbf{V} \rightarrow \mathbf{V}$ é

$$\text{Tr}(L) = \sum_{i=1}^n \langle \mathbf{e}_i^*, L\mathbf{e}_i \rangle$$

Para compreender o significado geométrico do traço (relacionado a volumes e derivadas) é preciso esperar a introdução do “determinante” e a definição/cálculo do “exponencial de uma matriz”. Isto será feito na disciplina de Complementos de Cálculo e de Geometria Analítica.

ex: Determine o traço da matriz ou do operador identidade num espaço linear de dimensão n .

ex: Determine o traço da homotetia $\mathbf{x} \mapsto \lambda \mathbf{x}$ em $\mathbb{R}^n$.

ex: Mostre que o traço de uma projeção P é igual à ordem de P , ou seja, a dimensão da imagem $\text{Im}(P)$.

ex: Mostre que o traço de um produto finito de matrizes quadradas não depende da permutação cíclica dos fatores, ou seja, $\text{Tr}(ABC \dots Z) = \text{Tr}(BC \dots ZA) = \text{Tr}(C \dots ZAB) = \dots$

ex: Mostre que, por outro lado, $\text{Tr}(ABC)$ pode ser diferente de $\text{Tr}(BAC)$.

ex: Mostre que o traço de uma matriz (quadrada) anti-simétrica é nulo.

ex: Mostre que não existem duas matrizes quadradas A e B tais que o comutador seja um múltiplo não nulo da identidade, ou seja, $[A, B] = \lambda I$ com $\lambda \neq 0$ (calcule o traço do comutador ...). Os operadores posição Q e momento P na mecânica quântica satisfazem $[Q, P] = i\hbar I \dots$

Produto escalar de Frobenius. O traço permite definir um produto escalar, logo uma norma, interessante nos espaços das matrizes, não necessariamente quadradas. Sejam $A = (a_{ij})$ e $B = (b_{ij})$ duas matrizes $m \times n$, reais ou complexas. Então

$$\langle A, B \rangle_F := \text{Tr}(\bar{A}^\top B) = \sum_{i=1}^m \sum_{j=1}^n \bar{a}_{ij} b_{ij}$$

Pela (10.7) esta expressão é “hermítica”, ou seja, satisfaz

$$\langle A, B \rangle_F = \overline{\langle B, A \rangle_F}$$

Em particular, é simétrica se A e B são matrizes reais. É também definida positiva, pois

$$\langle A, A \rangle_F = \sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2$$

é uma soma de quadrados, logo não negativa e nula sse A é a matriz nula. Consequentemente, $\|A\|_F := \sqrt{\langle A, A \rangle_F}$ é uma norma no espaço linear das matrizes $m \times n$, chamada *norma de Frobenius*. Quando $m = 1$, este é o produto escalar hermítico usual entre vetores de $\text{Mat}_{1,n}(\mathbb{C}) \approx \mathbb{C}^n$.

11 Sistemas lineares

ref: [Ap69] Vol 2, 2.17-18 ; [La97] Ch. II

11.1 Sistemas lineares no plano e no espaço

Peppermint Patty's problems.

24 nov 2022



Copyright © 2002 United Feature Syndicate, Inc.

ex: “In driving from town A to town D you pass first through town B and then through town C. It is 10 miles farther from A to B than from B to C and 10 miles farther from B to C than from C to D. If it is 390 miles from A do D, how far is it from A to B?”²⁵

ex: “A man has a daughter and a son.. The son is three years older than the daughter ... In one year the man will be six times as old as the daughter is now, and in ten years he will be fourteen years older than the combined ages of his children ... What is the man's present age?”

ex: “A man has twenty coins consisting of dimes and quarters²⁶ ... If the dime were quarters and the quarters were dimes, he would have ninety cents more than he has now ... How many dimes and quarters does he have?”

e.g. Equações lineares na reta. Uma equação linear

$$ax = b$$

na reta real $\mathbb{R}$ (ou na reta complexa $\mathbb{C}$, ou, em geral, num corpo), com $a \neq 0$ (caso contrário é apenas a afirmação $b = 0$), admite uma única solução $x = b/a$.

e.g. Equações lineares no plano. Uma equação linear

$$ax + by = c$$

no plano cartesiano $\mathbb{R}^2$, com $\mathbf{n} = (a, b) \neq (0, 0)$, define uma reta afim $R \subset \mathbb{R}^2$. A equação homogênea associada

$$ax + by = 0$$

define uma reta que passa pela origem, ou seja, um subespaço vetorial $\mathbf{n}^\perp = \mathbb{R}\mathbf{v} \subset \mathbb{R}^2$ de dimensão 1 (por exemplo, com $\mathbf{v} = (b, -a)$). Se $\mathbf{r}_0 = (x_0, y_0)$ é um ponto de R , ou seja, (apenas) uma solução de $ax + by = c$, então o espaço de todas as soluções é $R = \mathbf{r}_0 + \mathbb{R}\mathbf{v}$. Ou seja, as soluções de $ax + by = c$ são dadas por

$$(x, y) = \mathbf{r}_0 + t\mathbf{v} = (x_0, y_0) + t(b, -a)$$

ao variar o parâmetro $t \in \mathbb{R}$.

²⁵Peppermint Patty, in *Peanuts*, by Charles M. Schulz, December 6th, 1968.

²⁶A *dime* is a 10 cents coin, and a *quarter* is a 25 cents coin.

Um sistema de duas equações lineares

$$\begin{cases} ax + by = c \\ a'x + b'y = c' \end{cases}$$

descreve a interseção $(R \cap R') \subset \mathbb{R}^2$ entre duas retas afins R e R' . Esta interseção pode ser vazia (retas paralelas e distintas), pode ser uma reta $ax + by = c$ (equações proporcionais/equivalentes), ou pode ser um único ponto.

A última possibilidade é o caso genérico, e acontece quando os vetores normais $\mathbf{n} = (a, b)$ e $\mathbf{n}' = (a', b')$ são linearmente independentes. A menos de reordenar as equações, podemos assumir que $a \neq 0$. Então o sistema é equivalente (eliminando x na segunda equação) ao sistema “em escada de linhas”

$$\begin{cases} ax + by = c \\ b''y = c'' \end{cases}$$

com $b'' = b' - a'b/a$ e $c'' = c' - a'c/a$, e portanto (pois no caso genérico também $b'' \neq 0$) ao sistema “diagonal”

$$\begin{cases} x = \alpha \\ y = \beta \end{cases}$$

com $\alpha = (c - b\beta)/a$ e $\beta = c''/b''$. A única solução é neste caso o ponto (α, β) .

ex: Verifique que a solução do sistema genérico é o ponto (α, β) de coordenadas

$$\alpha = \frac{b'c - bc'}{ab' - a'b} \quad \beta = \frac{ac' - a'c}{ab' - a'b}$$

Interprete estas fórmulas usando determinantes de matrizes 2×2 .

ex: Resolva, se possível, os seguintes sistemas lineares

$$\begin{cases} x + y = 0 \\ x - y = 0 \end{cases} \quad \begin{cases} x + y = 20 \\ x - y = 2 \end{cases} \quad \begin{cases} 2x - y = 13 \\ -x + 2y = 7 \end{cases}$$

e.g. Equações lineares no espaço. Uma equação linear

$$ax + by + cz = d$$

no espaço $\mathbb{R}^3$, com $\mathbf{n} = (a, b, c) \neq (0, 0, 0)$, define um plano afim $P = \{ \mathbf{n} \cdot \mathbf{r} = d \} \subset \mathbb{R}^3$. A equação homogênea associada

$$ax + by + cz = 0$$

define o supespaço vetorial $\mathbf{n}^\perp \subset \mathbb{R}^3$.

Um sistema de duas equações lineares

$$\begin{cases} ax + by + cz = d \\ a'x + b'y + c'z = d' \end{cases}$$

descreve a interseção $(P \cap P') \subset \mathbb{R}^3$ entre dois planos afins P e P' . Esta interseção pode ser vazia (dois planos paralelos e distintos), pode ser um plano $ax + by + cz = d$ (duas equações proporcionais/equivalentes), ou pode ser uma reta. A última possibilidade é o caso genérico, e o sistema é equivalente (eliminando x na segunda equação, se $a \neq 0$) ao sistema “em escada de linhas”

$$\begin{cases} ax + by + cz = d \\ b''y + c''z = d'' \end{cases}$$

A última variável pode ser pensada como um parâmetro $z = t$ da reta:

$$t \mapsto (\alpha t + \gamma, \beta t + \delta, t).$$

Um sistema de três equações lineares

$$\begin{cases} ax + by + cz &= d \\ a'x + b'y + c'z &= d' \\ a''x + b''y + c''z &= d'' \end{cases}$$

descreve a interseção $(P \cap P' \cap P'') \subset \mathbb{R}^3$ entre três planos afins P , P' e P'' . Esta interseção pode ser vazia (dois planos paralelos e distintos, ou um plano paralelo à reta de interseção entre os outros dois), pode ser um plano $ax + by + cz = d$ (equações proporcionais/equivalentes), pode ser uma reta (sistema equivalente a um sistema de duas equações), ou pode ser um único ponto. A última possibilidade é o caso genérico, e o sistema é equivalente (eliminando x na segunda e na terceira equação, se $a \neq 0$, e depois y na terceira, se $b''' \neq 0$) ao sistema “em escada de linhas”

$$\begin{cases} ax + by + cz &= d \\ b'''y + c'''z &= d''' \\ c'''z &= d'''' \end{cases}$$

com $a \neq 0$, $b''' \neq 0$ e $c''' \neq 0$, e portanto ao sistema “diagonal”

$$\begin{cases} x &= \alpha \\ y &= \beta \\ z &= \gamma \end{cases}$$

com $\gamma = d''''/c'''$, $\beta = (d''' - c'''\gamma)/b'''$ e $\alpha = (d - c\gamma - b\beta)/a$. A única solução é neste caso o ponto (α, β, γ) .

ex: Resolva, se possível, os seguintes sistemas lineares

$$\begin{cases} 3x - y &= 0 \\ x + y + z &= 1 \end{cases} \quad \begin{cases} x + y + z &= 1 \\ x + y - z &= 0 \end{cases} \quad \begin{cases} x &= 3 \\ x + y &= 2 \\ x + y + z &= 1 \end{cases}$$

11.2 Sistemas lineares

Para tratar sistemas de um número grande de equações num número arbitrário de incógnitas é necessário introduzir uma notação conveniente.

Sistemas lineares. Um *sistema* de m equações lineares nas *incógnitas* $x_1, x_2, \dots, x_n$ é um conjunto de m equações lineares

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= b_1 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n &= b_2 \\ \vdots &\vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= b_m \end{cases} \quad (11.1)$$

com a_{ij} e b_k números reais (ou complexos). A matriz $A = (a_{ij}) \in \text{Mat}_{m \times n}(\mathbb{R})$ é dita *matriz dos coeficientes* do sistema, e os números b_k são chamados *termos independentes*. Em notação matricial, o sistema é

$$AX = B$$

se $B = (b_1, b_2, \dots, b_m)^\top \in \mathbb{R}^m$ e $X = (x_1, x_2, \dots, x_n)^\top \in \mathbb{R}^n$ são vetores coluna.

O *sistema homogêneo* correspondente/associado é o sistema

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= 0 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n &= 0 \\ \vdots &\vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= 0 \end{cases} \quad (11.2)$$

onde todos os termos independentes são nulos, ou seja, em notação matricial,

$$AX = 0$$

Uma *solução* do sistema linear é um n -úpla $(x_1, x_2, \dots, x_n)$ de números, ou seja, um vetor $\mathbf{x} \in \mathbb{R}^n$, que satisfazem as m equações (11.1). Um sistema linear pode ter uma solução única (sistema *possível e determinado*), ter uma família (uma reta afim, um plano afim, ...) de soluções (sistema *possível e indeterminado*), ou não ter nenhuma solução (sistema *impossível*). O sistema homogêneo (11.2) admite pelo menos a *solução trivial* $(0, 0, \dots, 0)$.

Soluções de um sistema linear. Também útil é o ponto de vista “funcional”. A matriz $A = (a_{ij}) \in \text{Mat}_{m \times n}(\mathbb{R})$ dos coeficientes de um sistema linear define uma transformação linear $L_A : \mathbb{R}^n \rightarrow \mathbb{R}^m$, que em notação matricial envia $X \mapsto Y = AX$. O sistema linear (11.1) é portanto equivalente à equação linear

$$L_A(\mathbf{x}) = \mathbf{b}$$

onde $\mathbf{b} = (b_1, b_2, \dots, b_m) \in \mathbb{R}^m$. Uma solução é portanto um vetor $\mathbf{x} \in \mathbb{R}^n$ cuja imagem pela transformação L_A é o vetor $\mathbf{b}$. O espaço das soluções é a imagem inversa $L_A^{-1}(\{\mathbf{b}\})$.

Do ponto de vista “geométrico”, o vetor $\mathbf{x} = (x_1, x_2, \dots, x_n)^\top \in \mathbb{R}^n$ é solução do sistema $AX = B$ se

$$x_1 \begin{pmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{m1} \end{pmatrix} + x_2 \begin{pmatrix} a_{12} \\ a_{22} \\ \vdots \\ a_{m2} \end{pmatrix} + \dots + x_n \begin{pmatrix} a_{1n} \\ a_{2n} \\ \vdots \\ a_{mn} \end{pmatrix} = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{pmatrix}$$

ou seja, se

$$x_1 A_{*1} + x_2 A_{*2} + \dots + x_n A_{*n} = B$$

onde $A_{*j} = (a_{1j}, a_{2j}, \dots, a_{mj})^\top \in \mathbb{R}^m$ é a j -ésima coluna da matriz A . Portanto, o sistema admite (pelo menos) uma solução (ou seja, é *possível*) sse B é uma combinação linear das colunas da matriz A , ou seja, sse $B \in \text{Span}(A_{*1}, A_{*2}, \dots, A_{*n})$. Em termos abstratos,

Teorema 11.1. *O sistema $L_A(\mathbf{x}) = \mathbf{b}$ é possível, ou seja, admite (pelo menos) uma solução, sse $\mathbf{b} \in \text{Im}(L_A)$.*

A dimensão da imagem da transformação linear L_A , ou seja, o número de colunas linearmente independentes da matriz A , é $r = \text{Rank}(A) = \dim(\text{Im}(L_A))$, a ordem de L_A . Podemos acrescentar à matriz A uma coluna formada pelo vetor B , e considerar a “matriz aumentada”

$$(A|B) := \left(\begin{array}{cccc|c} a_{11} & a_{12} & \dots & a_{1n} & b_1 \\ a_{21} & a_{22} & \dots & a_{2n} & b_2 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} & b_m \end{array} \right)$$

O teorema anterior diz então

Teorema 11.2 (Kronecker-Capelli). *O sistema $AX = B$ é possível sse a característica da matriz aumentada $(A|B)$ é igual a característica da matriz A .*

O sistema linear homogêneo (11.2) é equivalente à equação homogênea

$$L_A(\mathbf{x}) = 0$$

Por definição, as suas soluções são os vetores do núcleo $\text{Ker}(L_A) \subset \mathbb{R}^n$, que é um subespaço de dimensão $k = \dim(\text{Ker}(L_A))$, a nulidade de L_A . Do ponto de vista geométrico, o vetor $X = (x_1, x_2, \dots, x_n)^\top \in \mathbb{R}^n$ é solução do sistema homogêneo se

$$A_{1*} \cdot X = 0 \quad A_{2*} \cdot X = 0 \quad \dots \quad A_{m*} \cdot X = 0,$$

onde $A_{i*} = (a_{i1}, a_{i2}, \dots, a_{in}) \in \mathbb{R}^n$ é a i -ésima linha da matriz A , ou seja, se é ortogonal ao espaço vetorial $\text{Span}(A_{1*}, A_{2*}, \dots, A_{m*})$ gerado pelas linhas de A . A dimensão do espaço das soluções do sistema homogêneo é portanto também igual a $k = n - r'$, se r' é o número de linhas linearmente independentes de A . Pelo teorema 8.7, $k + r = n$, assim que o número de linhas linearmente independentes de A é igual a $r' = r = \text{Rank}(A)$.

Se X e X' são soluções do sistema $AX = B$, então a diferença $Z = X - X'$ é solução do sistema homogêneo $AZ = 0$. Vice-versa, se X é uma solução do sistema $AX = B$ e Z é uma solução do sistema homogêneo $AZ = 0$ então $X' = X + Z$ é também uma solução do sistema linear $AX' = B$. Portanto, o conjunto das soluções de um sistema linear possível é caracterizado pelo seguinte “princípio de sobreposição”

Teorema 11.3 (princípio de sobreposição). *Se $\mathbf{x}$ é uma (apenas uma!) das soluções do sistema linear possível $L_A(\mathbf{x}) = \mathbf{b}$, então o espaço d(e todas)as soluções é o subespaço afim*

$$\mathbf{x} + \text{Ker}(L_A)$$

A maneira tradicional de enunciar o princípio de sobreposição é a seguinte: “a solução geral de uma equação linear é obtida somando a uma solução particular a solução geral da equação homogênea associada”. Ou seja, se X é uma solução do sistema possível $AX = B$, e se os vetores $F_1, F_2, \dots, F_k$ formam uma base de $\text{Ker}(L_A)$, assim que satisfazem o sistema homogêneo $AF_i = 0$, então a “solução geral” do sistema é

$$X + t_1 F_1 + t_2 F_2 + \dots + t_k F_k$$

com $t_1, t_2, \dots, t_k$ parâmetros reais.

A solução do sistema possível $AX = B$ é única quando $k = 0$, ou seja, quando o núcleo de L_A é trivial e portanto $\text{Rank}(A) = n$. Neste caso a transformação linear $L_A : \mathbb{R}^n \rightarrow \text{Im}(L_A)$ admite uma inversa $L_A^{-1} : \text{Im}(L_A) \rightarrow \mathbb{R}^n$, e a solução é dada por $X = L_A^{-1}(B)$.

Particularmente importante é o caso em que $m = n$, de um sistema linear $AX = B$ definido por uma matriz quadrada $A \in \text{Mat}_{n \times n}(\mathbb{R})$, equivalente ao problema $L_A(\mathbf{x}) = \mathbf{b}$ para o operador $L_A : \mathbb{R}^n \rightarrow \mathbb{R}^n$. Neste caso L_A é sobrejetiva sse é injetiva, logo sse é invertível. Consequentemente,

Teorema 11.4. *Seja $L_A : \mathbb{R}^n \rightarrow \mathbb{R}^n$ um operador. O sistema linear $L_A(\mathbf{x}) = \mathbf{b}$ admite uma (única) solução para cada $\mathbf{b} \in \mathbb{R}^n$ sse o sistema homogêneo $L_A(\mathbf{x}) = 0$ apenas admite a solução trivial, ou seja, sse $\text{Ker}(L_A) = 0$.*

Na prática, isto significa que $AX = B$ admite uma solução para cada B sse a matriz quadrada A é invertível, e neste caso a solução é

$$X = A^{-1}B$$

Também útil é reformular este teorema no seguinte princípio geral sobre problemas lineares, chamado “alternativa de Fredholm” (importante para certos operadores definidos em certos espaços de dimensão infinita chamados “espaços de Banach”). Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador definido no espaço linear $\mathbf{V}$, e seja $L^\top : \mathbf{V}^* \rightarrow \mathbf{V}^*$ o operador transposto. Dado $\mathbf{b} \in \mathbf{V}$, ou o sistema linear não homogêneo $L\mathbf{v} = \mathbf{b}$ admite uma solução $\mathbf{v}$, ou existe uma forma $\xi \in \mathbf{V}^*$ tal que $\langle \xi, \mathbf{b} \rangle \neq 0$ (logo não nula) e $L^*\xi = 0$. De fato, se $L\mathbf{v} = \mathbf{b}$ e se $\langle \xi, \mathbf{b} \rangle \neq 0$ então $\langle L^*\xi, \mathbf{v} \rangle = \langle \xi, L\mathbf{v} \rangle = \langle \xi, \mathbf{b} \rangle \neq 0$, e isto significa que $L^*\xi$ não é a forma nula. Consequentemente,

Teorema 11.5 (alternativa de Fredholm). *Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador, e $L^\top : \mathbf{V}^* \rightarrow \mathbf{V}^*$ o operador transposto. Então,*

ou o problema não homogêneo $L\mathbf{v} = \mathbf{b}$ admite uma solução para todo $\mathbf{b} \in \mathbf{V}$,

ou o problema homogêneo adjunto $L^\top \xi = 0$ admite uma solução não nula.

Em particular, dada uma matriz quadrada A , o sistema $AX = B$ admite soluções para todo B sse o sistema homogêneo $A^\top \Xi = 0$ admite apenas a solução nula.

ex: Estude os seguintes sistemas (ou seja, diga se são possíveis e, caso afirmativo, determine o espaço das soluções)

$$\begin{cases} 2x + y = -1 \\ x - y = 3 \end{cases} \quad \begin{cases} x + y = 11 \\ x - y = 33 \end{cases} \quad \begin{cases} 2x - 3y = -1 \\ -6x + 9y = 0 \end{cases}$$

$$\begin{cases} x + y - z = 1 \end{cases} \quad \begin{cases} 2x - 5y + 4z = -3 \\ x - 2y + z = 5 \\ x - 4y + 6z = 10 \end{cases} \quad \begin{cases} 3x + y - 10z = 1 \\ -2x - 5y + 7z = 2 \\ x + 3y - z = 0 \end{cases}$$

ex: [Ap69] 16.20.

11.3 Eliminação de Gauss-Jordan

É importante dispor de um algoritmo que resolva, quando possível, um sistema linear.

Eliminação de Gauß-Jordan. Consideramos o sistema linear $AX = B$ de m equações em n incógnitas

$$\begin{aligned} a_{11}x_1 + a_{12}x_2 + \cdots + a_{1n}x_n &= b_1 \\ a_{21}x_1 + a_{22}x_2 + \cdots + a_{2n}x_n &= b_2 \\ \vdots &\vdots \\ a_{m1}x_1 + a_{m2}x_2 + \cdots + a_{mn}x_n &= b_m \end{aligned}$$

O método de eliminação de Gauß-Jordan consiste em efectuar uma série de operações algébricas sobre as equações, e portanto sobre as linhas da “matriz aumentada”

$$(A|B) := \left(\begin{array}{cccc|c} a_{11} & a_{12} & \cdots & a_{1n} & b_1 \\ a_{21} & a_{22} & \cdots & a_{2n} & b_2 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} & b_m \end{array} \right)$$

até chegar a um sistema equivalente, ou seja, com as mesmas soluções, mais simples.

As operações permitidas, chamadas *operações elementares*, são as seguintes:

EG1 trocar duas equações,

EG2 multiplicar (todos os termos de) uma equação por um escalar não nulo $\lambda \neq 0$,

EG3 somar a uma equação um múltiplo α de outra equação.

Estas operações são reversíveis, e as operações inversas são também operações elementares (a inversa de EG1 é a própria, a inversa de EG2 com escalar λ é uma EG2 com escalar λ^{-1} , e a inversa de EG3 com escalar α consiste numa EG3 com escalar $-\alpha$). É claro portanto que não alteram o espaço das soluções. Dois sistemas lineares, $AX = B$ e $A'X = B'$, obtidos um do outro por meio de operações elementares são ditos *equivalentes*. As matrizes A e A' são também chamadas “equivalentes”.

O objetivo do método de eliminação de Gauss é efetuar uma série de operações elementares até obter um sistema equivalente $A'X = B'$, com A' “matriz em escada de linhas” (ou “escalorada”, em inglês, *row echelon form*), ou seja, da forma

$$A' = \left(\begin{array}{cccccc} \star & \star & \star & \star & \star & \cdots & \star \\ 0 & \star & \star & \star & \star & \cdots & \star \\ 0 & 0 & 0 & \star & \star & \cdots & \star \\ 0 & 0 & 0 & 0 & \star & \cdots & \star \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{array} \right)$$

onde os “pivots” $\star$ (ou “leading coefficients”) são os elementos $\neq 0$ mais à esquerda de cada linha não nula, e o pivot de cada linha é situado mais à direita do pivot da linha superior. As

linhas nulas são as últimas. Este processo, que transforma a matriz A do sistema numa matriz em escada de linhas A' , é chamado *eliminação de Gauss*, ou também *condensação*. É claro que é sempre possível.

Também a matriz A' é dita equivalente à matriz A , e as operações elementares podem ser interpretadas como operações sobre as linhas de A . É evidente que as operações elementares não mudam a característica $r = \text{Rank}(A)$ de uma matriz (pois não alteram o espaço das soluções da equação, e em particular da equação homogênea, que é um subespaço de dimensão $k = n - r$). Consequentemente, a característica da matriz A é igual à característica da matriz em escada de linhas A' , que é claramente igual ao número de linhas não nulas, ou seja ao número dos pivots. Em particular, o método de eliminação de Gauss pode ser usado como um algoritmo para calcular a característica de uma matriz, ou seja, como um algoritmo para calcular o número de vetores independentes numa família finita de vetores, as linhas de A .

É finalmente possível mudar a ordem das variáveis, assim que a matriz A' assume, na base reordenada, a forma

$$A' = \begin{pmatrix} \star & \star & \star & \star & \star & \dots & \star \\ 0 & \star & \star & \star & \star & \dots & \star \\ 0 & 0 & \star & \star & \star & \dots & \star \\ 0 & 0 & 0 & \star & \star & \dots & \star \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

A partir desta matriz em escada de linhas, a solução, ou a família paramétrica de soluções, pode ser facilmente calculada como nos exemplos anteriores em dimensão um, dois e três.

Se o número n das incógnitas é superior à característica r da matriz (ou seja, ao número dos pivots), então as r variáveis que correspondem aos pivots são funções das restantes $k = n - r$ variáveis, os parâmetros das soluções.

O caso mais importante é quando o número n das incógnitas é igual ao número m das equações, e a característica da matriz é $r = n$, assim que a solução é única. A matriz escada de linhas A' é então uma matriz quadrada diagonal superior, com entradas diagonais a'_{kk} diferentes de zero (os pivots). A matriz aumentada do sistema tem a forma

$$\left(\begin{array}{cccc|c} a'_{11} & \star & \dots & \star & b'_1 \\ 0 & a'_{22} & \dots & \star & b'_2 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \dots & a'_{nn} & b'_n \end{array} \right)$$

A última equação, $a'_{nn}x_n = b'_n$, determina x_n . A penúltima equação, $a'_{n-1,n-1}x_{n-1} + a'_{n-1,n}x_n = b'_{n-1}$, determina x_{n-1} em função do valor já calculado de x_n e assim a seguir, até obter a solução única do sistema. Esta última série de operações é chamada “back-substitution”, e é claro que também consiste em operações elementares. O resultado final é um sistema definido pela matriz identidade, e portanto por uma matriz aumentada do género

$$\left(\begin{array}{cccc|c} 1 & 0 & \dots & 0 & b''_1 \\ 0 & 1 & \dots & 0 & b''_2 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \dots & 1 & b''_n \end{array} \right)$$

A única solução é $(x_1, x_2, \dots, x_n) = (b''_1, b''_2, \dots, b''_n)$.

Cálculo da matriz inversa. Quando A é uma matriz quadrada, o cálculo da solução de $AX = B$, enquanto função do vetor B , é equivalente à inversão da matriz A . O método de Gauss pode ser usado então para calcular a matriz inversa. De fato, as colunas de A^{-1} são as soluções dos sistemas $AX = E_k$, se E_k , com $k = 1, 2, \dots, n$, denotam os vetores colunas da base canónica de $\mathbb{R}^n$. Podemos então aplicar a eliminação de Gauss a todos estes problemas $AX = E_k$, logo à matriz aumentada

$$\left(\begin{array}{cccc|cccc} a_{11} & a_{12} & \dots & a_{1n} & 1 & 0 & \dots & 0 \\ a_{21} & a_{22} & \dots & a_{2n} & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} & 0 & 0 & \dots & 1 \end{array} \right)$$

até chegar ao problema equivalente

$$\left(\begin{array}{cccc|cccc} 1 & 0 & \dots & 0 & b_{11} & b_{12} & \dots & b_{1n} \\ 0 & 1 & \dots & 0 & b_{21} & b_{22} & \dots & b_{2n} \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 & b_{n1} & b_{n2} & \dots & b_{nn} \end{array} \right)$$

Os b_{ij} 's são então os elementos da matriz inversa, ou seja, $A^{-1} = (b_{ij})$.

e.g. Um sistema com solução única. Por exemplo, queremos resolver o sistema

$$\begin{cases} 2x & +y & -2z & = 1 \\ x & -2y & +z & = 2 \\ 3x & +y & -3z & = -5 \end{cases}$$

Se trocamos primeira e segunda linhas, ficamos com o sistema equivalente

$$\begin{cases} x & -2y & +z & = 2 \\ 2x & +y & -2z & = 1 \\ 3x & +y & -3z & = -5 \end{cases}$$

Se retiramos da segunda linha o dobro da primeira, e depois da terceira linha o triplo da primeira, obtemos o sistema equivalente

$$\begin{cases} x & -2y & +z & = 2 \\ & 5y & -4z & = -3 \\ & 7y & -6z & = -11 \end{cases}$$

Se agora retiramos 7 vezes a segunda linha de 5 vezes a terceira linha, ficamos com o sistema equivalente em escada de linhas

$$\begin{cases} x & -2y & +z & = 2 \\ & 5y & -4z & = -3 \\ & & -2z & = -34 \end{cases}$$

A terceira equação diz que $z = 34/2 = 17$, então a segunda equação diz que $y = (-3+4 \cdot 17)/5 = 13$, e a finalmente a primeira equação diz que $x = (2 + 2 \cdot 13 - 17) = 11$.

e.g. Um sistema com uma família de soluções. Por outro lado, consideramos o sistema

$$\begin{cases} x & -2y & +z & = 2 \\ 2x & +y & -2z & = 1 \\ 3x & -y & -z & = 3 \end{cases}$$

Se retiramos da segunda linha o dobro da primeira, e depois da terceira linha o triplo da primeira, obtemos o sistema equivalente

$$\begin{cases} x & -2y & +z & = 2 \\ & 5y & -4z & = -3 \\ & 5y & -4z & = -3 \end{cases}$$

A segunda e a terceira linhas são iguais! Então, retirando da terceira a segunda linha, ficamos com um sistema de apenas duas equações em escada de linhas

$$\begin{cases} x & -2y & +z & = 2 \\ & 5y & -4z & = -3 \end{cases}$$

A variável z pode então ser considerada um parâmetro, assim que o sistema determina x e y em função de z , pois podemos escrever

$$\begin{cases} x & -2y & & = 2 - z \\ & 5y & & = -3 + 4z \end{cases}$$

Finalmente, se chamamos $z = 5t$ (o parâmetro das soluções), então a segunda equação diz que $y = (-3 + 4z)/5 = -\frac{3}{5} + 4t$, e a primeira equação diz que $x = 2 - z + 2y = \frac{4}{5} + 3t$. Temos assim uma reta de soluções, a reta $(4/5, -3/5, 0) + \mathbb{R}(3, 4, 5)$.

ex: Usando operações elementares sobre as linhas, transforme a matriz A dada numa matriz em escada de linhas, e calcule a característica de A .

$$A = \begin{pmatrix} 2 & 2 & 1 \\ 1 & 3 & 1 \\ 1 & 2 & 2 \end{pmatrix} \quad A = \begin{pmatrix} 1 & 2 \\ 0 & 1 \\ 3 & 4 \end{pmatrix}$$

$$A = \begin{pmatrix} 0 & 1 & 3 & -2 \\ 2 & 1 & -4 & 3 \\ 2 & 3 & 2 & -1 \end{pmatrix} \quad A = \begin{pmatrix} 1 & 2 & -1 & 2 \\ 2 & 4 & 1 & -2 \end{pmatrix}$$

ex: Resolva os seguintes sistemas lineares usando o método de eliminação de Gauss.

$$\begin{cases} 3x + 2y + z = 1 \\ 5x + 3y + 3z = 2 \\ -x + y + z = -1 \end{cases} \quad \begin{cases} 3x + 2y + z = 1 \\ 2x - 6y + 4z = 3 \\ x + y + z = -2 \\ 2x - 5y + 5z = -1 \end{cases} \quad \begin{cases} 2x + y + 4z = 2 \\ 6x + y = -10 \\ -x + 2y - 10z = -4 \end{cases}$$

$$\begin{cases} y + z = 1 \\ x + 2y - z = 3 \\ x + y + z = 1 \end{cases} \quad \begin{cases} x + 2y + 3z + 4w = 3 \\ 5z + 6w = 0 \\ z + 3w = 1 \\ x - y + 8w = 0 \end{cases} \quad \begin{cases} -x + y - z = 1 \\ x + 3z = -3 \\ y - z = 0 \end{cases}$$

ex: [Ap69] 16.20.

ex: Dê exemplos de

- um sistema de 2 equações lineares com 2 incógnitas com solução única,
- um sistema de 2 equações lineares com 2 incógnitas sem nenhuma solução,
- um sistema de 3 equações lineares com 3 incógnitas tal que o espaço das soluções seja uma reta afim.
- um sistema de 3 equações lineares com 3 incógnitas com solução única,
- um sistema de 2 equações lineares com 3 incógnitas tal que o espaço das soluções seja um plano afim.
- um sistema de 2 equações lineares com 3 incógnitas tal que o espaço das soluções seja um subespaço vetorial de dimensão 1.

Operações sobre as linhas e matrizes elementares. No caso de sistemas de n equações em n incógnitas, as operações elementares do método de Gauss podem convenientemente ser interpretadas e calculadas como produtos da matriz quadrada A que define o sistema por certas “matrizes elementares”.

Fixada a dimensão n , denotamos por I_{ij} a matriz quadrada $n \times n$ definida em (9.1), cuja única entrada não nula é o elemento da i -ésima linha e a j -ésima coluna, que é igual a 1 (já observamos na seção 9 que estas matrizes formam uma base do espaço linear das matrizes $n \times n$).

Se A é uma matriz quadrada, então o produto $I_{ij}A$ é uma matriz quadrada cuja única linha não nula é a i -ésima, e esta linha é igual a j -ésima linha da matriz A . Consequentemente, a operação EG1, trocar as linhas i e j (naturalmente com $i \neq j$), corresponde a substituir a matriz A pela matriz $A' = EA$, onde E é a matriz elementar

$$E_{ij} = I_{ij} + I_{ji} + \sum_{k \neq i, j} I_{kk} \quad (11.3)$$

obtida da matriz identidade I ao trocar as linhas i e j .

A operação EG2, multiplicar a i -ésima linha por uma constante não nula λ , corresponde a substituir a matriz A pela matriz $A' = EA$, onde E é a matriz elementar

$$M_i(\lambda) = \lambda I_{ii} + \sum_{k \neq i} I_{kk} \quad (11.4)$$

obtida da matriz identidade I ao multiplicar a i -ésima linha (ou seja, a i -ésima entrada diagonal) por λ .

Finalmente, a operação EG3, somar à i -ésima linha α vezes a j -ésima linha (com $i \neq j$), corresponde a substituir a matriz A pela matriz $A' = EA$, onde E é a matriz elementar

$$S_{ij}(\alpha) = I + \alpha E_{ij} \quad (11.5)$$

obtida da matriz identidade somando α vezes E_{ij} .

As matrizes do género (11.3), (11.4) e (11.5) são chamadas *matrizes elementares*. São invertíveis, e a inversa de uma matriz elementar é também uma matriz elementar (pois as operações elementares são reversíveis, e as inversas são operações elementares).

A matriz escada de linha A' obtida ao aplicar a eliminação de Gauss é portanto um produto $A' = E_k \dots E_2 E_1 A$ de um certo número de matrizes elementares E_i 's pela matriz original A . Consequentemente, A é invertível sse A' é invertível. Por outro lado, a matriz quadrada em escada de linha A' , que é uma matriz diagonal superior, é invertível sse os termos diagonais a'_{ii} forem todos diferentes de zero. Neste caso, é também fácil ver que mais umas operações elementares (que fazem a back substitution) podem transformar a matriz A' na matriz identidade I . Finalmente,

Teorema 11.6. *Uma matriz quadrada A é invertível sse existem matrizes elementares $E_1, E_2, \dots, E_k$ tais que*

$$E_k \dots E_2 E_1 A = I$$

ou seja, sse é um produto $A = E'_1 E'_2 \dots E'_k$ de matrizes elementares.

ex: Seja A uma matriz quadrada. Descreva os produtos $I_{ij}A$ e AI_{ij} em termos das linhas e das colunas de A .

ex: Calcule as inversas das matrizes elementares E_{ij} , $M_i(\lambda)$ e $S_{ij}(\alpha)$, definidas em (11.3), (11.4) e (11.5) (com $\lambda \neq 0$), e verifique que também são elementares.

Computational cost of Gaussian elimination. Real world linear systems of interest in engineering involve large matrices, with hundreds or even thousands of entries. Gaussian elimination is clearly the simplest algorithm that can be implemented in a computer. It is a good exercise to write such a code, in your favourite programming language. Below, I show the most simple-minded code in [Python](#) that solves a linear system of n equations in n unknown with a unique solution, written without worrying about possible divisions by zero (but a real code must treat such cases!).

```

1  # scientific libraries
2  import numpy as np
3
4  # data
5  n = 3 # number of variables
6  a = np.array([[1, 2, 3], [5, 6, 7], [13, 17, 66]], float) # matrix of the linear map
7  b = np.array([1, 11, 111], float) # vector of the r.h.s.
8  x = np.zeros(n) # vector of the unknowns / solution
9  ratio = np.zeros(n)
10
11 # Gaussian elimination
12 for i in range(n-1):
13     for j in range(i+1, n):
14         ratio = a[j,i] / a[i,i]
15         b[j] = b[j] - ratio * b[i]
16         for k in range(n):
17             a[j,k] = a[j,k] - ratio * a[i,k]
18
19 # back substitution
20 x[n-1] = b[n-1] / a[n-1,n-1] # last coordinate of the solution
21 for i in range(n-2, -1, -1):
22     somma = b[i] # partial sum
23     for j in range(i+1, n):
24         somma = somma - a[i,j] * x[j]
25     x[i] = somma / a[i,i]
26
27 # output
28 print('\n The upper-diagonal equivalent matrix is \n A = ' + str(a))
29 print('the r.h.s. is B = ' + str(b))
30 print('and the solution is X = ' + str(x) )

```

It is important to estimate the time needed for the algorithm to solve the problem. In a first approximation, this amounts to estimate the number of elementary algebraic operations (sums, multiplications and divisions) that a machine must perform. The loop in lines 16 and 17 requires about $n - i$ multiplications and $n - i$ sums, hence $\sim 2(n - i)$ operations (we may safely disregard the 3 operations of lines 14 and 15). The loop at line 13 performs therefore something like $2(n - i)^2$ operations. Finally, the initial loop at line 12, hence the whole Gaussian elimination algorithm, performs something like

$$\sum_{i=1}^{n-1} 2(n-i)^2 \sim \frac{2}{3} n^3$$

(up to leading order) elementary operations. A similar analysis shows that back substitution, which consists of two loops, only requires something of order $\sim n^2$ operations. Therefore, the entire Gauss-Jordan algorithm that solves a linear system of n equations in n unknowns has an “algebraic cost” (which is different from the real “computational cost”, since we are disregarding the actual complexity of making multiplications or divisions ...) of order $\mathcal{O}(n^3)$. This means that if solving a system of N equations in N unknowns requires τ seconds, then solving a similar system of $10 \cdot N$ equations in $10 \cdot N$ unknowns, hence only a factor 10 larger, requires something like 1000τ seconds (a ratio that makes the difference between one second and three hours!).

There exist more sophisticated algorithms that reduce the cost to something like $\mathcal{O}(n^{2.5})$ or even $\mathcal{O}(n^{2.376\dots})$, and also Monte Carlo (i.e. probabilistic, originally proposed by John von Neumann and Stan Ulam) algorithms that run in much smaller time. The computational cost can be substantially reduced if we know (as is the case in many real situations) that the involved matrix A is a “sparse matrix”, i.e. has a lot of zeroes in some precise sense. This is an important theme in contemporary computational mathematics ...

Equações diferenciais lineares com coeficientes constantes. O núcleo do operador derivação $(Df)(t) = f'(t)$ é o espaço de dimensão um das funções constantes $f(t) = c$. As soluções da equação diferencial linear

$$Df = g,$$

onde $g(t)$ é uma função integrável dada, são $f(t) = (Pg)(t) + c$, onde $(Pg)(t) = \int_0^t g(s) ds$ é uma solução particular, e $c = f(0)$ é uma constante arbitrária, solução geral da equação homogênea $Df = 0$.

Uma equação diferencial linear com coeficientes constantes é uma equação do género

$$a_n D^n f + \dots + a_1 Df + a_0 = g,$$

para a função $f(t)$, onde os a_k são coeficientes (reais ou complexos) e $g(t)$ é uma função dada (uma força quando $n = 2$). Pode ser pensada como $Lf = g$ se L é operador diferencial

$$L := a_n D^n + \cdots + a_1 D + a_0 I.$$

O núcleo de L , o espaço das soluções da *equação homogênea associada*

$$a_n D^n f + \cdots + a_1 Df + a_0 f = 0,$$

é um espaço vetorial de dimensão n . De fato, $h(t) = e^{zt}$ é uma solução de $Lh = 0$ se z é uma raiz do *polinômio característico* $a_n z^n + \cdots + a_1 z + a_0 = 0$. No caso genérico, este polinômio tem n raízes complexas distintas $z_1, z_2, \dots, z_n$ (conjugadas em pares se os coeficientes a_k forem reais, como acontece às equações típicas que descrevem o mundo real). Os exponenciais $h_k(t) = e^{z_k t}$, com $k = 1, 2, \dots, n$, formam então uma base de $\mathbf{H} = \ker(L) \approx \mathbb{C}^n$. Se $f(t)$ é uma “solução particular” (ou seja, apenas uma!) da equação $Lf = g$, então o espaço das “todas as soluções” de $Lf = g$ é o espaço afim $f + \mathbf{H}$. Ou seja, a “solução geral” é uma combinação linear

$$f(t) + c_1 e^{z_1 t} + c_2 e^{z_2 t} + \cdots + c_n e^{z_n t}$$

com $c_1, c_2, \dots, c_n$ coeficientes arbitrários (determinados pelas condições iniciais).

O caso não genérico em que algumas raízes do polinômio característico têm multiplicidade > 1 é tratado na UC de Complementos de Cálculo e de Geometria Analítica, e envolve “quase-polinômios”, ou seja, produtos $p(t)e^{zt}$ de exponenciais vezes polinômios. Determinar uma solução particular é também simples quando a força é um quase-polinômio $g(t) = p(t)e^{zt}$, pois neste caso é sempre possível arranjar um quase-polinômio $f(t) = q(t)e^{zt}$, com $\deg(q) \leq \deg(p) + n$, que resolve $Lf = g$ (método dos coeficientes indeterminados).

12 Volume e determinantes

ref: [Ap69] Vol 2, 3.1-17 ; [La97] Ch. VII

12.1 Formas alternadas e volumes

Motivações geométricas. A maneira mais transparente de definir determinantes é, na minha opinião, a maneira axiomática de Emil Artin, adotada, por exemplo, em [Wa91] e [Ap69].

7 dez 2022

O volume de uma região “razoável” $\Omega \subset \mathbb{R}^n$ é definido/calculado aproximando a região com reuniões de hipercubos de lados r suficientemente pequenos, cujo volume é r^n (as regiões para as quais estas aproximações fazem sentido, sem entrar em detalhes difíceis de análise, são chamadas “mensuráveis”). Um operador linear L envia hipercubos em paralelepípedos. A linearidade de L diz que estes paralelepípedos apenas dependem das dimensões dos hipercubos, e não das posições. Consequentemente, é claro que a razão entre o volume da imagem $L(\Omega)$ e o volume de Ω é um número que não depende da região Ω , mas apenas do operador L .



É suficiente portanto, pela homogeneidade de L , compreender o volume da imagem do hipercubo unitário $Q = [0, 1]^n$. Esta imagem $L(Q)$ é um paralelepípedo cujos lados são as colunas da matriz A que define L na base canônica. Queremos portanto calcular o volume de um paralelepípedo, ou seja, definir uma função que associa um volume a uma n -úpla de vetores, os seus lados. O problema é que se um lado do paralelepípedo é multiplicado por um fator λ , então o volume deve ser multiplicado por $|\lambda|$. É mais conveniente definir um “volume orientado”, ou seja, um volume com um sinal que toma conta da ordem dos lados, pois desta forma as propriedades naturais dos volumes podem ser traduzidas nos simples axiomas algébricos que caracterizam as formas alternadas, de acordo com Emil Artin. O determinante de um operador é, finalmente, um escalar

$$\text{Det}(L) = \pm \frac{\text{Vol}(L(\Omega))}{\text{Vol}(\Omega)}$$

que mede a distorção que o operador produz nos volumes, com um sinal que toma conta da orientação.

Por exemplo, se L é diagonal e positivo, ou seja, é representado na base canônica por uma matriz diagonal com entradas positivas $\lambda_1, \lambda_2, \dots, \lambda_n$, então a imagem do hipercubo unitário tem volume igual ao produto dos lados, e portanto o determinante de L será $\text{Det}(L) = \lambda_1 \lambda_2 \dots \lambda_n$. Uma isometria, um operador que preserva as distâncias euclidianas, preserva também os volumes (pois envia o hipercubo unitário num hipercubo de lado um), logo deve necessariamente ter determinante de módulo um. Por exemplo, uma rotação do plano, que também preserva a orientação, deve ter determinante 1. Por outro lado, uma reflexão do plano, que muda a orientação preservando as distâncias, deve ter determinante -1 .

Uma consequência natural desta interpretação geométrica é que o determinante é multiplicativo: o determinante da composição de dois operadores, ou seja, do produto de duas matrizes, é igual ao produto dos determinantes dos dois operadores. Em particular, deve ser possível calcular determinantes fatorizando um operador genérico como produto de isometrias e operadores diagonais e positivos...

Formas alternadas. Uma n -forma linear no espaço vetorial $\mathbb{R}^n$ é uma função escalar

$$F : \underbrace{\mathbb{R}^n \times \mathbb{R}^n \times \cdots \times \mathbb{R}^n}_{n \text{ vezes}} \rightarrow \mathbb{R}$$

que envia n vetores ordenados $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$ de $\mathbb{R}^n$ num escalar $F(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n)$, que é *multilinear*, ou seja, homogênea e aditiva em cada variável. Homogênea significa que

$$F(\dots, \lambda \mathbf{v}, \dots) = \lambda F(\dots, \mathbf{v}, \dots) \quad (12.1)$$

e aditiva significa que

$$F(\dots, \mathbf{v} + \mathbf{w}, \dots) = F(\dots, \mathbf{v}, \dots) + F(\dots, \mathbf{w}, \dots) \quad (12.2)$$

Em particular, a homogeneidade implica que F assume o valor nulo se uma das variáveis for igual ao vetor nulo, ou seja, $F(\dots, \mathbf{0}, \dots) = 0$.

Uma n -forma F é *alternada* se é nula quando duas variáveis são iguais (ou, pela homogeneidade, proporcionais), ou seja, se

$$F(\dots, \mathbf{v}, \dots, \mathbf{v}, \dots) = 0 \quad (12.3)$$

Uma forma alternada é também *anti-simétrica*, pois uma consequência da (12.3) é que o valor da forma muda de sinal ao trocar duas variáveis, ou seja,

$$F(\dots, \mathbf{v}, \dots, \mathbf{w}, \dots) = -F(\dots, \mathbf{w}, \dots, \mathbf{v}, \dots) \quad (12.4)$$

(a prova consiste em calcular $F(\dots, \mathbf{v} + \mathbf{w}, \dots, \mathbf{v} + \mathbf{w}, \dots)$ e usar a linearidade ...) Vice-versa, é claro que (12.4) com $\mathbf{v} = \mathbf{w}$ implica (12.3) (porque no corpo dos reais, assim como no corpo dos complexos, podemos dividir por 2), assim que as duas propriedades são equivalentes.

Uma consequência da aditividade (12.2) e da (12.3) é também que o valor de uma forma alternada não muda se somamos uma variável a outra, ou seja,

$$F(\dots, \mathbf{v}, \dots, \mathbf{w}, \dots) = F(\dots, \mathbf{v} + \mathbf{w}, \dots, \mathbf{w}, \dots) \quad (12.5)$$

Finalmente, usando repetidamente a homogeneidade e a (12.5), o valor de uma forma alternada é nulo se uma das variáveis é uma combinação linear das outras, ou seja, se as n variáveis são linearmente dependentes.

Ainda não sabemos se formas alternadas não triviais, ou seja, não identicamente nulas, existem.

Fixada uma base de $\mathbb{R}^n$, por exemplo a base canônica $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$, uma n -forma F é determinada pelas suas “coordenadas”

$$f_{ijk\dots} := F(\mathbf{e}_i, \mathbf{e}_j, \mathbf{e}_k, \dots)$$

pois, pela linearidade em cada variável,

$$F(\mathbf{x}, \mathbf{y}, \mathbf{z}, \dots) = \sum_{ijk\dots} x_i y_j z_k \dots F(\mathbf{e}_i, \mathbf{e}_j, \mathbf{e}_k, \dots)$$

Se a n -forma F é alternada, então as coordenadas com (pelo menos) dois índices repetidos são nulas, ou seja, $f_{\dots i \dots i \dots} = 0$. Assim, apenas podem ser diferentes de zero as coordenadas cujos índices são permutações (funções bijetivas) do conjunto $\{1, 2, \dots, n\}$. Pela (12.4), as coordenadas são também anti-simétricas, ou seja, $f_{\dots i \dots j \dots} = -f_{\dots j \dots i \dots}$.

É um fato que toda permutação $\sigma : \{1, 2, \dots, n\} \rightarrow \{1, 2, \dots, n\}$ pode ser obtida da permutação identidade, definida por $\sigma_0(k) = k$, usando transposições (trocas entre as posições de apenas dois elementos) repetidas. O número N de transposições que envia a permutação identidade σ_0 em σ não é, naturalmente, único, mas é bem definida a sua “paridade” $\text{sgn}(\sigma) := (-1)^N$, que assume os valores ± 1 dependendo se N é par ou ímpar (isto é intuitivamente óbvio, podem ler uma ideia da prova no próximo parágrafo).

Consequentemente, se a coordenada de uma forma alternada F sobre os vetores ordenados da base canônica é $f_{12\dots n} = \lambda$, então as outras coordenadas não nulas são $\pm \lambda$, dependendo da paridade da permutação, ou seja, são iguais a $f_{\sigma(1)\sigma(2)\dots\sigma(n)} = \lambda \cdot \text{sgn}(\sigma)$.

Em outras palavras, o espaço linear $\text{Alt}^n(\mathbb{R}^n)$ das n -formas alternadas em $\mathbb{R}^n$ tem dimensão igual a 1. Temos portanto o seguinte teorema de unicidade.

Teorema 12.1. *Existe uma única n -forma alternada D em $\mathbb{R}^n$ normalizada de maneira tal que $D(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n) = 1$.*

Uma forma alternativa e útil de formular este teorema de unicidade é a seguinte. Toda n -forma alternada F em $\mathbb{R}^n$ é proporcional a D , ou seja, é igual a $F = \lambda D$ se $\lambda = F(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n)$ é o seu valor nos vetores da base canônica, ou seja,

$$F(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n) = F(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n) D(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n) \quad (12.6)$$

A “ n -forma alternada canônica” D é também chamada “determinante”. As suas coordenadas na base canônica são os números

$$\varepsilon_{ijk\dots} := D(\mathbf{e}_i, \mathbf{e}_j, \mathbf{e}_k, \dots),$$

definidos por $\varepsilon_{ijk\dots} = 0$ se dois índices são iguais, e $\varepsilon_{ijk\dots} = \text{sgn}(\sigma)$ se $ijk\dots = \sigma(1)\sigma(2)\sigma(3)\dots$ onde σ é uma permutação de $\{1, 2, 3, \dots, n\}$. A coleção dos números $\varepsilon_{ijk\dots}$ é chamada *símbolo de Levi-Civita*. O valor da forma canônica D sobre os vetores $\mathbf{x}, \mathbf{y}, \mathbf{z}, \dots$ é

$$D(\mathbf{x}, \mathbf{y}, \mathbf{z}, \dots) = \sum_{ijk\dots} \varepsilon_{ijk\dots} x_i y_j z_k \dots$$

ex: Mostre que, se F é uma forma bi-linear, então $F(\mathbf{x}, \mathbf{x}) = 0$ implica $F(\mathbf{x}, \mathbf{y}) + F(\mathbf{y}, \mathbf{x}) = 0$ (calcule $F(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) \dots$)

ex: Verifique que a única 2-forma alternada no plano $\mathbb{R}^2$ normalizada de maneira tal que $D(\mathbf{e}_1, \mathbf{e}_2) = 1$ é

$$D(a\mathbf{e}_1 + b\mathbf{e}_2, c\mathbf{e}_1 + d\mathbf{e}_2) = ad - bc.$$

ex: Mostre que o valor de uma n -forma alternada é nulo se uma das variáveis é uma combinação linear das outras, ou seja, se as n variáveis são linearmente dependentes.

Permutations, transpositions and parity. Existence of a non-trivial alternating forms depends, in our discussion, on existence of the parity of a permutation of a finite set. Here we sketch a possible proof. The first observation is that any permutation of the set $\{1, 2, \dots, N\}$ is a composition of transpositions, permutations that only exchange two positions, say i and $j \neq i$ (you may try to prove it by induction). We also note that transpositions are involutions, i.e. they are their own inverses. The second observation is that we may restrict to “adjacent transpositions”, those between two successive positions, say i and $i + 1$. Indeed, any transposition may be obtained as a composition of a minimal number of adjacent transpositions, which is an odd number (if you want to exchange i and j , you may perform $|j - i|$ adjacent transpositions to bring i to the position of j , and then back $|j - i| - 1$ adjacent transpositions to bring j at the position of i , for a total of $2|j - i| - 1$ adjacent transpositions). Therefore, a generic permutation σ may be written as a composition $\sigma = \tau_n \dots \tau_2 \tau_1$ of a certain number of adjacent transpositions τ_k . This representation is clearly not unique. So, suppose that the same permutation may also be written as a possibly different product $\sigma = \tau'_m \dots \tau'_2 \tau'_1$ of adjacent transpositions. We want to prove that n and m share the same parity, i.e. that they are both even or odd. We then observe that the composition $\sigma\sigma^{-1} = \tau_n \dots \tau_2 \tau_1 (\tau'_m \dots \tau'_2 \tau'_1)^{-1} = \tau_n \dots \tau_2 \tau_1 \tau'_1 \tau'_2 \dots \tau'_m$ represents the identity permutation as a composition of $n + m$ adjacent transpositions. So, we want to prove that $n + m$ is even. This amounts to prove the following.

Teorema 12.2. *If the identity permutation is represented as a composition $\tau_k \dots \tau_2 \tau_1$ of k adjacent transpositions, the k is even.*

To help seeing what is going on, we may associate to such a composition a diagram, with the planar trajectories of N numbered “particles” (here a picture would help). The i -th particle starts at the position i . Every unit of time, the two particles which occupy the positions involved in the transposition τ_1 , then τ_2 , $\dots$, and finally τ_k , flip. Since the final permutation is the trivial one, the position of the i -th particle at the final time k is again i . Now, for any chosen pair of particles, say the i -th and the j -th with $i < j$, the number k_{ij} of flips between the two must be even (possibly zero), since the two particles start and end at the same positions, in particular with the same order. But the total number of flips, i.e. the total number of adjacent transpositions, is the sum $k = \sum_{i < j} k_{ij}$. This number is even, being the sum of even integers.

Formas alternadas e volumes de paralelepípedos. O *paralelepípedo* de lados $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$ no espaço $\mathbb{R}^n$ é o conjunto

$$P = \{t_1\mathbf{v}_1 + t_2\mathbf{v}_2 + \dots + t_n\mathbf{v}_n \text{ com } t_1, t_2, \dots, t_n \in [0, 1]\}.$$

A n -forma alternada canônica D calcula o *volume orientado* dos paralelepípedos, um volume com sinal positivo ou negativo dependendo da maneira em que os lados são ordenados. De fato, uma pequena reflexão mostra que um volume orientado satisfaz as propriedades que definem a forma canônica, ou seja, é linear em cada variável/lado, é nulo quando dois lados são iguais ou proporcionais, e é normalizado de maneira tal que o hipercubo cujos lados são os vetores ordenados da base canônica tem volume igual a um. Em dimensão dois, isto pode ser observado muito facilmente. É interessante observar que assim como o produto escalar, uma forma bilinear simétrica, contém informação sobre projeções, ou seja, sobre o cosseno do ângulo entre dois vetores, o determinante, uma forma bilinear anti-simétrica, contém informações sobre área de paralelogramos, ou seja, sobre o seno do ângulo entre dois vetores.

Finalmente, o volume (n -dimensional) do paralelepípedo P é igual ao valor absoluto do valor da n -forma canônica sobre os lados $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$, ou seja,

$$\text{Vol}(P) = |D(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n)|$$

12.2 Determinante

Determinante de uma matriz. Uma matriz quadrada $A = (a_{ij}) \in \text{Mat}_{n \times n}(\mathbb{R})$ pode ser pensada como uma lista de n vetores de $\mathbb{R}^n$, as suas n colunas

$$A_{*1} = \begin{pmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{n1} \end{pmatrix} \quad A_{*2} = \begin{pmatrix} a_{12} \\ a_{22} \\ \vdots \\ a_{n2} \end{pmatrix} \quad \dots \quad A_{*n} = \begin{pmatrix} a_{1n} \\ a_{2n} \\ \vdots \\ a_{nn} \end{pmatrix}$$

O teorema de unicidade 12.1 implica então que existe uma única forma multilinear alternada nas colunas da matriz e normalizada de maneira tal que o seu valor sobre a matriz identidade I seja (pois a matriz identidade é a matriz cujas colunas são os vetores da base canônica, ordenados da maneira natural). Esta função, definida por

$$\text{Det} A := D(A_{*1}, A_{*2}, \dots, A_{*n}) \quad (12.7)$$

é denotada por $\text{Det} : \text{Mat}_{n \times n} \rightarrow \mathbb{R}$ e chamada *determinante (de ordem n)*. Outra notação também utilizada é $\text{Det} A = |A|$. Usando o símbolo de Levi-Civita, o determinante pode ser definido pela fórmula

$$\text{Det} A = \sum_{ijk\dots} \varepsilon_{ijk\dots} a_{i1} a_{j2} a_{k3} \dots \quad (12.8)$$

Mais explícita é a fórmula

$$\text{Det} A = \sum_{\sigma} \text{sgn}(\sigma) a_{\sigma(1)1} a_{\sigma(2)2} \dots a_{\sigma(n)n} \quad (12.9)$$

onde a soma é sobre toda as permutações $\sigma : \{1, 2, \dots, n\} \rightarrow \{1, 2, \dots, n\}$.

O determinante de ordem n , sendo linear em cada coluna (ou linha), é homogêneo de grau n . Ou seja, se $A \in \text{Mat}_{n \times n}(\mathbb{R})$ e $\lambda \in \mathbb{R}$, então

$$\text{Det}(\lambda A) = \lambda^n \text{Det} A$$

o que é natural sendo também um volume orientado.

Teorema 12.3. *O determinante de uma matriz quadrada A é igual ao determinante da matriz transposta $A^\top$, ou seja,*

$$\boxed{\text{Det} A^\top = \text{Det} A} \quad (12.10)$$

Demonstração. Se usamos a fórmula (12.9) para o determinante de $A^\top$ e reordenamos os termos da soma de maneira tal que o segundo índice dos elementos da matriz seja crescente, obtemos

$$\begin{aligned} \text{Det} A^\top &= \sum_{\sigma} \text{sgn}(\sigma) a_{1\sigma(1)} a_{2\sigma(2)} \cdots a_{n\sigma(n)} \\ &= \sum_{\sigma} \text{sgn}(\sigma) a_{\sigma^{-1}(1)1} a_{\sigma^{-1}(2)2} \cdots a_{\sigma^{-1}(n)n} \end{aligned}$$

onde σ^{-1} é a permutação inversa de σ . Observamos finalmente que as paridades de uma permutação e da sua inversa são iguais, sendo a composição a permutação identidade, e que somar sobre todas as permutações σ é equivalente a somar sobre todas as permutações σ^{-1} , assim que

$$\text{Det} A^\top = \sum_{\sigma^{-1}} \text{sgn}(\sigma^{-1}) a_{\sigma^{-1}(1)1} a_{\sigma^{-1}(2)2} \cdots a_{\sigma^{-1}(n)n} = \text{Det} A$$

□

Assim, o determinante de uma matriz pode também ser definido como o valor da forma alternada canônica D sobre as suas linhas, ou seja,

$$\boxed{\text{Det} A = D(A_{1*}, A_{2*}, \dots, A_{n*})} \quad (12.11)$$

Se as colunas (ou as linhas) de uma matriz quadrada A são vetores linearmente dependentes, então $\text{Det} A = 0$. Por exemplo, o determinante de uma matriz quadrada com uma coluna ou uma linha nula é igual a zero.

Cálculo de determinantes. O determinante da matriz identidade, cujas colunas são os vetores $E_1, E_2, \dots, E_n$ da base canônica, é igual a

$$\text{Det} I = 1$$

É possível calcular determinantes de algumas matrices usando as propriedades (12.1), (12.2), (12.3) e a normalização $D(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n) = 1$. As fórmulas ficam logo compridas, pois o número das permutações de n elementos é $n!$, e o fatorial cresce muito rapidamente.

Por exemplo, é um exercício verificar que o determinante de uma matriz 2×2 é

$$\boxed{\text{Det} \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = a_{11} a_{22} - a_{12} a_{21}}$$

(ou seja, que esta é a única 2-forma alternada nas colunas de A que assume valor unitário se as colunas são os vetores ordenados da base canônica). Também é relativamente simples verificar que o determinante de uma matriz 3×3 é

$$\boxed{\text{Det} \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} = a_{11} \text{Det} \begin{pmatrix} a_{22} & a_{23} \\ a_{32} & a_{33} \end{pmatrix} - a_{12} \text{Det} \begin{pmatrix} a_{21} & a_{23} \\ a_{31} & a_{33} \end{pmatrix} + a_{13} \text{Det} \begin{pmatrix} a_{21} & a_{22} \\ a_{31} & a_{32} \end{pmatrix}}$$

Mais importante é que esta fórmula sugere a possibilidade de dar uma definição recursiva do determinante ...

O cálculo do determinante de uma matriz diagonal é uma consequência da homogeneidade e da normalização, e é igual ao produto dos termos diagonais, ou seja

$$\text{Det} \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix} = \lambda_1 \lambda_2 \dots \lambda_n$$

(quando os λ_k são positivos, este é o volume de um paralelepípedo de lados dois a dois ortogonais de comprimentos λ_k 's).

Usando repetidamente a (12.5), este é também o caso do determinante de uma matriz diagonal superior (ou inferior), ou seja,

$$\text{Det} \begin{pmatrix} \lambda_1 & * & \dots & * \\ 0 & \lambda_2 & \dots & * \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix} = \lambda_1 \lambda_2 \dots \lambda_n$$

De fato, se uma das entradas diagonais for zero, então uma coluna da matriz é dependente de outras, e portanto o determinante é nulo. Se, por outro lado, todas as entradas diagonais forem diferentes de zero, então é possível subtrair de cada coluna umas oportunas combinações lineares de outras colunas até obter uma matriz diagonal com o mesmo determinante.

ex: Calcule o determinante das seguintes matrizes quadradas

$$\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \quad \begin{pmatrix} 2 & 3 \\ 0 & 5 \end{pmatrix} \quad \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$$

$$\begin{pmatrix} 3 & 0 & 0 \\ 0 & 5 & 0 \\ 0 & 0 & 7 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 1 & 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 0 & 2 & 3 \\ 0 & 0 & 7 \\ 4 & 5 & 6 \end{pmatrix} \quad \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix}$$

ex: Mostre que

$$\text{Det} \begin{pmatrix} 1 & 1 & 1 \\ a & b & c \\ a^2 & b^2 & c^2 \end{pmatrix} = (b-a)(c-a)(c-b)$$

ex: Calcule o determinante de $2A$ e $-A$ sabendo que A é uma matriz 5×5 com determinante $\text{Det}A = -3$.

ex: Verifique que uma equação cartesiana da reta que passa pelos pontos (a, b) e (c, d) de $\mathbb{R}^2$ é

$$\text{Det} \begin{pmatrix} x-a & y-b \\ c-a & d-b \end{pmatrix} = 0 \quad \text{ou} \quad \text{Det} \begin{pmatrix} x & y & 1 \\ a & b & 1 \\ c & d & 1 \end{pmatrix} = 0$$

ex: [Ap69] 3.6.

Determinante e produtos. A propriedade mais importante do determinante, consequência natural da sua interpretação geométrica, é a multiplicatividade.

Teorema 12.4. *O determinante é uma função multiplicativa. Ou seja, se A e B são duas matrizes $n \times n$, então*

$$\boxed{\text{Det}(AB) = (\text{Det}A)(\text{Det}B)} \quad (12.12)$$

Demonstração. Dada uma matriz quadrada $A = (a_{ij}) \in \text{Mat}_{n \times n}(\mathbb{R})$, a função

$$X_1, X_2, \dots, X_n \mapsto D(AX_1, AX_2, \dots, AX_n)$$

é uma n -forma alternada dos n vetores coluna $X_i \in \mathbb{R}^n$. Pelo teorema de unicidade 12.1, é proporcional à forma canônica D , ou seja, pela fórmula (12.6),

$$D(AX_1, AX_2, \dots, AX_n) = D(AE_1, AE_2, \dots, AE_n) \cdot D(X_1, X_2, \dots, X_n)$$

onde E_i são os vetores coluna da base canônica. Por outro lado, os AE_i são as colunas da matriz A . Portanto, pela definição de determinante (12.7), $D(AE_1, AE_2, \dots, AE_n) = \text{Det} A$. Se $B \in \text{Mat}_{n \times n}(\mathbb{R})$ denota a matriz cujas colunas são os vetores X_i , então os vetores AX_i são as colunas da matriz AB . Então, sempre pela definição (12.7), $D(X_1, X_2, \dots, X_n) = \text{Det} B$ e $D(AX_1, AX_2, \dots, AX_n) = \text{Det}(AB)$. $\square$

Observe que $\text{Det}(AB) = \text{Det}(BA)$, embora AB pode ser diferente de BA . Por indução, o determinante de um produto (finito) é igual ao produto dos determinantes

$$\text{Det}(ABC \dots) = \text{Det}(A) \text{Det}(B) \text{Det}(C) \dots$$

e portanto independente da ordem dos fatores.

Em particular, se A é invertível então $\text{Det}(A)\text{Det}(A^{-1}) = \text{Det} I = 1$, e portanto $\text{Det}(A) \neq 0$ e

$$\boxed{\text{Det}(A^{-1}) = \frac{1}{\text{Det} A}} \quad (12.13)$$

Outra consequência útil da multiplicatividade do determinante é a seguinte. Se $A \in \text{Mat}_{n \times n}(\mathbb{R})$ e $B \in \text{Mat}_{m \times m}(\mathbb{R})$, então é possível construir a matriz “diagonal por blocos”

$$\begin{pmatrix} A & 0 \\ 0 & B \end{pmatrix} \in \text{Mat}_{(n+m) \times (n+m)}(\mathbb{R})$$

que define o operador $(X, Y) \mapsto (AX, BY)$ de $\mathbb{R}^n \times \mathbb{R}^m \simeq \mathbb{R}^{n+m}$. O seu determinante é

$$\boxed{\text{Det} \begin{pmatrix} A & 0 \\ 0 & B \end{pmatrix} = (\text{Det} A)(\text{Det} B)} \quad (12.14)$$

De fato, a matriz diagonal em blocos é um produto

$$\begin{pmatrix} A & 0 \\ 0 & B \end{pmatrix} = \begin{pmatrix} A & 0 \\ 0 & I \end{pmatrix} \begin{pmatrix} I & 0 \\ 0 & B \end{pmatrix}$$

e os determinantes dos fatores são claramente $\text{Det}(A)$ e $\text{Det}(B)$, respectivamente (sempre pelo teorema de unicidade, porque são formas alternadas nas colunas de A e B normalizadas de maneira correta).

ex: Calcule o determinante de A^5 e de AB^2 sabendo que A e B são matrizes 3×3 com determinante $\text{Det} A = -2$ e $\text{det} B = 3$.

ex: Verdadeiro ou falso? Dê uma demonstração ou um contra-exemplo.

$$\text{Det}(A + B) = \text{Det} A + \text{Det} B \quad \text{Det}((A + B)^2) = (\text{Det}(A + B))^2$$

$$\text{Det}(ABA^{-1}) = \text{Det}(B) \quad \text{Det}(AB - BA) = 0$$

ex: Observe que

$$\text{Det}(A^n) = (\text{Det}A)^n.$$

Deduz a que o determinante de uma matriz nilpotente é nulo.

ex: Uma matriz quadrada A é dita “ortogonal” se $A^\top A = A A^\top = I$, ou seja, se é invertível e a sua inversa é $A^\top$ (se reais, são as isometrias lineares do espaço euclidiano $\mathbb{R}^n$). Mostre que o determinante de uma matriz ortogonal é ± 1 .

ex: [Ap69] 3.11.

Cálculo de determinantes pelo método de eliminação de Gauss. O determinante de uma matriz triangular superior (ou triangular inferior) é igual ao produto dos termos diagonais. Consequentemente, é possível calcular um determinante transformando uma matriz quadrada genérica A numa matriz triangular superior A' pelo método de Gauss, ou seja, por meio de um número finito de operações elementares. Isto acontece porque as operações elementares têm efeitos simples no determinante. A operação EG1, permutar duas linhas, muda apenas o sinal do determinante (pela (12.4)). A operação EG2, multiplicar uma linha por um escalar $\lambda \neq 0$, transforma $\text{Det}A$ em $\lambda \text{Det}A$ (pela (12.1)). A operação EG3, somar a uma linha um múltiplo de uma outra linha, não muda o determinante (pelas (12.1) e (12.5)).

Em outras palavras, o algoritmo de eliminação de Gauss permite transformar uma matriz quadrada A numa matriz equivalente $A' = E_k \dots E_2 E_1 A$ que é diagonal superior, sendo as E_i 's matrizes elementares que correspondem a certas operações elementares sobre as linhas de A . Os determinantes das matrizes elementares são simples de calcular, usando as propriedades da n -forma alternada canônica. Finalmente, a multiplicatividade do determinante permite calcular o determinante de A como quociente entre o determinante de A' , que é o produto $a'_{11} a'_{22} \dots a'_{nn}$ dos termos diagonais, e os produtos dos determinantes das E_i 's.

Em particular, apesar da fórmula (12.8) para o determinante de uma matriz $n \times n$ envolver $n!$ produtos e somas (pela fórmula de Stirling, o fatorial cresce como $\sim n^{n+1/2}$), na prática um algoritmo como a eliminação de Gauss reduz este cálculo a apenas $\sim n^2$ operações.

e.g. Por exemplo, queremos calcular o determinante da matriz

$$A = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \\ 3 & 1 & 2 \end{pmatrix}$$

Se retiramos da segunda linha de A o dobro da primeira linha, e da terceira linha de A o triplo da primeira linha, ficamos com a matriz equivalente

$$A' = \begin{pmatrix} 1 & 2 & 3 \\ 0 & -1 & -5 \\ 0 & -5 & -7 \end{pmatrix}$$

que tem o mesmo determinante de A . Se retiramos da terceira linha de A' o quádruplo da segunda linha, ficamos com a matriz equivalente

$$A'' = \begin{pmatrix} 1 & 2 & 3 \\ 0 & -1 & -5 \\ 0 & 0 & 18 \end{pmatrix}$$

que tem o mesmo determinante de A' e portanto de A . Finalmente, $\text{Det}A = -18$.

ex: Verifique, usando a definição de determinante (12.7) e as propriedades da n -forma alternada D , que as matrizes elementares definidas em (11.3), (11.4) e (11.5) têm determinantes

$$\text{Det}(E_{ij}) = -1 \quad \text{Det}(M_i(\lambda)) = \lambda \quad \text{Det}(S_{ij}(\alpha)) = 1$$

ex: Use o método de eliminação de Gauss para calcular o determinante das matrizes

$$\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{pmatrix} \quad \begin{pmatrix} 1 & -1 & 1 & 1 \\ 1 & -1 & -1 & -1 \\ 1 & 1 & -1 & -1 \\ 1 & 1 & 1 & -1 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 & 1 & 1 \\ a & b & c & d \\ a^2 & b^2 & c^2 & d^2 \\ a^3 & b^3 & c^3 & d^3 \end{pmatrix}$$

Determinante e independência linear. As operações elementares sobre as linhas de uma matriz não alteram a sua característica. Em particular, se a característica de uma matriz quadrada A de ordem n é igual a $\text{Rank } A = n$, ou seja, se as colunas de A são vetores linearmente independentes, então a matriz é equivalente a uma matriz triangular superior com pivots λ_k não nulos. O seu determinante é portanto diferente de zero. Vice-versa, é evidente que o determinante de uma matriz cujas colunas são linearmente dependentes é nulo (porque uma coluna é combinação linear das outras). Consequentemente,

Teorema 12.5. *As colunas ou as linhas de uma matriz quadrada A são linearmente independentes sse $\text{Det } A \neq 0$.*

12.3 Fórmula de Laplace

14 dez 2022

A fórmula explícita para determinantes de ordem 3 sugere a possibilidade de determinar fórmulas recursivas que permitam calcular determinantes de ordem n à custa de determinantes de ordem $n - 1$. Esta possibilidade depende apenas das propriedades das formas alternadas e da unicidade da forma alternada canônica.

Complementos algébricos e fórmula de Laplace. Seja $A = (a_{ij}) \in \text{Mat}_{n \times n}(\mathbb{R})$ uma matriz quadrada de ordem $n \geq 2$. Fixados os índices i e j , entre 1 e n , o *menor* ij de A é a matriz $A_{ij} \in \text{Mat}_{(n-1) \times (n-1)}(\mathbb{R})$ obtida da matriz A suprimindo a linha i e a coluna j . O *complemento algébrico* do elemento a_{ij} de A é o número

$$\text{Cal}(a_{ij}) := (-1)^{i+j} \text{Det } A_{ij}.$$

Também útil é definir a *matriz dos complementos algébricos* (ou *dos co-fatores*) de A como sendo a matriz $\text{Cal } A \in \text{Mat}_{n \times n}(\mathbb{R})$ cujo elemento ij é $\text{Cal}(a_{ij})$, ou seja

$$\text{Cal } A := (\text{Cal}(a_{ij}))$$

Finalmente, as *fórmulas de Laplace* mostram como desenvolver o determinante de uma matriz à partir dos elementos da sua i -ésima linha ou da sua j -ésima coluna.

Teorema 12.6 (fórmulas de Laplace). *O determinante de uma matriz quadrada $n \times n$ A é*

$$\boxed{\text{Det } A = \sum_{j=1}^n a_{ij} \text{Cal}(a_{ij}) \quad \text{ou também} \quad \text{Det } A = \sum_{i=1}^n a_{ij} \text{Cal}(a_{ij})} \quad (12.15)$$

onde i ou j são índices (linha ou coluna) arbitrários entre 1 e n .

Demonstração. A ideia da prova é bastante simples, embora custe escrever todos os detalhes, e consiste em usar repetidamente as propriedades básicas das formas alternadas. Consideramos inicialmente a primeira linha de A , que é

$$A_{1*} = a_{11}\mathbf{e}_1 + a_{12}\mathbf{e}_2 + \cdots + a_{1n}\mathbf{e}_n$$

Pela definição de determinante (12.11), e pela linearidade da n -forma alternada canônica D , o determinante de A é uma soma

$$\begin{aligned}\text{Det} A &= D(a_{11}\mathbf{e}_1 + a_{12}\mathbf{e}_2 + \cdots + a_{1n}\mathbf{e}_n, A_{2*}, \dots, A_{n*}) \\ &= \sum_{j=1}^n a_{1j} D(\mathbf{e}_j, A_{2*}, \dots, A_{n*}) \\ &= \sum_{j=1}^n a_{1j} \text{Det} A'_{1j}\end{aligned}$$

onde

$$A'_{1j} = \begin{pmatrix} 0 & \cdots & 1 & \cdots & 0 \\ a_{21} & \cdots & a_{2j} & \cdots & a_{2n} \\ \vdots & & \vdots & & \vdots \\ a_{n1} & \cdots & a_{nj} & \cdots & a_{nn} \end{pmatrix}$$

é a matriz $n \times n$ obtida da matriz A ao substituir a primeira linha pelas coordenadas do vetor $\mathbf{e}_j$ (ou seja, 1 na posição j e zero nas outras). Pela (12.10) e a (12.5), o seu determinante é igual ao determinante da matriz

$$\begin{pmatrix} 0 & \cdots & 1 & \cdots & 0 \\ a_{21} & \cdots & 0 & \cdots & a_{2n} \\ \vdots & & \vdots & & \vdots \\ a_{n1} & \cdots & 0 & \cdots & a_{nn} \end{pmatrix}$$

obtida de A'_{1j} ao substituir por 0 todas as outras entradas da coluna j , exceto a primeira (ou seja somando oportunos múltiplos da primeira linha às outras linhas). Finalmente, pela (12.4), podemos passar a coluna j (formada por um 1 inicial e todos 0's) na primeira posição, multiplicando o determinante por $(-1)^{1+j}$. O determinante da matriz resultante

$$\begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & a_{21} & \cdots & a_{2n} \\ \vdots & & & \vdots \\ 0 & a_{n1} & \cdots & a_{nn} \end{pmatrix}$$

é claramente uma $(n-1)$ -forma alternada nas linhas da matriz $A_{1j} \in \text{Mat}_{(n-1) \times (n-1)}(\mathbb{R})$, obtida da matriz A suprimindo a primeira linha e a coluna j , e assume o valor 1 quando as linhas desta matriz são os vetores da base canônica de $\mathbb{R}^{n-1}$. Pelo teorema de unicidade, este determinante é $\text{Det} A_{1j}$. Consequentemente,

$$\text{Det} A = \sum_{j=1}^n a_{1j} (-1)^{1+j} \text{Det} A_{1j}$$

As fórmulas de Laplace são uma consequência da (12.4), pois podemos trocar a linha i pela primeira linha à custa de um fator $(-1)^{1+i}$ no seu determinante, e depois da (12.10). $\square$

A fórmula de Laplace pode ser (e de fato é) usada, em alternativa, como definição recursiva dos determinantes. As propriedades das formas alternadas são, então, consequências desta definição.

Regra de Sarrus. O determinante de uma matriz $n \times n$ é uma soma de $n!$ produtos de n das entradas a_{ij} , com certos sinais $+$ ou $-$. Acontece que quando $n = 3$, e apenas quando $n = 3$,

$$n! = 2n$$

Esta coincidência dá origem a uma mnemônica para calcular o determinante de uma matriz 3×3 , chamada “regra de Sarrus”. Consiste em construir uma matriz 3×5 (ou 5×3) repetindo

duas vezes a matriz A (exceto a última coluna, ou linha), e observar que os produtos com sinal $+$ correspondem às diagonais descendentes

$$\begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{11} & a_{12} \\ a_{21} & a_{22} & a_{23} & a_{21} & a_{22} \\ a_{31} & a_{32} & a_{33} & a_{31} & a_{32} \end{pmatrix} \Rightarrow a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32}$$

e os produtos com sinal $-$ correspondem às diagonais ascendentes

$$\begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{11} & a_{12} \\ a_{21} & a_{22} & a_{23} & a_{21} & a_{22} \\ a_{31} & a_{32} & a_{33} & a_{31} & a_{32} \end{pmatrix} \Rightarrow -a_{13}a_{22}a_{31} - a_{11}a_{23}a_{32} - a_{11}a_{21}a_{33}$$

ex: Calcule a matriz dos complementos algébricos das matrizes

$$\begin{pmatrix} 1 & 2 & 3 \\ 0 & 3 & 0 \\ -7 & 0 & 0 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 1 & 0 & 1 \end{pmatrix}$$

ex: Calcule o determinante das matrizes

$$\begin{pmatrix} 1 & 0 & -1 & 3 \\ -5 & 0 & 2 & -1 \\ 4 & 0 & 1 & 1 \\ 3 & -2 & -1 & 0 \end{pmatrix} \quad \begin{pmatrix} 2 & 1 & 0 & 0 \\ 1 & 3 & 0 & 0 \\ 0 & 0 & -1 & 1 \\ 0 & 0 & 4 & -2 \end{pmatrix}$$

$$\begin{pmatrix} 2 & 0 & -1 & 0 \\ 1 & 3 & 3 & 0 \\ 1 & 2 & 2 & 0 \\ 3 & 6 & 6 & 0 \end{pmatrix} \quad \begin{pmatrix} 2 & 0 & 1 & 2 \\ 1 & 3 & 1 & -2 \\ 0 & 0 & 2 & 1 \\ 0 & 0 & 0 & -2 \end{pmatrix}$$

12.4 Determinante de um operador

Determinante de um operador. Seja A a matriz quadrada que define a transformação linear $L : \mathbb{R}^n \rightarrow \mathbb{R}^n$, relativamente à base canónica (ou seja, $y_i = \sum_j a_{ij} x_j$). É claro que $(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n) \mapsto D(L\mathbf{v}_1, L\mathbf{v}_2, \dots, L\mathbf{v}_n)$ é uma n -forma alternada. Pelo teorema de unicidade, ou seja, pela (12.6),

$$\begin{aligned} D(L\mathbf{v}_1, L\mathbf{v}_2, \dots, L\mathbf{v}_n) &= D(L\mathbf{e}_1, L\mathbf{e}_2, \dots, L\mathbf{e}_n) D(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n) \\ &= (\text{Det}A) \cdot D(\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n) \end{aligned}$$

pois os vetores $L\mathbf{e}_k$ são as colunas da matriz A . Isto mostra que o determinante da matriz A depende apenas da transformação linear L , e não da base usada para calcular a matriz!

Uma maneira mais pedante de ver isto é a seguinte. Uma mudança de coordenadas envia a matriz A na matriz $A' = U^{-1}AU$, e a multiplicatividade implica que

$$\text{Det}A' = \text{Det}(U^{-1}AU) = \text{Det}(U^{-1}) \text{Det}A \text{Det}U = \text{Det}A$$

Faz portanto sentido definir o “determinante” de um operador $L : \mathbb{R}^n \rightarrow \mathbb{R}^n$ como

$$\boxed{\text{Det}L := \text{Det}A}$$

onde A é a matriz que representa L numa base arbitrária de $\mathbb{R}^n$. Naturalmente, o operador identidade tem determinante $\text{Det}I = 1$. É claro, pelo teorema 12.4, que o determinante de uma composição é o produto dos determinantes, ou seja,

$$\text{Det}(LM) = (\text{Det}L)(\text{Det}M)$$

Em particular, se $L : \mathbb{R}^n \rightarrow \mathbb{R}^n$ é invertível, então o seu determinante não é nulo, e o determinante da transformação inversa é

$$\text{Det}(L^{-1}) = \frac{1}{\text{Det}L}$$

Orientação. O determinante de um automorfismo de $\mathbb{R}^n$, sendo diferente de zero, dá origem uma dicotomia e permite finalmente dar uma definição quantitativa/rigorosa de “orientação”, a menos, naturalmente, de uma escolha inicial arbitrária.

O sinal do determinante divide o grupo dos automorfismos de $\mathbb{R}^n$ em duas partes complementares. Dizemos que um automorfismo $L \in \text{Aut}(\mathbb{R}^n)$ *preserva a orientação* (sem nem ter definido o que é a orientação!) se o seu determinante é positivo, e que *inverte a orientação* se o seu determinante é negativo. A composição de dois automorfismos que preservam a orientação também preserva a orientação, como natural. Da mesma forma, o inverso de um automorfismo que preserva a orientação também preserva a orientação. Por outro lado, a composição de dois automorfismos que invertem a orientação preserva a orientação.

Mas o que é, afinal, a orientação? Uma possibilidade é a seguinte. É natural chamar “positiva” a orientação da base ordenada canônica $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ de $\mathbb{R}^n$. Seja $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ uma base ordenada de $\mathbb{R}^n$. Existe um único automorfismo $L \in \text{Aut}(\mathbb{R}^n)$ que envia a base canônica nesta base, ou seja, tal que $L\mathbf{e}_k = \mathbf{b}_k$ para todo $k = 1, 2, \dots, n$. A *orientação* de $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ é então *positiva ou negativa* (ou a base em questão “é orientada positivamente ou negativamente”) dependendo se $\text{Det}L$ é positivo ou negativo. Isto significa que uma base é orientada positivamente sse é obtida da base canônica por meio de um automorfismo que preserva a orientação.

Por exemplo, a orientação natural de $\mathbb{R}^3$ é a orientação da base ordenada $\mathbf{i}, \mathbf{j}, \mathbf{k}$. É imediato verificar que são também orientadas positivamente as bases ordenadas $\mathbf{j}, \mathbf{k}, \mathbf{i}$ e $\mathbf{k}, \mathbf{i}, \mathbf{j}$ (ou seja, com a mesma ordem cíclica), por ser obtida da base canônica por meio dos automorfismos representados (na base canônica) pelas matrizes

$$\begin{pmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix} \quad \text{e} \quad \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{pmatrix}$$

respetivamente. Por outro lado, são orientadas negativamente as bases ordenadas $\mathbf{i}, \mathbf{k}, \mathbf{j}$, ou $\mathbf{j}, \mathbf{i}, \mathbf{k}$, ou $\mathbf{k}, \mathbf{j}, \mathbf{i}$.

Determinante e volume. Se

$$Q := [0, 1]^n = \{t_1 E_1 + t_2 E_2 + \dots + t_n E_n \in \mathbb{R}^n \text{ com } 0 \leq t_k \leq 1\}$$

denota o hiper-cubo unitário, ou seja, o paralelepípedo de lados $E_1, E_2, \dots, E_n$, então a imagem $L(Q)$ é o paralelepípedo de lados $L(E_1) = A_1, L(E_2) = A_2, \dots, L(E_n) = A_n$, que são as colunas da matriz A que define L na base canônica. O determinante do operador L é portanto o quociente

$$\text{Det}L = \pm \frac{\text{Vol}(L(Q))}{\text{Vol}(Q)},$$

o sinal sendo positivo ou negativo dependendo se L preserva ou não a orientação. Em geral, se $R \subset \mathbb{R}^n$ é uma região suficientemente regular, e portanto o seu volume pode ser aproximado com precisão arbitrária usando somas de volumes de hipercubos, então $\text{Det}L$ é igual a $\pm$ o quociente $\text{Vol}(L(R))/\text{Vol}(R)$. A conclusão é que o determinante de uma matriz A é o fator pelo qual a transformação L_A multiplica os volumes orientados. A multiplicatividade do determinante é uma consequência evidente desta interpretação.

ex: Calcule o determinate das transformações lineares

$$T(x, y) = (2x, 3y) \quad T(x, y) = (x, -y) \quad T(x, y) = (y, x)$$

$$T(x, y, z) = (3z, 2y, x) \quad T(x, y, z) = (x, x + y, x + y + z) \quad T(x, y, z) = (y, x, 0)$$

ex: O que pode dizer do determinante de uma projeção, um operador que satisfaz $P^2 = P$?

ex: O que pode dizer do determinante de uma involução, um operador que satisfaz $R^2 = I$?

12.5 Regra de Cramer

Regra de Cramer. Seja $A \in \text{Mat}_{n \times n}(\mathbb{R})$ uma matriz quadrada com $\text{Det}A \neq 0$. Então as suas colunas $A_{*1}, A_{*2}, \dots, A_{*n}$ geram o espaço $\mathbb{R}^n$, e a transformação linear $X \mapsto AX$ é invertível. Para todo vetor coluna $B \in \mathbb{R}^n$ existe portanto uma única solução do sistema linear

$$AX = B.$$

As coordenadas desta solução são os únicos coeficientes x_k 's tais que

$$x_1 A_{*1} + x_2 A_{*2} + \dots + x_n A_{*n} = B.$$

Se substituimos a k -ésima coluna A_{*k} da matriz A com o vetor B e calculamos o determinante, acontece que

$$\begin{aligned} D(A_{*1}, \dots, B, \dots, A_{*n}) &= D(A_{*1}, \dots, x_1 A_{*1} + x_2 A_{*2} + \dots + x_n A_{*n}, \dots, A_{*n}) \\ &= x_k D(A_{*1}, \dots, A_{*k}, \dots, A_{*n}) \end{aligned}$$

pela multilinearidade e a anti-simetria de D . Consequentemente,

Teorema 12.7 (regra de Cramer). *Seja $A \in \text{Mat}_{n \times n}(\mathbb{R})$ uma matriz quadrada com $\text{Det}A \neq 0$. As coordenadas x_k da única solução do sistema linear $AX = B$ são os quocientes*

$$x_k = \frac{\text{Det} \begin{pmatrix} a_{11} & \dots & b_1 & \dots & a_{1n} \\ a_{21} & \dots & b_2 & \dots & a_{2n} \\ \vdots & & \vdots & & \vdots \\ a_{n1} & \dots & b_n & \dots & a_{nn} \end{pmatrix}}{\text{Det} \begin{pmatrix} a_{11} & \dots & a_{1k} & \dots & a_{1n} \\ a_{21} & \dots & a_{2k} & \dots & a_{2n} \\ \vdots & & \vdots & & \vdots \\ a_{n1} & \dots & a_{nk} & \dots & a_{nn} \end{pmatrix}}$$

Apesar da sua elegância aparente, e da sua importância teórica, a regra de Cramer não é um método eficiente para resolver sistemas lineares, se comparada com outros métodos computacionais.

ex: Resolva os seguintes sistema utilizando a regra de Cramer

$$\begin{cases} 3x + 2y + z = 1 \\ 5x + 3y + 3z = 2 \\ -x + y + z = -1 \end{cases} \quad \begin{cases} 3x + 2y + z = 1 \\ 2x - 6y + 4z = 3 \\ x + y + z = -2 \end{cases}$$

$$\begin{cases} 2x + y = -6 \\ -x + 2y + 4z = 1 \\ -x + z = 3 \end{cases} \quad \begin{cases} 3x + y + z = 0 \\ 2x - y + 3z = 1 \\ x + y + z = 1 \end{cases}$$

Determinante de Vandermonde. Dados n números reais ou complexos $z_1, z_2, \dots, z_n$, a *matriz de Vandermonde* é a matriz $n \times n$ cujas linhas (ou colunas) são as progressões geométricas (até o grau $n - 1$) dos z_k 's, ou seja,

$$V := \begin{pmatrix} 1 & z_1 & z_1^2 & \dots & z_1^{n-1} \\ 1 & z_2 & z_2^2 & \dots & z_2^{n-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & z_n & z_n^2 & \dots & z_n^{n-1} \end{pmatrix}$$

O determinante de Vandermonde é o produto

$$\text{Det}V = \prod_{i < j} (z_j - z_i)$$

Determinante e matrizes inversas. Um matriz quadrada é invertível sse o seu determinante não é nulo. A regra de Cramer 12.7 sugere uma maneira de calcular a inversa.

Teorema 12.8. Uma matriz quadrada A é invertível sse $\text{Det}A \neq 0$, e a sua inversa é dada por

$$A^{-1} = \frac{1}{\text{Det}A} (\text{Cal}A)^{\top}$$

Demonstração. Se A é invertível, a sua matriz inversa $A^{-1} = (x_{ij})$ satisfaz

$$A A^{-1} = I.$$

As colunas da matriz identidade I são os vetores $E_1, E_2, \dots, E_n$ da base canónica. A identidade acima então diz que as colunas $X_1, X_2, \dots, X_n$ da matriz inversa A^{-1} são as soluções dos sistemas lineares $AX_k = E_k$. Pela regra de Cramer 12.7, as coordenadas x_{ik} de X_k , com $i = 1, 2, \dots, n$, são obtidas ao dividir por $\text{Det}A$ os determinantes das matrizes obtidas ao substituir à i -ésima coluna da matriz A o vetor E_k da base canónica. É imediato verificar que x_{ik} é então o elemento ki da matriz $\text{Cal}A$ dos complementos algébricos de A . $\square$

As mesmas considerações que fizemos sobre a regra de Cramer são aplicáveis a esta fórmula: não é um método eficiente para calcular a inversa de uma matriz de dimensão grande.

ex: Calcule a inversa das seguintes matrizes

$$\begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 1 & 3 \\ 7 & 5 \end{pmatrix}$$

ex: Determine os valores de λ para os quais $\lambda I - A$ é singular, quando

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 2 \end{pmatrix} \quad A = \begin{pmatrix} 1 & 1 \\ 1 & 2 \end{pmatrix}$$

$$A = \begin{pmatrix} 1 & 0 & 2 \\ 0 & -1 & -2 \\ 2 & -2 & 0 \end{pmatrix} \quad A = \begin{pmatrix} 11 & -2 & 8 \\ 19 & -3 & 14 \\ -8 & 2 & -5 \end{pmatrix}$$

ex: [Ap69] 3.17.

13 Valores e vetores próprios

ref: [Ap69] Vol 2, 4.1-10 ; [La87] Ch. VIII

13.1 Subespaços invariantes

15 dez 2022

Subespaços invariantes. Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador definido no espaço vetorial $\mathbf{V}$, real ou complexo. Um subespaço linear $W \subset \mathbf{V}$ é *-invariante*, ou *estável*, (ou, melhor, *L-invariante* quando o operador não é subentendido) se

$$L(W) \subset W$$

ou seja, se $\mathbf{w} \in W$ implica $L\mathbf{w} \in W$.

Subespaços invariantes triviais são o subespaço nulo 0 e o próprio $\mathbf{V}$. Também são subespaços invariantes o núcleo $\text{Ker}(L)$ e a imagem $\text{Im}(L)$, assim como núcleos e imagens das potências L^k .

Existe uma maneira canónica de produzir subespaços invariantes não triviais. A *órbita* do vetor $\mathbf{v}$ pelo operador L é o conjunto

$$\mathcal{O}_L^+(\mathbf{v}) := \{L^k \mathbf{v}, k = 0, 1, 2, 3, \dots\} = \{\mathbf{v}, L\mathbf{v}, L^2\mathbf{v}, L^3\mathbf{v}, \dots\}$$

Naturalmente, a órbita do vetor nulo é o conjunto formado apenas pelo próprio vetor nulo. Em geral, uma órbita é um conjunto não vazio e numerável, possivelmente finito. É finito se o vetor $\mathbf{v}$ é periódico, ou seja, se existe um inteiro minimal $k \geq 1$ tal que $L^k \mathbf{v} = \mathbf{v}$, ou se alguma imagem $L^n \mathbf{v}$ de $\mathbf{v}$ é periódica. Por exemplo, as órbitas de um operador nilpotentes são todas finitas. O subespaço $\text{Span}(\mathcal{O}_L^+(\mathbf{v}))$ gerado pela órbita do vetor $\mathbf{v}$ é chamado *subespaço cíclico* gerado por $\mathbf{v}$. É evidente que é um subespaço invariante (pois $L(L^k \mathbf{v}) = L^{k+1} \mathbf{v}$, logo a própria órbita é um conjunto invariante), e contém vetores não nulos se $\mathbf{v} \neq \mathbf{0}$. Pode ser também caracterizado como o “menor” subespaço invariante que contém o vetor $\mathbf{v}$, ou seja, a interseção de todos os subespaços invariantes que contém $\mathbf{v}$ (a família destes subespaços não é vazia porque contém pelo menos o espaço total).

ex: Mostre que o núcleo $\text{Ker}(L)$ e a imagem $\text{Im}(L)$ são subespaços L -invariantes.

ex: Mostre que $\text{Ker}(L^k)$ e $\text{Im}(L^k)$ são subespaços L -invariantes.

ex: Uma interseção de subespaços invariantes é também invariante?

ex: Verifique que se W é um subespaço invariante para o operador L , então também é invariante para todas as suas potências L^k , com $k \geq 1$.

ex: Se $p(z)$ é um polinómio, então $\text{Ker}(p(L))$ é L -invariante (observe que $p(L)$ comuta com L).

ex: Verifique que $\text{Span}(\mathcal{O}_L^+(\mathbf{v}))$ é um subespaço invariante.

ex: Determine os subespaços invariantes da transformação $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ que transforma cada ponto (x, y) no seu simétrico em relação à reta $y = x$.

ex: Determine os subespaços invariantes da transformação $T : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ que transforma cada ponto (x, y, z) no seu simétrico em relação ao plano $z = 0$.

ex: Determine os subespaços invariantes da transformação $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ que transforma cada ponto (x, y) na sua projecção ortogonal sobre a reta $y = x$.

ex: Para quais valores de θ existem subespaços invariantes não triviais de uma rotação do plano de um ângulo θ ?

Polinômios e quase-polinômios. O subespaço $\text{Pol}(\mathbb{R}) \subset \mathcal{C}^\infty(\mathbb{R})$ dos polinômios é um subespaço invariante do operador derivação, definido por $(Df)(t) := f'(t)$, e do operador multiplicação, definido por $(Xf)(t) = tf(t)$. O subespaço $\text{Pol}_n(\mathbb{R}) \subset \text{Pol}(\mathbb{R})$ dos polinômios de grau $\leq n$ é também um subespaço invariante do operador derivação mas não do operador multiplicação.

Fixado um expoente λ , o subespaço $\text{QPol}^\lambda(\mathbb{R}) \subset \mathcal{C}^\infty(\mathbb{R})$ dos *quase-polinômios* é o subespaço formado pelos produtos $p(t)e^{\lambda t}$, com $p(t)$ polinômio. Pode ser pensado como o subespaço cíclico para o operador multiplicação gerado pelo vetor $e^{\lambda t}$. É também um subespaço invariante para o operador derivação, e portanto para qualquer operador diferencial com coeficientes constantes $L = a_n D^n + \dots + a_1 D + a_0$. Esta observação justifica o “método dos coeficientes indeterminados” para encontrar uma solução particular de uma EDO linear $Lf = g$ quando a “força” $g(x)$ é um quase-polinômio. As mesmas considerações se aplicam a quase-polinômios com coeficientes complexos, e portanto com expoentes complexos $\lambda + i\omega$, cuja parte imaginária representa uma “frequência”.

Espaço de Schwartz. No estudo das EDPs da física matemática, é oportuno considerar funções integráveis, ou de quadrado integrável (ou seja, de energia finita). Isto quer dizer que as funções devem decair, ou seja, ter limite $|f(t)| \rightarrow 0$ quando $t \rightarrow \pm\infty$, suficientemente rápido. Na análise de Fourier, o instrumento básico para tratar EDPs lineares, são particularmente importantes os operadores derivação e multiplicação (que também correspondem aos operadores momento linear e posição da mecânica quântica), pois estes dois operadores são “entrelaçados” pela transformada de Fourier. É então útil considerar um espaço de funções integráveis que seja invariante para estes dois operadores e no entanto suficientemente grande para poder aproximar arbitrariamente bem funções integráveis arbitrárias. Este é o famoso *espaço de Schwartz* (Laurent, um francês). Tecnicamente é definido como o espaço $\mathcal{S}(\mathbb{R})$ das funções infinitamente deriváveis na reta real tais que para todos inteiros $m, n \geq 0$

$$\|f\|_{m,n} := \sup_{t \in \mathbb{R}} |t^m f^{(n)}(t)| < \infty$$

ou seja, que decaem, juntamente com todas as derivadas $f^{(n)}$, mais rápido do que qualquer polinômio. É claro que os operadores X e D deixam este espaço invariante.

Subespaços invariantes e matrizes em blocos. Se $W \subset \mathbf{V}$ é um subespaço invariante não trivial (para que a ideia seja interessante) do operador $L : \mathbf{V} \rightarrow \mathbf{V}$, então é possível definir a restrição $L|_W : W \rightarrow W$ do operador L ao subespaço W , que é naturalmente um operador linear. Se o subespaço invariante tem dimensão finita, por exemplo $W \approx \mathbb{R}^n$, e se $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ é uma sua base, então a restrição $L|_W$ é definida, nesta base, por uma matriz $A \in \text{Mat}_{n \times n}(\mathbb{R})$. Se também $\mathbf{V}$ tem dimensão finita, por exemplo $\mathbf{V} \approx \mathbb{R}^{n+m}$, é então possível, pelo teorema 4.5, completar o sistema a uma base $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n, \mathbf{e}_{n+1}, \dots, \mathbf{e}_{n+m}$ do próprio $\mathbf{V}$. Nesta base, o operador L é definido por uma “matriz em blocos”

$$\begin{pmatrix} A & C \\ 0 & B \end{pmatrix} \quad (13.1)$$

onde $C \in \text{Mat}_{n \times m}(\mathbb{R})$ e $B \in \text{Mat}_{m \times m}(\mathbb{R})$, pois as suas primeiras n colunas, que são as imagens dos primeiros n vetores da base, devem ser vetores de W .

O espaço quociente $\mathbf{V}/W$, o espaço das classes de equivalência $[\mathbf{v}] := \mathbf{v} + W$, admite uma base formada pelas classes $[\mathbf{e}_{n+1}], [\mathbf{e}_{n+2}], \dots, [\mathbf{e}_{n+m}]$, e é portanto isomorfo a $\mathbf{V}/W \approx \mathbb{R}^m$. O operador L induz um “operador quociente” $M : \mathbf{V}/W \rightarrow \mathbf{V}/W$, definido por

$$M[\mathbf{v}] := [L\mathbf{v}] \quad \text{ou seja,} \quad M(\mathbf{v} + W) := L\mathbf{v} + W$$

(a definição é bem posta justamente porque $L(W) \subset W$, e portanto a classe de $L\mathbf{v}$ não depende do representante $\mathbf{v}$ da classe de equivalência $[\mathbf{v}]$). A matriz que representa o operador quociente M na base $[\mathbf{e}_{n+1}], [\mathbf{e}_{n+2}], \dots, [\mathbf{e}_{n+m}]$ é precisamente a matriz B em (13.1).

e.g. Deslizamentos. O “deslizamento/cisalhamento” (em inglês, *shear*) horizontal é a transformação $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ definida por $T(x, y) = (x + y, y)$, ou seja, induzida pela matriz em blocos

$$\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$$

A reta $\mathbb{R}\mathbf{i}$ é um espaço invariante, o único não trivial, e a restrição de T a esta reta é a transformação identidade $T(x, 0) = (x, 0)$

ex: Determine os subespaços invariantes de um deslizamento vertical, a transformação $T(x, y) = (x, x + y)$ induzida pela matriz

$$\begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix}$$

Somas diretas de operadores e matrizes diagonais em blocos. Pode acontecer que o espaço total é uma soma direta $\mathbf{V} = W \oplus W'$ de dois subespaços invariantes, W e W' , para o endomorfismo L . Neste caso, existe uma base $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n, \mathbf{e}_{n+1}, \dots, \mathbf{e}_{n+m}$ de $\mathbf{V}$ formada por uma base $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ de W e uma base $\mathbf{e}_{n+1}, \mathbf{e}_{n+2}, \dots, \mathbf{e}_{n+m}$ de W' . A matriz que representa o operador L nesta base é então uma matriz “diagonal em blocos”

$$\begin{pmatrix} A & 0 \\ 0 & B \end{pmatrix}$$

Os blocos $A \in \text{Mat}_{n \times n}(\mathbb{R})$ e $B \in \text{Mat}_{m \times m}(\mathbb{R})$ representam as restrições de L aos subespaços invariantes W e W' , respectivamente. O operador é uma soma direta $L = L_A \oplus L_B$ dos operadores $L|_W : W \rightarrow W$ e $L|_{W'} : W' \rightarrow W'$, definidos pelas matrizes A e B nas bases escolhidas.

Da mesma forma, se o espaço $\mathbf{V}$ é uma soma direta $\mathbf{V} = V_1 \oplus V_2 \oplus \dots \oplus V_m$ de um número finito de subespaços invariantes V_k 's, então o operador é uma soma direta $L = L_1 \oplus L_2 \oplus \dots \oplus L_m$ das suas restrições $L_k = L|_{V_k} : V_k \rightarrow V_k$, e é definido, numa base oportuna composta por bases dos diferentes V_k 's, por uma matriz diagonal em blocos,

$$\begin{pmatrix} A_1 & 0 & \dots & 0 \\ 0 & A_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & A_m \end{pmatrix}$$

formada por m blocos.

É uma estratégia natural tentar representar desta forma um operador genérico, assim reduzindo a sua complexidade a um número finito de operadores mais simples, associados a blocos A_k 's que já não podem ser decompostos, por exemplo de dimensão menor possível. O caso mais desejável é, naturalmente, o caso de blocos de dimensão apenas um, que correspondem a operadores L_k que são homotetias de retas, do género $x \mapsto \lambda x$. Esta é a ideia dos ...

13.2 Valores e vetores próprios

Valores e vetores próprios. Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador definido num espaço vetorial $\mathbf{V}$, real ou complexo. Um *vetor próprio* ou *autovetor* (em inglês, *proper vector* ou *eigenvector*, do alemão *eigen* = próprio) de L é um vetor não nulo $\mathbf{v} \in \mathbf{V}$ cuja imagem é proporcional ao próprio vetor, ou seja, tal que

$$L\mathbf{v} = \lambda\mathbf{v}$$

onde λ é um escalar, real ou complexo. O escalar λ , tal que existe um vetor não nulo $\mathbf{v}$ que satisfaz $L\mathbf{v} = \lambda\mathbf{v}$, é chamado *valor próprio* ou *autovalor* (em inglês *eigenvalue*) do operador L (associado ao vetor próprio $\mathbf{v}$). Se $\mathbf{v}$ é um vetor próprio com valor próprio λ , então todo vetor não nulo $\mathbf{v}' = \alpha\mathbf{v}$ proporcional a $\mathbf{v}$ é também um vetor próprio com valor próprio λ . Portanto, um vetor próprio de L é um vetor (necessariamente não nulo) $\mathbf{v}$ que gera uma reta $\mathbb{R}\mathbf{v}$ (ou $\mathbb{C}\mathbf{v}$ se o espaço é complexo) que é um subespaço invariante de dimensão um para L .

Um vetor próprio de $\mathbf{v}$ com valor próprio λ é um vetor não nulo do núcleo do operador

$$L_\lambda = \lambda I - L$$

Portanto, λ é um valor próprio de L sse L_λ não é injetivo. Em particular, 0 é um valor próprio de L sse o operador L não é injetivo.

Se o espaço $\mathbf{V}$ tem dimensão finita, então o escalar λ é um valor próprio do operador L sse o operador $\lambda - L$ não é invertível (que, em dimensão finita, é equivalente a não ser injetivo ou a não ser sobrejetivo). Esta afirmação é falsa em dimensão infinita.

Se $\mathbf{V}$ tem dimensão finita, fixada uma base, um operador L é representado por uma matriz quadrada A . Então um vetor próprio é um vetor coluna X que satisfaz

$$AX = \lambda X$$

Neste sentido é também usual falar de vetores e valores próprios de matrizes quadradas.

Espaços próprios e multiplicidade geométrica. Se λ é um valor próprio do operador $L : \mathbf{V} \rightarrow \mathbf{V}$, então o núcleo

$$E_\lambda := \text{Ker}(\lambda - L) = \{\mathbf{v} \in \mathbf{V} \text{ t.q. } L\mathbf{v} = \lambda\mathbf{v}\}$$

é um subespaço invariante não nulo para L , dito *espaço próprio* associado ao valor próprio λ . Um caso particular é $E_0 = \text{Ker}(L)$. A restrição do operador linear L a cada espaço próprio E_λ é uma homotetia $\mathbf{v} \mapsto \lambda\mathbf{v}$.

Se $\mathbf{V}$ tem dimensão finita, então também cada espaço próprio E_λ em dimensão finita. A dimensão $\dim(E_\lambda)$ é chamada *multiplicidade geométrica* do valor próprio λ . Se o operador L é representado, fixada uma base, pela matriz quadrada A , então E_λ é o espaço das soluções do sistema homogêneo

$$(\lambda I - A)X = 0$$

que é possível (pois contém pelo menos um vetor próprio com valor próprio λ) e indeterminado (pois todo múltiplo não nulo do vetor próprio também é um vetor próprio).

Os diferentes operadores $\lambda - L$, ao variar o escalar λ (independentemente de λ ser ou não um valor próprio), comutam, pois L comuta com o próprio L e com todos os múltiplos da identidade. Em particular, se $\mathbf{v}$ é um vetor próprio com valor próprio λ_k , então $\mathbf{v}$ está no núcleo de todo produto $(\lambda_1 - L) \dots (\lambda_k - L) \dots (\lambda_m - L)$. Esta observação é a chave do seguinte resultado, que mostra que vetores próprios associados a valores próprios distintos são independentes.

Teorema 13.1. Se $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ são vetores próprios do operador L , e se os correspondentes valores próprios $\lambda_1, \lambda_2, \dots, \lambda_m$ são dois a dois distintos (ou seja, $\lambda_i \neq \lambda_j$ se $i \neq j$), então os vetores $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ são linearmente independentes.

Demonstração. Seja

$$c_1\mathbf{v}_1 + c_2\mathbf{v}_2 + \dots + c_m\mathbf{v}_m = \mathbf{0}$$

uma combinação linear nula dos vetores próprios $\mathbf{v}_k$ associados a valores próprios diferentes λ_k . Se aplicamos o operador $(\lambda_2 - L)(\lambda_3 - L) \dots (\lambda_m - L)$ aos dois membros desta identidade, temos

$$c_1(\lambda_2 - \lambda_1)(\lambda_3 - \lambda_1) \dots (\lambda_m - \lambda_1)\mathbf{v}_1 = \mathbf{0}$$

e portanto $c_1 = 0$. Se depois aplicamos o operador $(\lambda_1 - L)(\lambda_3 - L) \dots (\lambda_m - L)$, obtemos $c_2 = 0$. Continuando assim, obtemos $c_k = 0$ para todos os k . $\square$

ex: Todo vetor não nulo é um vetor próprio da transformação nula, com valor próprio $\lambda = 0$.

ex: Todo vetor não nulo é um vetor próprio da transformação identidade, com valor próprio $\lambda = 1$. Em geral, todo vetor $\mathbf{v} \neq \mathbf{0}$ é um vetor próprio de uma homotetia $\mathbf{v} \mapsto \lambda\mathbf{v}$ (cuja matriz é λI em qualquer base) com valor próprio λ .

ex: Determine valores e vetores próprios da transformação $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ que transforma cada ponto (x, y) no seu simétrico em relação à reta $y = 2x$.

ex: Determine valores e vetores próprios da transformação $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ que transforma cada ponto (x, y) na sua projeção ortogonal sobre a reta $3y = x$.

ex: Existem operadores sem valores próprios. Por exemplo, uma rotação $R_\theta : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, definida no espaço vetorial real $\mathbb{R}^2$ por

$$(x, y) \mapsto (x \cos \theta - y \sin \theta, x \sin \theta + y \cos \theta),$$

não admite vetores próprios se o ângulo θ não é um múltiplo inteiro de π . No entanto, a mesma rotação pode ser pensada como a transformação $z \mapsto e^{i\theta}z$ definida no espaço vetorial complexo $\mathbb{C} \approx \mathbb{R}^2$, onde $z = x + iy \approx (x, y)$. Neste caso, todo vetor $z \neq 0$ é um vetor próprio, de valor próprio $e^{i\theta}$.

ex: Determine os valores e os vetores próprios das transformações

$$\begin{aligned} L(x, y) &= (x, 0) & L(x, y) &= (x/2, 3y) & L(x, y) &= (-y, x) \\ L(x, y) &= (x, x + y) & L(x, y) &= (x + \lambda y, y) & L(x, y) &= (x + \alpha y, y) \\ L(x, y, z) &= (0, y, -z) & L(x, y, z) &= (y, z, x) \end{aligned}$$

ex: Sejam $\mathbf{v}$ e $\mathbf{w}$ vetores próprios do operador L , associados a valores próprios λ e μ , respetivamente. Mostre que se também a soma $\mathbf{v} + \mathbf{w}$ é um vetor próprio de L então necessariamente $\lambda = \mu$.

ex: Se λ é um valor próprio do operador L e c é um escalar, então $c\lambda$ é um valor próprio do operador cL .

ex: Se λ é um valor próprio do operador L e k é um inteiro positivo, então λ^k é um valor próprio da potência L^k (o vetor próprio é o mesmo). Consequentemente, se $p(z) = a_0 + a_1z + \dots + a_mz^m$ é um polinómio, então $p(\lambda)$ é um valor próprio do operador $p(L) = a_0I + a_1L + \dots + a_mL^m$.

ex: No entanto, uma potência L^k pode ter também outros valores próprios, que não são potências dos valores próprios de L . Por exemplo, uma rotação $R_{\pi/2}$ de um ângulo $\pi/2$ no plano não tem valores próprios, mas o seu quadrado $R_{\pi/2}^2 = R_\pi$ admite o valor próprio -1 (cuja raiz quadrada não é um número real!), e a sua quarta potência $R_{\pi/2}^4 = I$ admite um valor próprio 1 . Dê outros exemplos.

ex: Em particular, os valores próprios de um operador nilpotente (um operador tal que alguma potência $k \geq 1$ é o operador nulo $L^k = 0$) não podem ser diferentes de zero.

ex: Observe que $L^2 - \lambda^2 = (L - \lambda)(L + \lambda)$, e deduza que se λ^2 é um valor próprio de L^2 , então L admite um valor próprio igual a λ ou $-\lambda$.

ex: Se L é um operador invertível e $\mathbf{v}$ é um vetor próprio de L com valor próprio λ , necessariamente diferente de zero (pois L é injetivo), então $\mathbf{v}$ é um vetor próprio de L^{-1} com valor próprio λ^{-1} .

ex: Seja A uma matriz $n \times n$, e sejam D e P as matrizes $(2n) \times (2n)$ em blocos definidas por

$$D = \begin{pmatrix} 0 & A \\ A & 0 \end{pmatrix} \quad P = \begin{pmatrix} I & 0 \\ 0 & -I \end{pmatrix}$$

Verifique que $P^2 = I$ e que P e D anti-comutam, ou seja, satisfazem $PD + DP = 0$. Mostre que se $\mathbf{v}$ é um vetor próprio de D com valor próprio λ , ou seja, $D\mathbf{v} = \lambda\mathbf{v}$, então $\mathbf{w} = P\mathbf{v}$ é também um vetor próprio de D , com valor próprio $-\lambda$ (“anti-partículas”).

ex: [Ap69] Vol 2 4.4.

Exponencial e funções trigonométricas. Os operadores básicos do cálculo são, naturalmente, o operador derivação D e as suas potências, em particular o laplaciano $\Delta = D^2$, definidos por exemplo no espaço $\mathcal{C}^\infty(\mathbb{R})$ das funções infinitamente deriváveis de uma variável real. Não surpreende que as funções próprias destes operadores sejam as funções mais importantes da análise. As funções próprias da derivada resolvem $Df = \lambda f$, ou seja, a equação diferencial

$$f'(t) = \lambda f(t)$$

e são portanto proporcionais aos exponenciais $f(t) = e^{\lambda t}$. As funções próprias do laplaciano resolvem $\Delta f = \lambda f$, ou seja, a equação diferencial

$$f''(t) = \lambda f(t)$$

Se $\lambda = 0$, esta é a equação de Newton do movimento retilíneo uniforme, cujas soluções são as retas afins $f(t) = a + bt$. Se $\lambda = k^2$ é positivo, as soluções são os exponenciais $f(t) = e^{\pm kt}$ (ou, tomando combinações lineares oportunas, cossenos e senos hiperbólicos). Se, finalmente, $\lambda = -\omega^2$ é negativo, então esta é a equação do “oscilador harmónico” $f'' = -\omega^2 f$, cujas soluções são as funções trigonométricas $\cos(\omega t)$ e $\sin(\omega t)$, que descrevem oscilações com frequência ω .

Num espaço complexo (ou, melhor, no espaço vetorial complexo das funções complexas de uma variável real), não há diferença entre exponenciais e funções trigonométricas pois, de acordo com a fórmula de Euler (6.9) e (6.11), $e^{(k+i\omega)t} = e^{kt} (\cos(\omega t) + i \sin(\omega t))$.

ex: Se $\lambda_1, \dots, \lambda_n$ são distintos, então as funções $e^{\lambda_1 t}, e^{\lambda_2 t}, \dots, e^{\lambda_n t}$ são linearmente independentes.

ex: Mostre que $\cos(\omega t)$ e $\sin(\omega t)$ são linearmente independentes, se $\omega \neq 0$.

ex: Verifique que os exponenciais $e^{\lambda t}$ são também funções próprias dos operadores *translação*, definidos por $(T_s f)(t) = f(t + s)$, com $s \in \mathbb{R}$ (isto não é surpreendente, pois mostraremos que T_s é o exponencial de sD ...).

Laplaciano na circunferência. Considere o espaço vetorial $\mathcal{C}^\infty(\mathbb{R}/2\pi\mathbb{Z})$ das funções $f(t)$ infinitamente deriváveis na reta real e periódicas de período 2π , ou seja, satisfazendo $f(t + 2\pi) = f(t)$ para todo t (que portanto podem ser pensadas funções definidas na circunferência $\mathbb{R}/2\pi\mathbb{Z}$). Mostre que, para cada inteiro não nulo n , o plano gerado por $\cos(nt)$ e $\sin(nt)$ é um espaço próprio do operador *laplaciano* $\Delta = D^2$, com valor próprio $\lambda = -n^2$. Determine o espaço próprio associado ao valor próprio $\lambda = 0$. Existem outros valores próprios?

Operadores diferenciais, translações e ondas planas. Sejam $\Omega \subset \mathbb{R}^n$ um aberto e $\mathcal{C}^\infty(\Omega)$ o espaço vetorial complexo das funções $f : \Omega \rightarrow \mathbb{C}$ infinitamente diferenciáveis. Dado um multi-índice $\alpha = (\alpha_1, \alpha_2, \dots, \alpha_n) \in \mathbb{N}^n$, de grau $|\alpha| := \alpha_1 + \alpha_2 + \dots + \alpha_n$, o operador diferencial $\partial^\alpha : \mathcal{C}^\infty(\Omega) \rightarrow \mathcal{C}^\infty(\Omega)$ é definido por

$$(\partial^\alpha f)(\mathbf{x}) := \frac{\partial^{|\alpha|} f}{\partial x_1^{\alpha_1} \partial x_2^{\alpha_2} \dots \partial x_n^{\alpha_n}}(\mathbf{x}).$$

As ondas planas $e_\xi(\mathbf{x}) := e^{i\xi \cdot \mathbf{x}}$, com $\xi \in (\mathbb{R}^n)^* \approx \mathbb{R}^n$, são funções próprias dos operadores diferenciais ∂^α , com $\alpha \in \mathbb{N}^n$, com valores próprios $(i\xi)^\alpha := (i\xi_1)^{\alpha_1} (i\xi_2)^{\alpha_2} \dots (i\xi_n)^{\alpha_n}$, ou seja,

$$\partial^\alpha e_\xi = (i\xi)^\alpha e_\xi.$$

O operador de *translação* $T_{\mathbf{a}}$, com $\mathbf{a} \in \mathbb{R}^n$, é definido por

$$(T_{\mathbf{a}} f)(\mathbf{x}) := f(\mathbf{x} + \mathbf{a}).$$

As ondas planas $e_\xi(\mathbf{x}) = e^{i\xi \cdot \mathbf{x}}$ são também funções próprias dos operadores de translação com valores próprios $\lambda_{\mathbf{a}}(\xi) = e^{i\xi \cdot \mathbf{a}}$, ou seja,

$$T_{\mathbf{a}} e_\xi = e^{i\xi \cdot \mathbf{a}} e_\xi.$$

O operador de modulação M_{ξ} , com $\xi \in (\mathbb{R}^n)^* \approx \mathbb{R}^n$, é definido por

$$(M_{\xi}f)(\mathbf{x}) := e^{i\xi \cdot \mathbf{x}} f(\mathbf{x}).$$

É imediato verificar que $T_{\mathbf{a}}M_{\xi} = e^{i\xi \cdot \mathbf{a}}M_{\xi}T_{\mathbf{a}}$. Os operadores translação e modulação geram o *grupo de Heisenberg*.

Existência de valores próprios. As rotações do plano real são exemplos de operadores sem valores próprios. Também é fácil fazer exemplos de operadores sem valores próprios em dimensão infinita, como o operador primitivação. Por outro lado, se o espaço é complexo e se a dimensão é finita, então o teorema fundamental da álgebra 6.2 (um teorema de análise sobre zeros de polinômios complexos!) garante a existência de valores próprios. A prova mais popular deste resultado usa a noção de “polinômio caraterístico” de uma matriz, que será introduzido mais à frente e que é um determinante de uma matriz dependente de um parâmetro. Por outro lado, a prova mais transparente é a prova de Sheldon Axler²⁷, explicada em [Ax15], e que dispensa o uso dos determinantes.

Teorema 13.2. *Todo operador de um espaço vetorial complexo de dimensão finita admite (pelo menos) um valor próprio.*

Demonstração. Seja $\mathbf{V} \approx \mathbb{C}^n$ um espaço vetorial complexo de dimensão n , e seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador. A órbita $\mathcal{O}_L^+(\mathbf{v})$ de um vetor não nulo $\mathbf{v} \in \mathbf{V}$ não pode conter mais de n vetores independentes. Em particular, os primeiro $n + 1$ vetores desta órbita devem ser linearmente dependentes. Isto significa que existem coeficientes complexos $a_0, a_1, a_2, \dots, a_n$, não todos nulos, tais que

$$(a_0 + a_1L + a_2L^2 + \dots + a_nL^n)\mathbf{v} = \mathbf{0}$$

O polinômio $f(z) = a_0 + a_1z + a_2z^2 + \dots + a_nz^n$ não é constante se $\mathbf{v}$ não é o vetor nulo. Pelo teorema fundamental da álgebra 6.2, fatoriza num produto $f(z) = a_m(z - \lambda_1)(z - \lambda_2) \dots (z - \lambda_m)$, sendo λ_k 's as suas raízes, não necessariamente distintas, e $a_m \neq 0$ se $1 \leq m \leq n$ é o grau de f . Então

$$f(L)\mathbf{v} = (L - \lambda_1)(L - \lambda_2) \dots (L - \lambda_m)\mathbf{v} = \mathbf{0}$$

(como já observado, os $L - \lambda_k$ comutam, assim que a ordem é indiferente). Isto claramente diz que existe um vetor não nulo no núcleo de pelo menos um dos $L - \lambda_k$, e consequentemente este λ_k é um valor próprio de L . $\square$

ex: Mostre que se $\mathbf{v}$ é um vetor não nulo tal que $A_1A_2 \dots A_m\mathbf{v} = \mathbf{0}$, então existe um vetor não nulo no núcleo de um dos operadores A_k .

ex: A existência ou menos de valores próprios depende, naturalmente, do espaço onde os operadores estão definido. Por exemplo, os operadores derivação e multiplicação, definidos por $(Df)(t) = f'(t)$ e $(Xf)(t) := tf(t)$ respetivamente, não admitem vetores próprios no subespaço $\text{Pol}(\mathbb{R}) \subset \mathcal{C}^\infty(\mathbb{R})$ dos polinômios. Mostre que, por outro lado, os vetores próprios do operador $L = XD$ em $\text{Pol}(\mathbb{R})$ são os monômios $f(t) = t^n$.

ex: Considere o operador primitivação, definido por $(Pf)(x) := \int_0^x f(t) dt$ no espaço $\mathcal{C}^\infty(\mathbb{R})$. Mostre que P não tem valores próprios (derive a identidade $Pf = \lambda f$ para obter uma equação diferencial para o suposto vetor próprio $f(t)$, e observe que a mesma identidade também implica uma condição inicial $f(0) \dots$)

²⁷S. Axler, Down With Determinants!, *American Mathematical Monthly* **102** (1995), 139-154.

ex: Em dimensão infinita, é possível que $\lambda - L$ não seja invertível sem que λ seja um valor próprio. Por exemplo, o operador *deslocamento à direita* (em inglês, *right shift*)

$$S : (x_1, x_2, x_3, \dots) \mapsto (0, x_1, x_2, \dots),$$

definido no espaço $\mathbb{R}^{\mathbb{N}}$ das sucessões reais $\mathbf{x} = (x_1, x_2, x_3, \dots)$, não é invertível. No entanto, 0 não é um valor próprio de S , pois a única solução de $S\mathbf{x} = 0\mathbf{x}$ é a sucessão nula $(0, 0, 0, \dots)$.

13.3 Operadores diagonalizáveis

Operadores diagonalizáveis. Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador definido no espaço linear $\mathbf{V}$. O operador L é dito *diagonalizável* se existe uma decomposição de $\mathbf{V}$ como soma direta finita

$$\mathbf{V} = E_{\lambda_1} \oplus E_{\lambda_2} \oplus \dots \oplus E_{\lambda_m}$$

de espaços próprios E_{λ_k} associados aos valores próprios distintos λ_k de L . Isto significa que todo $\mathbf{v} \in \mathbf{V}$ é uma sobreposição única $\mathbf{v} = \mathbf{v}_1 + \mathbf{v}_2 + \dots + \mathbf{v}_m$ de vetores $\mathbf{v}_k \in E_{\lambda_k}$ tais que $L\mathbf{v}_k = \lambda_k \mathbf{v}_k$, e portanto o operador é uma soma direta $L = \oplus_k \lambda_k$ de homotetias nos espaços próprios. Em outras palavras, se $P_k : \mathbf{V} \rightarrow E_{\lambda_k}$ denota a projeção sobre E_{λ_k} (definida por $P_k(\mathbf{v}) = \mathbf{v}_k$), então

$$L = \sum_{k=1}^m \lambda_k P_k$$

Se $\mathbf{V}$ tem dimensão finita n , então um operador diagonalizável é representado por uma matriz diagonal

$$\Lambda = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix}$$

numa base formada por vetores próprios. Os elementos diagonais são os valores próprios, cada um repetido de acordo com a sua multiplicidade geométrica.

Lamentavelmente, nem todos os operadores, mesmo definidos em espaços de dimensão finita, são diagonalizáveis.

Semi-simplicidade. A caracterização algébrica do algoritmo indutivo que permite diagonalizar operadores é a seguinte. Um operador $L : \mathbf{V} \rightarrow \mathbf{V}$ é dito *semi-simples* se todo subespaço invariante $E \subset \mathbf{V}$ admite um complementar invariante, ou seja, um subespaço invariante $F \subset \mathbf{V}$ tal que $\mathbf{V} = E \oplus F$.

Teorema 13.3. *Um operador $L : \mathbf{V} \rightarrow \mathbf{V}$ de um espaço vetorial complexo de dimensão finita é diagonalizável sse é semi-simples.*

Demonstração. Pelo teorema 13.2, o operador admite um vetor próprio $\mathbf{v}$, logo um subespaço invariante de dimensão um $E = \mathbb{C}\mathbf{v}$. Se o operador é semi-simples, então o espaço é uma soma direta $\mathbf{V} = E \oplus F$, onde F é um subespaço invariante de dimensão $n - 1$. A implicação $\Rightarrow$ segue portanto por indução, sendo trivial em dimensão um. A implicação contrária $\Leftarrow$ é óbvia, pois subespaços invariantes de operadores diagonalizáveis são gerados por vetores próprios. $\square$

e.g. Projeções. Uma projeção, um operador $P : \mathbf{V} \rightarrow \mathbf{V}$ que verifica $P^2 = P$, é diagonalizável. Os valores próprios satisfazem $\lambda^2 = \lambda$, e portanto podem ser apenas 0 e 1. Os espaços próprios são $E_0 = \text{Ker}(P)$, que pode ser trivial se P é igual ao operador identidade (uma projeção trivial), e $E_1 = \text{Im}(P)$. O espaço total é uma soma direta $\mathbf{V} = E_0 \oplus E_1$ (eventualmente com $E_0 = 0$) e o operador é uma soma direta $P = 0 \oplus 1$.

De fato, um vetor genérico $\mathbf{v}$ é uma soma única $\mathbf{v} = \mathbf{v}_0 + \mathbf{v}_1$ de um vetor $\mathbf{v}_1 = P\mathbf{v}$ e de um vetor $\mathbf{v}_0 = \mathbf{v} - P\mathbf{v}$. Mas $P\mathbf{v}_0 = P(\mathbf{v} - P\mathbf{v}) = P\mathbf{v} - P^2\mathbf{v} = 0$, logo $\mathbf{v}_0 \in E_0 = \text{Ker}(P)$, e

$P\mathbf{v}_1 = P(P\mathbf{v}) = P\mathbf{v} = \mathbf{v}_1$, e portanto $\mathbf{v}_1 \in E_1 = \text{Im}(P)$. Se acontece que $\mathbf{v}_0$ é sempre o vetor nulo, então $P\mathbf{v} = \mathbf{v}$ para todo $\mathbf{v}$, o que significa que P é o operador identidade.

Em particular, o traço de uma matriz P que representa uma projeção é igual à dimensão da sua imagem, ou seja,

$$\text{Tr}(P) = \text{Rank}(P)$$

e.g. Involuções. Uma involução, um operador $R : \mathbf{V} \rightarrow \mathbf{V}$ que verifica $R^2 = I$, é também diagonalizável. Os seus valores próprios satisfazem $\lambda^2 = 1$, logo podem ser apenas ± 1 . Os operadores

$$P_{\pm} = \frac{1}{2}(I \pm R)$$

são projeções sobre os espaços próprios $E_{\pm 1}$, respetivamente, que satisfazem $P_+P_- = P_-P_+ = 0$. A involução é então uma diferença $R = P_+ - P_-$ entre duas projeções que comutam.

ex: Determine o operador diagonalizável $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ sabendo que os seus valores próprios são 2 e 3 e que os seus vetores próprios são $(1, 3)$ e $(-2, 1)$, respetivamente.

ex: Mostre que o operador $L : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, definido por $L(x, y) = (y, 0)$, não é diagonalizável.

ex: Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador tal que $L^3 = L$. Mostre que o espaço total é uma soma direta de espaços próprios $\mathbf{V} = E_0 \oplus E_1 \oplus E_{-1}$, e portanto L é diagonalizável.

ex: Dada uma matriz diagonal $\Lambda = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n) \in \text{Diag}_n(\mathbb{R})$, é possível considerar o operador $\mathcal{L}_{\Lambda} : \text{Mat}_{n \times n}(\mathbb{R}) \rightarrow \text{Mat}_{n \times n}(\mathbb{R})$ definido pelo comutador

$$\mathcal{L}_{\Lambda}A := [\Lambda, A] = \Lambda A - A\Lambda$$

Dada uma matriz diagonal invertível $D = \text{diag}(\delta_1, \delta_2, \dots, \delta_n) \in \text{Diag}_n(\mathbb{R})^{\times}$ (logo com entradas $\delta_k \neq 0$), é possível considerar o operador $\mathcal{M}_D : \text{Mat}_{n \times n}(\mathbb{R}) \rightarrow \text{Mat}_{n \times n}(\mathbb{R})$ definido em (10.6) pela conjugação

$$\mathcal{M}_DA := DAD^{-1}$$

Verifique que as matrizes I_{ij} definidas em (9.1) (com $1 \leq i, j \leq n$) são vetores próprios de $\mathcal{L}_{\Lambda}$, com valores próprios $\lambda_{ij} = \lambda_i - \lambda_j$, e também de $\mathcal{M}_D$, com valores próprios $\delta_{ij} = \delta_i/\delta_j$, ou seja,

$$\mathcal{L}_{\Lambda}I_{ij} = (\lambda_i - \lambda_j)I_{ij} \quad \text{e} \quad \mathcal{M}_DI_{ij} = \frac{\delta_i}{\delta_j}I_{ij}$$

Portanto, $\mathcal{L}_{\Lambda}$ e $\mathcal{M}_D$ são diagonalizáveis na base formada pelas I_{ij} 's.

13.4 Polinómio caraterístico

A estratégia usada para provar a existência no teorema 13.2 não é um método prático para encontrar todos os valores próprios de um operador, nem para decidir quantos são. Mais tradicional é a seguinte ideia, que utiliza os determinantes.

Polinómio caraterístico. Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador linear definido no espaço vetorial de dimensão finita $\mathbf{V}$, real ou complexo. Fixada uma base, o operador é definido pela equação matricial $X \mapsto AX$, onde A é uma matriz $n \times n$, real ou complexa. O escalar λ é um valor próprio do operador L sse operador $\lambda I - L$ não é injetivo, e portanto, sendo a dimensão finita, sse a matriz $\lambda I - A$ não é invertível. Pelo teorema 12.8, isto acontece sse

$$\text{Det}(\lambda I - A) = 0.$$

Esta observação sugere a seguinte definição. O *polinómio caraterístico* da matriz quadrada A é o polinómio

$$c_A(z) := \text{Det}(zI - A)$$

de uma variável z , com coeficientes reais ou complexos, dependendo se A é uma matriz real ou complexa. O polinómio caraterístico depende apenas do operador L e não da matriz que o representa. De fato, matrizes semelhantes (ou seja, matrizes que representam o mesmo operador em bases diferentes) têm o mesmo polinómio caraterístico, pois se $B = U^{-1}AU$ com U invertível, então pela multiplicatividade dos determinantes 12.4

$$\begin{aligned} \text{Det}(zI - B) &= \text{Det}(zI - U^{-1}AU) = \text{Det}(zU^{-1}IU - U^{-1}AU) \\ &= \text{Det}(U^{-1}(zI - A)U) = \text{Det}(zI - A) \end{aligned}$$

Finalmente, temos a seguinte caraterização dos valores próprios de um operador em dimensão finita.

Teorema 13.4. *Seja A a matriz quadrada que define o operador $L : \mathbf{V} \rightarrow \mathbf{V}$ numa base do espaço de dimensão finita $\mathbf{V}$, real ou complexo. O escalar λ é um valor próprio do operador L sse é uma raiz do polinómio caraterístico de A , ou seja, se $c_A(\lambda) = 0$.*

O polinómio caraterístico de uma matriz $n \times n$ é um polinómio mónico de grau n , com coeficientes reais ou complexos, dependendo se o espaço (e portanto a matriz) é real ou complexo. Se representa um operador definido num espaço real, pode não ter raízes, pois polinómios reais de grau par podem não ter raízes reais. Por outro lado, se a matriz representa um operador definido num espaço complexo, então o teorema de Gauss 6.1 garante a existência de valores próprios. Esta é a prova clássica do teorema 13.2. Ainda mais, pelo teorema fundamental da álgebra 6.2, o polinómio caraterístico fatoriza num produto

$$\begin{aligned} c_A(z) &= (z - \lambda_1)(z - \lambda_2) \dots (z - \lambda_n) \\ &= z^n - (\lambda_1 + \lambda_2 + \dots + \lambda_n)z^{n-1} + \dots + (-1)^n(\lambda_1 \lambda_2 \dots \lambda_n) \end{aligned}$$

onde os λ_k 's são as n raízes, não necessariamente distintas. Mas $c_A(0) = \text{Det}(-A) = (-1)^n \text{Det} A$, e portanto

$$\boxed{\text{Det} A = \lambda_1 \lambda_2 \dots \lambda_n} \quad (13.2)$$

Também uma comparação com a definição, ou um argumento por indução, mostra que

$$\text{Det} \begin{pmatrix} z - a_{11} & -a_{12} & \dots & -a_{1n} \\ -a_{21} & z - a_{22} & \dots & -a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ -a_{n1} & -a_{n2} & \dots & z - a_{nn} \end{pmatrix} = z^n - (a_{11} + a_{22} + \dots + a_{nn})z^{n-1} + \dots$$

e portanto o coeficiente de z^{n-1} é igual ao oposto do traço da matriz A . Consequentemente,

$$\boxed{\text{Tr} A = \lambda_1 + \lambda_2 + \dots + \lambda_n} \quad (13.3)$$

Estas fórmulas são também válidas no caso de matrizes reais, pois pelo teorema 6.3 as raízes complexas do polinómio caraterístico (pensado como um polinómio complexo, embora com coeficientes reais) ocorrem em pares de números complexos conjugados, assim que a soma e o produto de todas as raízes são sempre números reais. Finalmente, determinante e traço de um operador L , representado numa base arbitrária por uma matriz quadrada A , são iguais ao produto e a soma das raízes do polinómio caraterístico, respetivamente.

É claro que também os outros coeficientes do polinómio caraterístico são funções simétricas destas raízes, mas não têm a importância geométrica de determinante e traço ...

Também é possível juntar as raízes repetidas, e escrever

$$c_A(z) = (z - \lambda_1)^{n_1} (z - \lambda_2)^{n_2} \dots (z - \lambda_m)^{n_m}$$

onde agora os λ_k são os diferentes valores próprios (da matriz A pensada como matriz complexa), e os n_k 's são inteiros positivos que somam $\sum_k n_k = n$. O inteiro n_k é chamado *multiplicidade*

algébrica do valor próprio λ_k (embora exista uma definição mais “geométrica” deste número, que é a dimensão do “espaço próprio generalizado” $G_\lambda = \ker(\lambda - L)^n$, como explicado na próxima seção). O conjunto

$$\sigma(L) := \{\lambda_1, \lambda_2, \dots, \lambda_m\} \subset \mathbb{C}$$

das raízes do polinômio caraterístico é chamado *espectro* do operador L (ou, melhor, da “complexificação” do operador). O número

$$\rho(L) := \max_k |\lambda_k|$$

(ou seja, o menor raio $r \geq 0$ de um disco fechado $\overline{D_r(0)} \subset \mathbb{C}$ que contém todos os valores próprios de L) é chamado *raio espectral* do operador L , ou da matriz A que representa o operador numa base arbitrária.

Naturalmente, quando n não é muito pequeno, o próprio cálculo do determinante, e depois das raízes do polinômio caraterístico, continua sendo um problema difícil. Pode ser necessário usar métodos numéricos, como o método de Newton...

Se o operador é uma soma direta de dois ou mais operadores, então o seu polinômio caraterístico fatoriza no produto dos polinômios caraterísticos dos adendos. De fato, se $L = M \oplus N$ é representado, numa base oportuna, por uma matriz diagonal em blocos

$$A = \begin{pmatrix} B & 0 \\ 0 & C \end{pmatrix}$$

(onde B é a matriz de M e C é a matriz de N) então, pela (12.14), o seu polinômio caraterístico é um produto $c_A(z) = c_B(z) c_C(z)$ dos polinômios caraterísticos de B e C . O caso limite é o caso de um operador diagonalizável, que é uma soma direta $L = \lambda_1 \oplus \lambda_2 \oplus \dots \oplus \lambda_n$ de n homotetias da reta $x \mapsto \lambda_k x$. O seu polinômio caraterístico é então igual ao produto $(z - \lambda_1)(z - \lambda_2) \dots (z - \lambda_n)$.

ex: Mostre que A e $A^\top$ têm o mesmo polinômio caraterístico, e portanto os mesmos valores próprios.

ex: Mostre que uma matriz real $n \times n$ admite sempre um valor próprio se a dimensão n é ímpar.

ex: Verifique que o polinômio caraterístico da matriz 2×2 genérica

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

é $c_A(z) = z^2 - (a + d)z + (ad - bc)$. Verifique que $c_A(A) = 0$, ou seja, que

$$A^2 - (\text{Tr} A) A + (\text{Det} A) I = 0$$

ex: O único valor próprio de uma matriz nilpotente (uma matriz N tal que alguma potência $N^k = 0$) é $\lambda = 0$.

ex: O único valor próprio de uma matriz unipotente (uma matriz U tal que $U - I$ é nilpotente, assim que alguma potência $(U - I)^k = 0$) é $\lambda = 1$.

ex: Determine valores e vetores próprios dos endomorfismos definidos pelas seguintes matrizes

$$\begin{pmatrix} 2 & 0 \\ 0 & 1/3 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix} \quad \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$$

$$\begin{pmatrix} 2 & 5 & -1 \\ 0 & -3 & 1 \\ 0 & 0 & 7 \end{pmatrix} \quad \begin{pmatrix} 1 & 5 & -1 \\ 0 & -2 & 1 \\ -4 & 0 & 3 \end{pmatrix} \quad \begin{pmatrix} 2 & 1 & 1 \\ 2 & 3 & 2 \\ 3 & 3 & 4 \end{pmatrix}$$

$$\begin{pmatrix} 7 & 5 & 1 \\ 0 & -2 & 1 \\ 20 & 0 & 3 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{pmatrix}$$

ex: [Ap69] 4.10.

13.5 Diagonalização de matrizes

Diagonalização de matrizes. Uma matriz diagonal representa um operador numa base formada por vetores próprios. A manipulação de matrizes diagonais é particularmente simples. Por exemplo, as potências de uma matriz diagonal são também matrizes diagonais, e portanto polinômios, séries de potências, funções analíticas ... de matrizes diagonais são também matrizes diagonais, fáceis de calcular.

Pode ser útil, portanto, representar operadores diagonalizáveis com matrizes diagonais. Se um operador é definido por uma matriz quadrada A , por exemplo relativamente à base canónica de $\mathbb{R}^n$ ou $\mathbb{C}^n$, este processo é chamado “diagonalização”. A matriz quadrada A é dita *diagonalizável* se é semelhante a uma matriz diagonal. Isto acontece quando o operador linear definido na base canónica do espaço $\mathbb{R}^n$ ou $\mathbb{C}^n$ pela matriz quadrada A admite n vetores próprios linearmente independentes $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$, com valores próprios $\lambda_1, \lambda_2, \dots, \lambda_n$, respetivamente (não necessariamente distintos). Então, se U denota a matriz (invertível, pela independência dos $\mathbf{v}_k$'s) cujas colunas são os vetores $\mathbf{v}_k$, o produto $U^{-1}AU$ é a matriz diagonal

$$U^{-1}AU = \Lambda = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix}.$$

e portanto $A = U\Lambda U^{-1}$.

e.g. Exemplos de matrizes diagonalizáveis. As matrizes

$$A = \begin{pmatrix} 2 & 0 \\ 0 & 3 \end{pmatrix} \quad \text{e} \quad B = \begin{pmatrix} 2 & 1 \\ 0 & 3 \end{pmatrix}$$

são diagonalizáveis, têm os mesmos valores próprios, 2 e 3, e até o mesmo polinómio caraterístico $(z-2)(z-3)$. Os vetores próprios de A formam uma base ortogonal, a base canónica. Por outro lado, os vetores próprios de B , que são proporcionais a $(1,0)$ e $(1,1)$, respetivamente, não formam uma base ortogonal (relativamente a estrutura euclidiana natural do plano). No entanto, $U^{-1}BU = A$ se a mudança de coordenadas é definida pela matriz

$$U = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$$

e.g. Uma matriz 2×2 com apenas uma reta de vetores próprios. A matriz

$$C = \begin{pmatrix} 2 & 1 \\ 0 & 2 \end{pmatrix}$$

admite um valor próprio, 2, e apenas uma reta de vetores próprios, a reta dos vetores proporcionais a $(1,0)$. De fato, a imagem de todo vetor $\mathbf{v} = (x,y)$ é $C\mathbf{v} = 2\mathbf{v} + (0,y)$, que não pode ser proporcional a $\mathbf{v}$ se $y \neq 0$. Em particular, não é diagonalizável.

e.g. O polinómio caraterístico não deteta a diagonalizabilidade. As matrizes

$$D = \begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix} \quad \text{e} \quad C = \begin{pmatrix} 2 & 1 \\ 0 & 2 \end{pmatrix}$$

têm o mesmo polinómio caraterístico $(z-2)^2$. No entanto, a primeira é diagonal, logo diagonalizável, mas a segunda não é diagonalizável.

e.g. Uma matriz com nenhum vetor próprio. Finalmente, a matriz

$$R = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$$

que representa uma rotação de um ângulo $\pi/2$, não admite vetores próprios. Isto é evidente geometricamente. Do ponto de vista algébrico, basta observar que a imagem de $\mathbf{v} = (x, y)$ é $R\mathbf{v} = (-y, x)$, e pode ser proporcional a $\mathbf{v}$ sse $y = -\lambda x$ e $x = \lambda y$, o que é impossível se $\mathbf{v}$ não é o vetor nulo.

e.g. Diagonalização da matriz do “gato de Arnold”. Considere a transformação linear $L: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ definida, na base canônica, pela matriz

$$A = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$$

ou seja, $L(x, y) = (2x + y, x + y)$. O polinômio característico é $P_A(z) = z^2 - 3z - 1$, e portanto os valores próprios são $\lambda_{\pm} = (3 \pm \sqrt{5})/2$. Vetores próprios correspondentes, que satisfazem $L(\mathbf{v}_{\pm}) = \lambda_{\pm}\mathbf{v}_{\pm}$, são, por exemplo,

$$\mathbf{v}_+ = (\varphi, 1) \quad \text{e} \quad \mathbf{v}_- = (1, -\varphi)$$

onde

$$\varphi = \frac{1 + \sqrt{5}}{2} \simeq 1.6180339887 \dots$$

é a “razão” dos gregos (a “divina proportione” de Pacioli e Leonardo da Vinci, a “razão de ouro” de Kepler, ...), a raiz positiva do polinômio $\varphi^2 - \varphi - 1$ (assim, se de um retângulo de lados φ e 1 cortamos um quadrado de lado unitário, ficamos com um retângulo de lados 1 e $\varphi - 1$ que é similar ao retângulo inicial, pois $\varphi/1 = 1/(\varphi - 1)$). Observem que $\lambda_+ > 1$ e $0 < \lambda_- < 1$, e portanto L estica os vetores da reta $\mathbb{R}\mathbf{v}_+$ e contrae os vetores da reta $\mathbb{R}\mathbf{v}_-$. Observem também que $\text{Det} A = \lambda_+ \lambda_- = 1$. Em particular, L preserva as áreas. A matriz A é invertível, e a sua inversa A^{-1} tem também entradas inteira, pois $\text{Det} A = 1$. Observem também que $\text{Tr} A = 3 = \lambda_+ + \lambda_-$. A matriz que representa L na base formada pelos vetores próprios $\mathbf{v}_+$ e $\mathbf{v}_-$ é a matriz diagonal

$$\Lambda = \begin{pmatrix} \lambda_+ & 0 \\ 0 & \lambda_- \end{pmatrix} = U^{-1}AU$$

onde U é matriz mudança de base, a matriz ortogonal

$$U = \frac{1}{\sqrt{1 + \varphi^2}} \begin{pmatrix} \varphi & 1 \\ 1 & -\varphi \end{pmatrix}$$

cujas colunas são as coordenadas dos vetores próprios normalizados de A na base canônica.

e.g. Diagonalização das translações na circunferência discreta. Consideramos o operador “translação” no espaço das funções complexas definidas no grupo finito $\mathbb{Z}/N\mathbb{Z}$, ou seja, o operador $T: \mathbb{C}^N \rightarrow \mathbb{C}^N$ definido por

$$T(x_0, x_1, \dots, x_{N-1}) = (x_1, x_2, \dots, x_{N-1}, x_0)$$

(observem que é conveniente fazer variar as coordenadas dos vetores entre 0 e $N - 1$). A sua N -ésima potência é a identidade, $T^N = I$. Portanto, se $\mathbf{v}$ é um vetor próprio com valor próprio λ , assim que $T\mathbf{v} = \lambda\mathbf{v}$, então $T^N\mathbf{v} = \lambda^N\mathbf{v} = \mathbf{v}$. Concluimos que os valores próprios de T são raízes N -ésimas da unidade, pois satisfazem $\lambda^N = 1$. No plano complexo existem exatamente N raízes N -ésimas da unidade, as potências

$$\lambda_0 = 1 \quad \lambda_1 = \zeta \quad \lambda_2 = \zeta^2 \quad \dots \quad \lambda_{N-1} = \zeta^{N-1}$$

de uma raiz primitiva, por exemplo $\zeta = e^{i2\pi/N}$. Um cálculo elementar mostra que um vetor próprio associado ao valor próprio $\lambda_n = \zeta^n$ é

$$\xi_n = (1, \zeta^n, \zeta^{2n}, \dots, \zeta^{(N-1)n})$$

Os vetores próprios ξ_n 's, com $n = 0, 1, 2, \dots, N-1$, formam uma base de $\mathbb{C}^N$, sendo associados a valores próprios distintos. A possibilidade de representar todo vetor de $\mathbb{C}^N$ como sobreposição dos ξ_n 's dá origem a “transformada de Fourier discreta” (entre outras coisas um dos instrumentos básico para o tratamento de sinais digitais), que convém tratar com a ajuda da estrutura euclidiana natural do espaço (logo no segundo semestre, na UC de “Complementos de Cálculo e de Álgebra Linear”).

A matriz

$$P = \begin{pmatrix} 0 & 0 & \dots & 0 & 1 \\ 1 & 0 & \dots & 0 & 0 \\ 0 & 1 & \ddots & \vdots & \vdots \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & 1 & 0 \end{pmatrix}$$

que representa o operador T na base canónica é chamada *matriz permutação cíclica*. As matrizes que representam os polinómios $f(T) = a_0 + a_1T + a_2T^2 + \dots + a_{N-1}T^{N-1}$ com coeficientes complexos a_k 's no operador T são conhecidas como *matrizes circulares*. Os seus valores próprios são $f(\lambda_n)$, com $\lambda_n = \zeta^n$. Por exemplo, $D_+ = T - I$ e $D_- = I - S$ podem ser consideradas “derivadas discretas” (“forward” e “backward”, respetivamente), e consequentemente $D_+D_- = T - 2I + S$ um “laplaciano discreto” ...

ex: Escreva explicitamente as matrizes $f(P)$ que representam os operadores $f(T)$, onde $f(z)$ é um polinómio, em particular as matrizes que representam $D_\pm$ e D_+D_- . Calcule o determinante das matrizes $f(P)$.

ex: As matrizes

$$\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad \text{e} \quad \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$$

são semelhantes?

ex: Diagonalize as seguintes matrizes, ou mostre que não é possível.

$$\begin{pmatrix} 2 & 3 \\ 3 & 4 \end{pmatrix} \quad \begin{pmatrix} 2 & 1 \\ 0 & 3 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix} \quad \begin{pmatrix} 3 & -2 \\ 2 & -2 \end{pmatrix}$$

$$\begin{pmatrix} 2 & 1 \\ -1 & 4 \end{pmatrix} \quad \begin{pmatrix} 1 & -1 \\ -2 & 2 \end{pmatrix} \quad \begin{pmatrix} 2 & 1 \\ -1 & 0 \end{pmatrix} \quad \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$$

ex: Seja B uma matriz real 3×3 com valores próprios 2, 3 e 4. Calcule o determinante e o traço da matriz B^{-1} .

Linear homogeneous recursive equations. A linear homogeneous recursive system is a law

$$X_{k+1} = AX_k$$

for some vector valued sequence $X_k \in \mathbb{R}^n$, given a square matrix $A \in \text{Mat}_{n \times n}(\mathbb{R})$. The solution is

$$X_k = A^k X_0$$

where $X_0 \in \mathbb{R}^n$ is the initial condition. The computation of the powers A^k of a square matrix A is simplified if we can diagonalize it. Indeed, if $A = U^{-1}\Lambda U$ with $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$, then its k -th power is simply $A^k = U^{-1}\Lambda^k U$, where $\Lambda^k = \text{diag}(\lambda_1^k, \dots, \lambda_n^k)$.

Linear homogeneous recursive equations may arise from finite difference equations as in the following example.

ex: Consider the “Fibonacci sequence”, defined by the recursive equation

$$f_{k+2} = f_{k+1} + f_k$$

with initial conditions $f_0 = f_1 = 1$. Explicitly, the sequence reads

$$1, 1, 2, 3, 5, 8, 13, 21, 34, 55, 89, 144, 233, 377, 610, 987, 1597, 2584, \dots$$

These numbers soon become astronomically large. For example, $f_{120} \simeq 8.67 \times 10^{24}$, larger than the Avogadro number! If we define the vector valued sequence $F_k = (f_{k+1}, f_k)^\top \in \mathbb{R}^2$, the recursive equation for the f_k 's is equivalent to

$$F_{k+1} = AF_k \quad \text{with} \quad A = \begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix}$$

and initial condition $F_0 = (1, 1)^\top$. The solution is $F_k = A^k F_0$. Compute the eigenvalues of A , diagonalize it, compute its k -th power, and finally find a formula for the k -th Fibonacci number f_k (known as “Binet’s formula”).

Stochastic matrices and Markov chains. An important class of square matrices is the class of *stochastic matrices*. They are those square matrices $P = (p_{ij}) \in \text{Mat}_{n \times n}(\mathbb{R})$ that define endomorphisms of $\mathbb{R}^n$ which preserve the unit simplex Δ^{n-1} (defined in (3.2)). We think at the simplex as the space of probability measures on a space made of n atoms, column vectors $\mathbf{p} = (p_1, p_2, \dots, p_n)^\top$ with non-negative entries such that $\sum_{k=1}^n p_k = 1$, also called “stochastic vectors” (warning: probabilists use row probability vectors instead, hence act with matrices on the right!). The condition $P\mathbf{p} \in \Delta^{n-1}$ for all $\mathbf{p} \in \Delta^{n-1}$ requires that the matrix has non-negative entries, i.e. $p_{ij} \geq 0$, and that the sum of each column is one, i.e.

$$\sum_{i=1}^n p_{ij} = 1 \tag{13.4}$$

for all $j = 1, 2, \dots, n$. Stochastic matrices were introduced by the Russian mathematician Andrey Markov in his theory of “Markov chains”. In its simplest version, we may think that a system can be observed in one of the states $1, 2, \dots, n$, and that in each step of time (thought as discrete) may undergo a transition from the state i to the state j with “probability” $\text{Prob}(i \rightarrow j) = p_{ji}$ (hence the condition on the columns of P is simply the “law of total probability”). A simple reasoning shows that if a system at time 0 is described by a probability vector $\mathbf{p}$, i.e. it is observed in the state k with probability p_k , then its state at time n will be described by the probability vector

$$\mathbf{p}(n) = P^n \mathbf{p}$$

Therefore, the far future of the system is governed by the behaviour of large powers of P . An eigenvector $\mathbf{v}$ with eigenvalue 1 which is also a probability vector is a “stationary/equilibrium state/distribution” for the Markov chain, since $P\mathbf{v} = \mathbf{v}$ implies that if the system starts at state $\mathbf{v}$ then it will remain in such a state for all future times. It always exists (see exercises), but needs not be unique. It is also the case that 1 is the largest absolute value of the eigenvalues of a stochastic matrix (see exercises).

Today, Markov chains and processes, together with the ergodic theorem (dealing with the possible asymptotics of the $P^n \mathbf{p}$ as $n \rightarrow \infty$), are fundamental tools in probability and dynamical systems, as well as important tools in many areas of applied mathematics (typical examples are Monte Carlo methods to find approximate solutions to complex problems, and a striking example is the Google “PageRank algorithm” by Brin and Page²⁸).

ex: Verifique que produtos, e portanto potências, de matrizes estocásticas são também matrizes estocásticas.

²⁸S. Brin and L. Page, The anatomy of a large-scale hypertextual Web search engine, *Computer Networks and ISDN Systems* **30** (1998), 107-117.

ex: Se P é uma matriz estocástica, então a própria (13.4) diz que o vetor coluna $(1, 1, \dots, 1)^\top$ é um vetor próprio da matriz transposta $P^\top$, com valor próprio igual a 1. Consequentemente, também P admite um valor próprio igual a 1, pois os polinômios característicos de P e $P^\top$ são iguais (embora não seja evidente como calcular um vetor próprio).

ex: Mostre que uma matriz estocástica P não pode ter valores próprios λ com módulo $|\lambda| > 1$ (se $P\mathbf{v} = \lambda\mathbf{v}$ com $|\lambda| > 1$, então $P^n\mathbf{v} = \lambda^n\mathbf{v}$, e isto implica que alguma entrada de P^n , quando n é grande, é superior a um ...).

ex: Exemplos simples de matrizes estocásticas 2×2 são

$$\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix} \quad \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \begin{pmatrix} 1/2 & 1/2 \\ 1/2 & 1/2 \end{pmatrix}$$

Calcule os seus valores e vetores próprios.

ex: A matriz estocástica genérica 2×2 é

$$A = \begin{pmatrix} \varepsilon & \eta \\ 1 - \varepsilon & 1 - \eta \end{pmatrix}$$

com $0 \leq \varepsilon, \eta \leq 1$. Mostre que se $0 < \varepsilon < 1$ e $0 < \eta < 1$ (ou seja, se todas as entradas da matriz A são estritamente positivas) então o segundo valor próprio tem módulo $|\lambda| < 1$ (observe que a soma dos valores próprios é igual ao traço da matriz ...).

Transition matrices and Markov graphs. Another interesting class of square matrices, motivated by the theory of Markov chains but not directly interpreted as “transformations”, is that of *transition matrices*. They are square matrices $T = (t_{ij}) \in \text{Mat}_{n \times n}(\mathbb{R})$, with entries that only take two possible values, $t_{ij} = 0$ or 1. Those matrices clearly preserve the “cone of non-negative vectors”, those vectors having coordinates $x_k \geq 0$. Any transition matrix defines a directed graph $\mathcal{G}_T$, whose vertices are the possible states $1, 2, \dots, n$ of a system, and with a directed arrow joining the state i to the state j (hence a possible transition) iff $t_{ji} = 1$. Such “Markov graphs” have a natural metrics, and large powers of T determine the asymptotic of the lengths of its closed geodesics.

The Perron-Frobenius theorem, dealing with the relation between large powers T^n and the largest eigenvalue of T , has also important applications in modern technology.

14 Estrutura dos operadores

ref: [Ax15] 8.A-D 9.A

14.1 Espaços próprios generalizados

não lecionado

Operadores “genéricos” num espaço vetorial complexo de dimensão finita são diagonalizáveis. De fato, a menos de pequenas perturbações nas entradas da matriz, o polinómio característico admite raízes distintas. Por outro lado, é natural estar interessados em operadores específicos, rígidos. É necessário então dispor de resultados que descrevam a sua estrutura, na sua forma mais simples, ou seja, compreensível.

Decomposição de operadores. Se $P : \mathbf{V} \rightarrow \mathbf{V}$ é uma projeção, então o espaço total é uma soma direta $\mathbf{V} = \text{Ker}(P) \oplus \text{Im}(P)$ do seu núcleo e da sua imagem, e o operador é uma soma direta $P = 0 \oplus 1$ do operador nulo e do operador identidade. As coisas não são assim simples para operadores genéricos.

Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um endomorfismo de um espaço vetorial $\mathbf{V}$, e consideramos as suas potências L^k , com $k \geq 0$. É claro que $\text{Ker}(L^k) \subset \text{Ker}(L^{k+1})$, pois, se $L^k \mathbf{v} = 0$, então também $L^{k+1} \mathbf{v} = L(L^k \mathbf{v}) = 0$. Assim, temos umas inclusões de subespaços invariantes

$$0 = \text{Ker}(L^0) \subset \text{Ker}(L) \subset \text{Ker}(L^2) \subset \dots \subset \text{Ker}(L^k) \subset \text{Ker}(L^{k+1}) \subset \dots$$

Também temos inclusões opostas de subespaços invariantes

$$\dots \subset \text{Im}(L^{k+1}) \subset \text{Im}(L^k) \subset \dots \subset \text{Im}(L^2) \subset \text{Im}(L) \subset \text{Im}(L^0) = \mathbf{V}$$

Se o espaço $\mathbf{V}$ tem dimensão finita, as inclusões não podem ser todas estritas, pois

$$0 = \dim \text{Ker}(L^0) \leq \dim \text{Ker}(L) \leq \dim \text{Ker}(L^2) \leq \dots \leq \dim \mathbf{V} = n$$

Logo, existe um inteiro minimal $m \geq 0$ tal que $\text{Ker} L^m = \text{Ker} L^{m+1}$, e, consequentemente, $\text{Ker}(L^m) = \text{Ker}(L^{m+k})$ para todo $k \geq 1$ (exercício). Se n é a dimensão de $\mathbf{V}$, então necessariamente acontece que $m \leq n$.

Teorema 14.1. Se $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador definido no espaço vetorial $\mathbf{V}$ de dimensão finita $\dim \mathbf{V} = n$. Então o espaço é uma soma direta dos subespaços invariantes

$$\mathbf{V} = \text{Im}(L^n) \oplus \text{Ker}(L^n)$$

e o operador é uma soma direta $L = R \oplus N$ de um operador invertível $R = L|_{\text{Im}(L^n)}$ e um operador nilpotente $N = L|_{\text{Ker}(L^n)}$.

Demonstração. Pelo teorema 8.7, os subespaços $\text{Ker}(L^n)$ e $\text{Im}(L^n)$ têm dimensões complementares. Falta portanto provar que têm interseção vazia. Um vetor $\mathbf{v}$ nesta interseção é tal que $L^m \mathbf{v} = 0$ e existe $\mathbf{u} \in \mathbf{V}$ tal que $L^n \mathbf{u} = \mathbf{v}$. Ao aplicar L^n na segunda equação e usando a primeira observamos que $L^{2n} \mathbf{u} = L^n \mathbf{v} = 0$, ou seja, que $\mathbf{u} \in \text{Ker}(L^{2n})$. Mas $\text{Ker}(L^{2n}) = \text{Ker}(L^n)$, e consequentemente $\mathbf{v} = L^n \mathbf{u} = 0$.

Os subespaços $\text{Im}(L^n)$ e $\text{Ker}(L^n)$ são invariantes, assim que L é uma soma direta das suas restrições a estes dois subespaços. É tautológico que L é nilpotente em $\text{Ker}(L^n)$. Por outro lado, seja $\mathbf{v} \in \text{Im}(L^n)$, assim que $\mathbf{v} = L^n \mathbf{w}$ para algum $\mathbf{w} \in \mathbf{V}$. Se $L \mathbf{v} = 0$ também $0 = L(L^n \mathbf{w}) = L^{n+1} \mathbf{w}$, assim que $\mathbf{w} \in \text{Ker}(L^{n+1})$. Mas $\text{Ker}(L^{n+1}) = \text{Ker}(L^n)$, e portanto $\mathbf{v} = L^n \mathbf{w} = 0$. Sendo injetiva, a restrição de L a $\text{Im}(L^n)$ é invertível, pois estamos em dimensão finita. $\square$

ex: Mostre que se $\text{Ker}(L^m) = \text{Ker}(L^{m+1})$ então $\text{Ker}(L^m) = \text{Ker}(L^{m+k})$ para todo $k \geq 1$.

ex: Mostre que $\text{Im}(L^{k+1}) \subset \text{Im}(L^k)$, e que existe um inteiro minimal $m \geq 1$ tal que $\text{Im}(L^{m+1}) = \text{Im}(L^m)$ e, consequentemente, $\text{Im}(L^{m+k}) = \text{Im}(L^m)$ para todo $k \geq 1$.

ex: Dê um exemplo de um operador definido num espaço vetorial (de dimensão necessariamente infinita) tal que $\text{Ker}(L^k) \neq \text{Ker}(L^{k+1})$ para todo $k \geq 0$.

Espaços próprios generalizados. Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador linear definido num espaço linear complexo de dimensão finita $\dim \mathbf{V} = n$. Dado um escalar $\lambda \in \mathbb{C}$, denotamos por L_λ o operador $\lambda - L$. Se o núcleo de L_λ não é trivial, então λ é um valor próprio de L , e $\mathbf{V}_\lambda = \text{Ker}(L_\lambda)$ é o espaço próprio associado. Um vetor não nulo $\mathbf{v} \in \mathbf{V}$ é dito *vetor próprio generalizado* associado ao valor próprio λ se está no núcleo de alguma potência de L_λ , ou seja, se $L_\lambda^m \mathbf{v} = 0$ para algum inteiro minimal $m \geq 1$. Pelas considerações anteriores necessariamente $m \leq n$ e portanto vetores próprios generalizados estão no núcleo de L_λ^n .

O *espaço próprio generalizado* associado ao valor próprio λ é

$$G_\lambda := \text{Ker}(L_\lambda^n) = \{\mathbf{v} \in \mathbf{V} \text{ t.q. } (\lambda - L)^n \mathbf{v} = 0\}$$

Como L comuta com os L_λ e as suas potências, é claro que os espaços próprios generalizados são subespaços L -invariantes, que contêm os espaços próprios, $V_\lambda \subset G_\lambda$. Também útil é a seguinte observação. Se $\lambda \neq \mu$, então G_λ é um subespaço L -invariante de $\text{Im}(\lambda - L)$, e portanto, pelo teorema 14.1,

Lemma 14.2. *Se $\lambda \neq \mu$, então $\lambda - L$ é uma bijeção de G_μ .*

O teorema 13.1 generaliza assim.

Teorema 14.3. *Vetores próprios generalizados associados a valores próprios distintos são linearmente independentes.*

Demonstração. A prova é também uma variação da prova do teorema 13.1. Seja

$$c_1 \mathbf{v}_1 + c_2 \mathbf{v}_2 + \cdots + c_m \mathbf{v}_m = 0$$

uma combinação linear nula dos vetores próprios generalizados $\mathbf{v}_k \in G_{\lambda_k}$ associados a valores próprios diferentes λ_k . Se aplicamos o operador $(\lambda_2 - L)^n (\lambda_3 - L)^n \cdots (\lambda_m - L)^n$ aos dois membros desta identidade, temos

$$c_1 (\lambda_2 - L)^n (\lambda_3 - L)^n \cdots (\lambda_m - L)^n \mathbf{v}_1 = 0$$

porque cada $\mathbf{v}_k$ com $k \geq 2$ está no núcleo de um dos fatores do operador, que comutam. Pelo lemma 14.2, o produto $(\lambda_2 - L)^n (\lambda_3 - L)^n \cdots (\lambda_m - L)^n$ é bijetivo em G_{λ_1} (pois uma composição de bijeções é uma bijeção), e portanto $c_1 = 0$. Da mesma forma, se depois aplicamos o operador $(\lambda_1 - L)^n (\lambda_3 - L)^n \cdots (\lambda_m - L)^n$, obtemos $c_2 = 0$. Continuando assim, obtemos $c_k = 0$ para todos os k . $\square$

O seguinte resultado fundamental mostra que, embora os vetores próprios não sejam, em geral, suficientes para gerar o espaço total, os vetores próprios generalizados sim.

Teorema 14.4. *Sejam $\lambda_1, \lambda_2, \dots, \lambda_m$ os valores próprios (distintos) do operador $L : \mathbf{V} \rightarrow \mathbf{V}$, definido no espaço vetorial complexo de dimensão finita $\mathbf{V}$. Então o espaço total é uma soma direta*

$$\mathbf{V} = G_{\lambda_1} \oplus G_{\lambda_2} \oplus \cdots \oplus G_{\lambda_m}$$

de espaços próprios generalizados, e o operador é uma soma direta das suas restrições aos espaços próprios generalizados.

Demonstração. A prova é por indução sobre a dimensão. O teorema é trivial em dimensão 1. Assumimos o enunciado verdadeiro em dimensão $\leq n$, e consideramos um operador L definido num espaço vetorial complexo de dimensão $n + 1$. Pelo teorema 13.2, L admite um valor próprio λ . Pelo teorema 14.1 o espaço é uma soma direta dos subespaços invariantes

$$\mathbf{V} = \text{Ker}(L_\lambda^{n+1}) \oplus \text{Im}(L_\lambda^{n+1})$$

Como λ é um valor próprio, $\dim \text{Ker}(L_\lambda^{n+1}) \geq 1$ e consequentemente $\dim \text{Im}(L_\lambda^{n+1}) \leq n$. Mas $\text{Im}(L_\lambda^{n+1})$ é ou vazio ou, pela hipótese indutiva, uma soma direta de espaços próprios generalizados. $\square$

A dimensão $n_k = \dim G_{\lambda_k}$ é chamada *multiplicidade* do valor próprio λ_k (é um fato que coincide com a que chamamos antes “multiplicidade algébrica”, mas ainda não sabemos), e é superior ou igual à multiplicidade geométrica $\dim V_{\lambda_k}$. O teorema 14.4 então implica que a soma das multiplicidades é igual a $n_1 + n_2 + \dots + n_m = n$, a dimensão do espaço.

Um corolário interessante é que se $\lambda = 0$ é o único valor próprio do operador L definido num espaço vetorial complexo de dimensão finita, então L é nilpotente, pois neste caso o espaço total é igual ao espaço próprio generalizado $G_0 = \text{Ker}(L^n)$.

ex: Mostre que existem operadores definidos em espaços vetoriais reais de dimensão finita com único valor próprio 0 que não são nilpotentes (pense nas rotações ...).

14.2 Decomposição de Jordan-Chevalley

Estrutura dos operadores complexos. De acordo com o teorema 14.4, a compreensão do operador L passa pela descrição das suas restrições aos espaços próprios generalizados. Mas a restrição de $\lambda - L$ ao espaço próprio generalizado G_λ é, por definição, um operador nilpotente. Consequentemente, cada restrição é uma soma $L|_{G_\lambda} = \lambda + N$ de um múltiplo λ da identidade e um operador nilpotente N .

Teorema 14.5 (decomposição de Jordan-Chevalley). *Sejam $\lambda_1, \lambda_2, \dots, \lambda_m$ os valores próprios (distintos) do operador $L : \mathbf{V} \rightarrow \mathbf{V}$, definido no espaço vetorial complexo de dimensão finita $\mathbf{V}$. Então o espaço total é uma soma direta*

$$\mathbf{V} = G_{\lambda_1} \oplus G_{\lambda_2} \oplus \dots \oplus G_{\lambda_m}$$

de espaços próprios generalizados, e o operador é uma soma direta

$$L = (\lambda_1 + N_1) \oplus (\lambda_2 + N_2) \oplus \dots \oplus (\lambda_m + N_m)$$

de operadores $\lambda_k + N_k$ que são somas de múltiplos da identidade e operadores nilpotentes N_k . Em particular, o operador é uma soma $L = \Lambda + N$ de um operador semi-simples, ou seja, diagonalizável, $\Lambda = \lambda_1 \oplus \lambda_2 \oplus \dots \oplus \lambda_n$ e um operador nilpotente $N = N_1 \oplus N_2 \oplus \dots \oplus N_m$ que comutam.

Para compreender a estrutura das matrizes que definem o operador, é suficiente portanto compreender a estrutura das matrizes que definem os operadores nilpotentes. Uma primeira caracterização é particularmente simples.

Teorema 14.6. *Seja N um operador nilpotente definido num espaço vetorial complexo de dimensão finita. Então existe uma base na qual o operador é definido por uma matriz diagonal superior com apenas zeros na diagonal, ou seja, da forma*

$$\begin{pmatrix} 0 & * & \dots & * \\ 0 & 0 & \dots & * \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix}$$

Demonstração. Começamos por escolher uma base de $\text{Ker}(N)$. Depois completamos o sistema até ter uma base de $\text{Ker}(N^2)$. Depois completamos o sistema até ter uma base de $\text{Ker}(N^3)$, ... e assim a seguir. Na base resultante, a matriz de N é claramente diagonal superior com zeros na diagonal, porque $N(\text{Ker}(N^{k+1})) \subset \text{Ker}(N^k)$. $\square$

Finalmente, juntando a decomposição de Jordan-Chevalley 14.5 e o teorema 14.6, a matriz que define um operador L numa base oportuna é uma matriz diagonal em blocos

$$A = \begin{pmatrix} A_1 & 0 & \dots & 0 \\ 0 & A_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & A_m \end{pmatrix}$$

com blocos que correspondem aos valores próprios λ_k 's e são matrizes diagonais superiores da forma (14.1)

$$A_k = \begin{pmatrix} \lambda_k & * & \dots & * \\ 0 & \lambda_k & \dots & * \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_k \end{pmatrix} \quad (14.1)$$

e dimensão igual a multiplicidade n_k .

Consequentemente, o polinómio caraterístico de L fatoriza no produto

$$c_L(z) = (z - \lambda_1)^{n_1} (z - \lambda_2)^{n_2} \dots (z - \lambda_m)^{n_m} \quad (14.2)$$

Em particular, a multiplicidade n_k do valor próprio λ_k (ou seja, a dimensão de G_{λ_k}) coincide com a multiplicidade algébrica de λ_k (ou seja, a sua multiplicidade enquanto raiz do polinómio caraterístico). Como sugerido por Sheldon Axler²⁹, esta fórmula (14.2) pode ser usada para dar uma definição geométrica do polinómio caraterístico, independente da noção de determinante.

Numa outra base, a matriz de L será UAU^{-1} , sendo U uma matriz invertível (a mudança de coordenadas). Em particular, a matriz que define o operador numa base arbitrária pode ser representada como uma soma

$$\Lambda + N$$

de uma matriz diagonalizável $\Lambda = U(\lambda_1 \oplus \lambda_2 \oplus \dots \oplus \lambda_m)U^{-1}$ e de uma matriz nilpotente $N = U(N_1 \oplus N_2 \oplus \dots \oplus N_m)U^{-1}$ que comutam.

Os blocos diagonais superiores A_k são invertíveis sse $\lambda_k \neq 0$. Se juntamos todos os espaços próprios com valor próprio $\lambda \neq 0$ recuperamos a decomposição do teorema 14.1.

O teorema de estrutura 14.5 é suficiente para a grande parte das aplicações interessantes, como a seguinte. No entanto, por exemplo para compreender os exponenciais de operadores reais, é útil simplificar ainda mais a matriz de um operador nilpotente, e chegar à chamada “forma normal de Jordan”.

ex: Mostre que o traço de uma matriz nilpotente é nulo, e o traço de uma matriz unipotente é igual a dimensão da matriz.

Polinómio minimal e teorema de Cayley-Hamilton. Como $L - \lambda_k$ é nilpotente em G_{λ_k} , existe um inteiro minimal $q_k \leq n_k$ tal que $(L - \lambda_k)^{q_k} = 0$ em G_{λ_k} . Isto sugere definir o *polinómio minimal* de L (ou das matrizes que o representam numas bases arbitrárias) como

$$m_L(z) = (z - \lambda_1)^{q_1} (z - \lambda_2)^{q_2} \dots (z - \lambda_m)^{q_m} \quad (14.3)$$

É claro que o polinómio minimal $m_L(z)$ divide o polinómio caraterístico $c_L(z)$.

Se $\mathbf{v} \in G_{\lambda_k}$, então $(L - \lambda_k)^{q_k} \mathbf{v} = 0$. Pelo teorema 14.4, cada vetor é uma combinação linear de vetores próprios generalizados. Como os $(L - \lambda_k)$'s comutam, temos portanto

²⁹S. Axler, Down With Determinants!, *American Mathematical Monthly* **102** (1995), 139-154.

Teorema 14.7 (Cayley-Hamilton). *Todo operador L num espaço vetorial complexo de dimensão finita anula o seu polinómio minimal, e consequentemente também o seu polinómio caraterístico, ou seja, satisfaz as equações algébricas*

$$m_L(L) = 0 \quad e \quad c_L(L) = 0$$

O nome “minimal” é justificado pelas seguintes considerações. O espaço linear $\text{Mat}_{n \times n}(\mathbb{C})$ das matrizes $n \times n$ tem dimensão finita, igual a n^2 , assim que no máximo n^2 potências de uma matriz dada A podem ser linearmente independentes. Consequentemente, existe pelo menos um polinómio

$$f(z) = a_m z^m + a_{m-1} z^{m-1} + \cdots + a_1 z + a_0$$

de grau $m \leq n^2$ e com coeficientes complexos a_k 's que “anula” a matriz A , ou seja, tal que $f(A) = 0$, no sentido em que

$$a_m A^m + a_{m-1} A^{m-1} + \cdots + a_1 A + a_0 I = 0$$

onde o segundo membro é a matriz nula. Naturalmente, anular a matriz é equivalente a anular o operador L que a matriz define numa base arbitrária, pois $f(A) = U^{-1} f(B) U$ se $A = U B U^{-1}$ com U invertível. É claro também que se $a_m \neq 0$ (ou seja, se o grau de $f(z)$ é igual a m) podemos dividir por a_m e considerar o polinómio “mónico” (cujo termo de grau maior tem coeficiente igual a um) $p(z) = z^m + b_{m-1} z^{m-1} + \cdots + b_1 z + b_0$, com coeficientes $b_k = a_k/a_m$, que também anula a matriz A .

Teorema 14.8. *O polinómio minimal divide todo polinómio que anula L , e portanto é o polinómio mónico de grau menor que anula o operador L .*

Demonstração. Pelo teorema 14.4, todo vetor é uma soma direta de vetores próprios generalizados, logo um operador nulo deve anular todos os $\mathbf{v}_k \in G_{\lambda_k}$. De acordo com o lemma 14.2, se $\mu \neq \lambda_k$ então $(L - \mu)$ é uma bijeção de G_{λ_k} . Consequentemente, se $f(z) = (z - \mu_1)(z - \mu_2) \cdots (z - \mu_p)$ é um polinómio mónico tal que $f(L)\mathbf{v}_k = 0$, então $f(z)$ deve conter o fator $(z - \lambda_k)^{q_k}$. $\square$

O polinómio minimal é único porque se dois polinómios mónicos do mesmo grau m anulam a matriz, então a diferença é um polinómio de grau $< m$ que ainda anula a matriz, logo proporcional a um polinómio mónico de grau $< m$ que anula a matriz.

e.g. Polinómios minimais de múltiplos da identidade. O polinómio minimal da matriz identidade I , em dimensão arbitrária, é $p_I(z) = z - 1$. Mais em geral, o polinómio minimal do múltiplo λI da matriz identidade é o polinómio de grau um

$$m_{\lambda I}(z) = z - \lambda$$

e.g. Polinómios minimais de projeções. A matriz P que representa uma projeção satisfaz $P^2 = P$. O seu polinómio minimal é

$$m_P(z) = z(z - 1)$$

(porque é claro que um polinómio de grau um não pode anular uma projeção não trivial).

e.g. Polinômios minimais de matrizes diagonais. O polinômio minimal de uma matriz diagonal

$$\Lambda = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n) := \begin{pmatrix} \lambda_1 & 0 & 0 & \dots \\ 0 & \lambda_2 & 0 & \dots \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & \lambda_n \end{pmatrix}$$

com todos os λ_k 's diferentes, ou seja, tais que $\lambda_i \neq \lambda_j$ se $i \neq j$, é o polinômio de grau n

$$m_\Lambda(z) = (z - \lambda_1)(z - \lambda_2) \dots (z - \lambda_n)$$

e portanto coincide com o polinômio caraterístico $c_\Lambda(z)$. Em geral, é igual ao produto dos valores próprios.

14.3 Blocos de Jordan

Blocos de Jordan. A estrutura dos operadores nilpotentes pode ser ainda simplificada (ou seja, a matriz diagonal superior pode ter ainda mais zeros!). Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador linear definido num espaço linear complexo de dimensão finita, e seja λ um seu valor próprio. Um vetor próprio generalizado é um vetor não nulo $\mathbf{v} \in \mathbf{V}$ tal que $L_\lambda^n \mathbf{v} = 0$ para algum inteiro minimal $n \geq 1$. Este inteiro n é dito *período* de $\mathbf{v}$, e o vetor $\mathbf{v}$ é também dito L_λ -*cíclico* (a órbita de $\mathbf{v}$ pela aplicação L_λ é formada por n vetores não nulos distintos).

Se o período é $n = 1$, então $\mathbf{v}$ é um vetor próprio de L . Em geral, os n vetores

$$\mathbf{v}_1 = L_\lambda^{n-1} \mathbf{v} \quad \mathbf{v}_2 = L_\lambda^{n-2} \mathbf{v} \quad \dots \quad \mathbf{v}_{n-1} = L_\lambda \mathbf{v}_n \quad \mathbf{v}_n = \mathbf{v} \quad (14.4)$$

são vetores próprios generalizados, pois $L_\lambda^k \mathbf{v}_k = L_\lambda^k L_\lambda^{n-k} \mathbf{v} = L_\lambda^n \mathbf{v} = 0$, e $\mathbf{v}_1$ é um vetor próprio com valor próprio λ . O espaço gerado pelos $\mathbf{v}_k$'s é L -invariante porque

$$L\mathbf{v}_k = L(L_\lambda^{n-k} \mathbf{v}) = L_\lambda^{n-k+1} \mathbf{v} + \lambda L_\lambda^{n-k} \mathbf{v} = \mathbf{v}_{k-1} + \lambda \mathbf{v}_k$$

onde, naturalmente, $\mathbf{v}_0 = (\lambda - L)^n \mathbf{v} = 0$.

Teorema 14.9. Se $\mathbf{v}$ é um vetor L_λ -cíclico de período n , então os n vetores (14.4) são linearmente independentes e geram um subespaço L -invariante de vetores próprios generalizados.

Demonstração. Seja $a_1 \mathbf{v}_1 + a_2 \mathbf{v}_2 + \dots + a_n \mathbf{v}_n = 0$. Ao aplicar L_λ^{n-1} aos dois termos da igualdade, temos que $a_n \mathbf{v}_1 = 0$, logo $a_n = 0$. Ao aplicar L_λ^{n-2} aos dois termos da igualdade que sobra $a_1 \mathbf{v}_1 + a_2 \mathbf{v}_2 + \dots + a_{n-1} \mathbf{v}_{n-1} = 0$, temos que $a_{n-1} \mathbf{v}_1 = 0$, logo também $a_{n-1} = 0$. E assim a seguir. □

Os vetores (14.4) geram um subespaço invariante $\text{Span}(\mathcal{O}_L^+(\mathbf{v})) \subset \text{Ker}(L_\lambda^n)$ de dimensão n , chamado *espaço cíclico* gerado por $\mathbf{v}$. O cálculo acima mostra que

$$L\mathbf{v}_1 = \lambda \mathbf{v}_1 \quad \text{e} \quad L\mathbf{v}_k = \lambda \mathbf{v}_k + \mathbf{v}_{k-1} \quad \text{se } 2 \leq k \leq n$$

Então, a matriz que representa a restrição de L a este subespaço relativamente à base $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$ é

$$J_\lambda = \begin{pmatrix} \lambda & 1 & & & \\ & \lambda & 1 & & \\ & & \ddots & \ddots & \\ & & & \lambda & 1 \\ & & & & \lambda \end{pmatrix} \quad (14.5)$$

A matriz (14.5) é dita *bloco de Jordan* de dimensão n , e a base (14.4) é dita *base de Jordan*, ou *cadeia de Jordan* de comprimento n .

Observe que um bloco de Jordan de dimensão n é da forma

$$J_\lambda = \lambda I + N$$

onde N é a matriz nilpotente (triangular superior)

$$N = \begin{pmatrix} 0 & 1 & & & \\ & 0 & 1 & & \\ & & \ddots & \ddots & \\ & & & 0 & 1 \\ & & & & 0 \end{pmatrix} \quad (14.6)$$

que verifica $NE_k = E_{k-1}$, se E_k denotam os vetores coluna da base canônica de $\mathbb{C}^n$, e $N^n = 0$. O operador definido pela matriz N , que envia $(x_1, x_2, x_3, \dots, x_n) \mapsto (x_2, x_3, \dots, x_n, 0)$, é também chamado *deslocamento esquerdo*.

O polinômio característico de um bloco de Jordan (14.5) de comprimento n é $c_{J_\lambda}(z) = (z - \lambda)^n$, e portanto a multiplicidade algébrica do valor próprio λ é n . O polinômio minimal é também $m_{J_\lambda}(z) = (z - \lambda)^n$ (comparado com o polinômio minimal da matriz λI , que é $m_{\lambda I}(z) = (z - \lambda)$).

Quase-polinômios e derivada. O modelo de um bloco de Jordan é o operador derivação, definido por $(Df)(t) := f'(t)$, no espaço linear $\text{QPol}_n^\lambda(\mathbb{R})$ dos quase-polinômios $f(t) = p(t)e^{\lambda t}$ de grau $\deg(p) < n$. Se $\lambda = 0$, então $D^n = 0$, pois a n -ésima derivada de um polinômio de grau $< n$ é nula. Em geral, é imediato verificar que $(\lambda - D)^n = 0$. O vetor próprio é $f(t) = e^{\lambda t}$, com valor próprio λ . A base de Jordan é

$$e^{\lambda t} \quad t e^{\lambda t} \quad \frac{1}{2} t^2 e^{\lambda t} \quad \dots \quad \frac{1}{(n-1)!} t^{n-1} e^{\lambda t}$$

Nesta base, o operador D é representado pela matriz (14.5).

ex: Considere um bloco de Jordan (14.5) de dimensão n . Calcule as suas potências J_λ^k , com $k = 0, 1, 2, \dots$, e depois $f(J_\lambda)$ quando f é o polinômio $f(z) = a_0 + a_1 z + a_2 z^2 + \dots + a_n z^n$.

Forma normal de Jordan. Um espaço próprio generalizado G_λ não é sempre, naturalmente, um espaço cíclico para L_λ . No entanto, é possível provar que é uma soma direta de espaços cíclicos. Este é o conteúdo do teorema da forma canônica de Jordam, que é essencialmente um teorema sobre operadores nilpotentes em dimensão finita.

Seja N um operador nilpotente definido num espaço vetorial de dimensão finita $\mathbf{V}$ (podem pensar que N é restrição de L_λ ao espaço próprio generalizado G_λ), e seja m o seu índice de nilpotência, ou seja, o menor $k \geq 1$ tal que $N^k = 0$. Para simplificar as notações, denotamos por $K_k := \text{Ker}(N^k)$ se $k = 1, 2, \dots, m$, os núcleos da potências de N . É claro que são subespaços N -invariantes, e que $N(K_k) \subset K_{k-1}$. A primeira observação é que as inclusões

$$0 \subset K_1 \subset K_2 \subset \dots \subset K_{m-1} \subset K_m = \mathbf{V}$$

são estritas.

O teorema 14.9 (que é de fato um teorema sobre um genérico operador nilpotente $N = L_\lambda$) afirma que se $\mathbf{v} \in K_k \setminus K_{k-1}$, então a sua k -órbita

$$\mathbf{v}_k := \mathbf{v} \quad \mathbf{v}_{k-1} := N\mathbf{v} \quad \mathbf{v}_{k-2} := N^2\mathbf{v} \quad \dots \quad \mathbf{v}_1 := N^{k-1}\mathbf{v}$$

é um conjunto livre, chamado “cadeia de Jordan”, que gera um subespaço N -invariante, e nesta base ordenada $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k$ a restrição do operador N a este subespaço é definida por um bloco de Jordan nilpotente (14.6) de dimensão k . É claro também que um número finito de órbitas gera o espaço total, pois a dimensão é finita. Concatenações de órbitas diferentes podem, no entanto, não constituir um conjunto livre.

A chave do teorema de Jordan está nas seguintes observações elementares, que convém registrar como

Lemma 14.10. *Seja F é um subespaço tal que $F \cap K_k = 0$, e seja $1 \leq i \leq k-1$. Então*

$$N^i(F) \cap K_{k-i} = 0$$

e as restrições das potências N^i a F são injetivas, e portanto definem bijeções

$$N^i|_F : F \rightarrow N^i(F)$$

Demonstração. Se dois vetores $\mathbf{x}, \mathbf{y} \in F$ têm imagens $N^i\mathbf{x} = N^i\mathbf{y}$ então a diferença $\mathbf{x} - \mathbf{y} \in \text{Ker}(N^i) \subset K_k$, mas o único vetor da interseção $F \cap K_k$ é o vetor nulo. Por outro lado, se $\mathbf{x} \in F$ e $\mathbf{y} = N^i(\mathbf{x}) \in K_{k-i}$, então $0 = N^{k-i}\mathbf{y} = N^{k-i}N^i\mathbf{x} = N^k\mathbf{x}$, logo também $\mathbf{x} \in K_k$, mas $F \cap K_k = 0$. $\square$

Teorema 14.11 (forma normal de Jordan de operadores nilpotentes). *Seja N um operador nilpotente definido num espaço vetorial de dimensão finita $\mathbf{V}$. Então $\mathbf{V}$ admite uma base que é uma concatenação de órbitas*

$$\mathbf{v}_1, N\mathbf{v}_1, N^2\mathbf{v}_1, \dots, N^{k_1}\mathbf{v}_1, \mathbf{v}_2, N\mathbf{v}_2, N^2\mathbf{v}_2, \dots, N^{k_2}\mathbf{v}_2, \dots$$

Iso significa que existe uma base na qual o operador nilpotente é definido por uma matriz diagonal em blocos, onde cada bloco é do género (14.6). Como muitas vezes acontece com os teoremas de estrutura, as suas provas não são particularmente interessantes, nem fornecem algoritmos eficientes para calcular a base que o teorema garante existir.

Demonstração. Seja F_m um complementar de K_{m-1} em $K_m = \mathbf{V}$, ou seja, um subespaço de $F_m \subset K_m$ tal que $K_{m-1} \cap F_m = 0$ e que

$$K_m = K_{m-1} \oplus F_m$$

Se $\mathbf{v}_1^m, \mathbf{v}_2^m, \dots, \mathbf{v}_{d_m}^m$ é uma base de F_m , logo uma família livre de vetores de $K_m \setminus K_{m-1}$, então a concatenação das m -órbitas destes vetores é um sistema livre, que gera um subespaço N -invariante $E_m = \cup_{i=0}^{m-1} N^i(F_m)$ que é uma soma direta de d_m subespaços gerados por cadeias de Jordan de dimensão m . De fato, qualquer relação linear entre os vetores destas órbitas implica, aplicando repetidamente potências de N , relações lineares entre os primeiros, os segundos, os terceiros, ... vetores das órbitas, que são independentes por definição e pela injetividade das N^i 's.

Seja agora F_{m-1} um complementar de $K_{m-2} \oplus N(F_m)$ em K_{m-1} (possivelmente vazio), assim que pelo 14.10

$$K_{m-1} = N(F_m) \oplus F_{m-1} \oplus K_{m-2}$$

Se $\mathbf{v}_1^{m-1}, \mathbf{v}_2^{m-1}, \dots, \mathbf{v}_{d_{m-1}}^{m-1}$ é uma base de F_{m-1} , então a concatenação das $(m-1)$ -órbitas destes vetores é um sistema livre, que gera um subespaço N -invariante $E_{m-1} = \cup_{i=0}^{m-2} N^i(F_{m-1})$ formado por d_{m-1} cadeias de Jordan de dimensão $m-1$. Também o mesmo raciocínio anterior mostra que a concatenação destes dois sistemas livres é ainda um sistema livre, e portanto que $E_m \cap E_{m-1} = 0$.

Continuando assim, é claro que chegamos a representar cada K_{m-k} , com $k = 0, 1, \dots, m-1$, como soma direta

$$K_{m-k} = N^k(F_m) \oplus N^{k-1}(F_{m-1}) \oplus \dots \oplus N(F_{m-k+1}) \oplus F_{m-k} \oplus K_{m-k-1}$$

com o último $K_0 = 0$, e consequentemente o espaço total como uma soma direta

$$\mathbf{V} = E_m \oplus E_{m-1} \oplus \dots \oplus E_1$$

de subespaços invariantes $E_k = \cup_{i=0}^{k-1} N^i(F_k)$ (eventualmente vazios, exceto o primeiro E_m), e cada E_k é uma soma direta de d_k subespaços cíclicos de dimensão k . $\square$

É então uma consequência do teorema de estrutura 14.5 e do teorema 14.11

Teorema 14.12 (forma normal de Jordan). *Seja L é um operador definido num um espaço vetorial complexo $\mathbf{V}$ de dimensão finita. O espaço é uma soma direta $\mathbf{V} = V_1 \oplus V_2 \oplus \cdots \oplus V_p$ de espaços cíclicos V_k 's.*

Se em cada espaço cíclico escolhemos uma base de Jordan, a matriz que representa L na base resultante é uma matriz diagonal em blocos

$$J = \begin{pmatrix} J_{\mu_1} & 0 & \cdots & 0 \\ 0 & J_{\mu_2} & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & J_{\mu_p} \end{pmatrix} \quad (14.7)$$

onde cada $J_{\mu_k} = \mu_k I + N_k$ é um bloco de Jordan da forma (14.5) de dimensão $d_k \leq n$. Os μ_k 's, não necessariamente diferentes, são valores próprios de L . Sejam $\lambda_1, \lambda_2, \dots, \lambda_m$ os diferentes valores próprios de L . A multiplicidade n_k de cada valor próprio λ_k , ou seja, a dimensão do espaço próprio generalizado G_{λ_k} , é então igual à soma das dimensões d_i dos blocos de Jordan com $\mu_i = \lambda_k$. Por outro lado, a multiplicidade geométrica do valor próprio λ_k ou seja, a dimensão do espaço próprio $\text{Ker}(\lambda - L)$, é igual à cardinalidade dos blocos de Jordan com $d_i = \lambda_k$.

O polinómio caraterístico da matriz diagonal superior (14.7), logo do operador L , é um produto

$$\begin{aligned} c_L(z) &= (z - \mu_1)^{d_1} (z - \mu_2)^{d_2} \cdots (z - \mu_p)^{d_p} \\ &= (z - \lambda_1)^{n_1} (z - \lambda_2)^{n_2} \cdots (z - \lambda_m)^{n_m} \end{aligned}$$

Outra consequência da (14.7) é que o traço do operador L é igual a

$$\text{Tr}(L) = n_1 \lambda_1 + n_2 \lambda_2 + \cdots + n_m \lambda_m$$

Também interessante é observar que o polinómio minimal de L é o produto

$$m_L(z) = (z - \lambda_1)^{q_1} (z - \lambda_2)^{q_2} \cdots (z - \lambda_m)^{q_m}$$

onde o expoente q_k (que é o índice de nilpotência de $L - \lambda_k$ em G_{λ_k}) igual à dimensão do maior bloco de Jordan com $\mu_i = \lambda_k$.

Se A é a matriz que representa o operador L numa base de $\mathbf{V}$, então existe uma matriz invertível U (cujas colunas são os vetores das bases de Jordan) tal que $U^{-1}AU = J$. A “forma canónica” J é única a menos de permutações dos blocos.

e.g. Classes de equivalência de matrizes quadradas complexas. As formas canónicas de Jordan parametrizam então as classes de equivalência de matrizes quadradas semelhantes. Por exemplo, toda matriz complexa 2×2 é semelhante a uma das duas matrizes de Jordan

$$\begin{pmatrix} \lambda & 0 \\ 0 & \mu \end{pmatrix} \quad \begin{pmatrix} \lambda & 1 \\ 0 & \lambda \end{pmatrix}$$

onde λ e μ são os valores próprios (não necessariamente distintos). Toda matriz complexa 3×3 é semelhante a uma das três matrizes de Jordan

$$\begin{pmatrix} \lambda & 0 & 0 \\ 0 & \mu & 0 \\ 0 & 0 & \nu \end{pmatrix} \quad \begin{pmatrix} \lambda & 1 & 0 \\ 0 & \lambda & 0 \\ 0 & 0 & \mu \end{pmatrix} \quad \begin{pmatrix} \lambda & 1 & 0 \\ 0 & \lambda & 1 \\ 0 & 0 & \lambda \end{pmatrix}$$

onde λ , μ e ν são os valores próprios (não necessariamente distintos). Toda matriz complexa 4×4 é semelhante a uma das cinco matrizes de Jordan

$$\begin{pmatrix} \lambda & 0 & 0 & 0 \\ 0 & \mu & 0 & 0 \\ 0 & 0 & \nu & 0 \\ 0 & 0 & 0 & \theta \end{pmatrix} \quad \begin{pmatrix} \lambda & 1 & 0 & 0 \\ 0 & \lambda & 0 & 0 \\ 0 & 0 & \mu & 0 \\ 0 & 0 & 0 & \nu \end{pmatrix} \quad \begin{pmatrix} \lambda & 1 & 0 & 0 \\ 0 & \lambda & 0 & 0 \\ 0 & 0 & \mu & 1 \\ 0 & 0 & 0 & \mu \end{pmatrix}$$

$$\begin{pmatrix} \lambda & 1 & 0 & 0 \\ 0 & \lambda & 1 & 0 \\ 0 & 0 & \lambda & 0 \\ 0 & 0 & 0 & \mu \end{pmatrix} \quad \begin{pmatrix} \lambda & 1 & 0 & 0 \\ 0 & \lambda & 1 & 0 \\ 0 & 0 & \lambda & 1 \\ 0 & 0 & 0 & \lambda \end{pmatrix}$$

onde λ, μ, ν e θ são os valores próprios (não necessariamente distintos) ...

ex: Mostre que o traço da k -ésima potência de um operador L definido num espaço vetorial complexo de dimensão finita é

$$\text{Tr}(L^k) = n_1 \lambda_1^k + n_2 \lambda_2^k + \cdots + n_m \lambda_m^k$$

onde os λ_i 's são os valores próprios e os n_i 's as multiplicidades respetivas.

14.4 Complexificação

Complexificação. Se A é uma matriz quadrada real que define um operador $L : \mathbb{R}^n \rightarrow \mathbb{R}^n$ relativamente à base canónica, então a mesma matriz A , pensada como matriz complexa, define um operador $L^{\mathbb{C}} : \mathbb{C}^n \rightarrow \mathbb{C}^n$, chamado “complexificação” do operador L . Esta operação admite interpretação mais conceptual, que não depende das bases utilizadas (e que, em espaços diferentes de $\mathbb{R}^n$ e $\mathbb{C}^n$, não são canónicas!).

Seja $\mathbf{V}$ um espaço vetorial real. A sua *complexificação* é o espaço vetorial complexo $\mathbf{V}^{\mathbb{C}} := \mathbf{V} \oplus i\mathbf{V}$ formado pelas somas formais $\mathbf{z} = \mathbf{v} + i\mathbf{w}$, com $\mathbf{v}, \mathbf{w} \in \mathbf{V}$ e munido das operações de soma e produto por um escalar definidas por

$$(\mathbf{v} + i\mathbf{w}) + (\mathbf{v}' + i\mathbf{w}') := (\mathbf{v} + \mathbf{v}') + i(\mathbf{w} + \mathbf{w}')$$

e

$$(a + ib)(\mathbf{v} + i\mathbf{w}) := (a\mathbf{v} - b\mathbf{w}) + i(b\mathbf{v} + a\mathbf{w})$$

se $a + ib \in \mathbb{C}$, respetivamente. O vetor nulo é $0 := 0 + i0$. Todo vetor $\mathbf{v} \in \mathbf{V}$ pode ser identificado de maneira natural com um vetor $\mathbf{v} \approx \mathbf{v} + i0$ de $\mathbf{V}^{\mathbb{C}}$. Os vetores $\mathbf{v} \approx \mathbf{v} + i0$ e $i\mathbf{w} \approx 0 + i\mathbf{w}$ são chamados “parte real” e “parte imaginária” do vetor $\mathbf{v} + i\mathbf{w}$, respetivamente. A involução

$$\mathbf{v} + i\mathbf{w} \mapsto \overline{\mathbf{v} + i\mathbf{w}} := \mathbf{v} - i\mathbf{w}$$

que fixa os vetores reais e muda os sinais aos vetores imaginários puros, é chamada “conjugação”. É imediato verificar que $\overline{\lambda z} = \bar{\lambda} \bar{z}$ se λ é um escalar complexo.

Se $L : \mathbf{V} \rightarrow \mathbf{V}$ é um operador, então a sua *complexificação* é o operador $L^{\mathbb{C}} : \mathbf{V}^{\mathbb{C}} \rightarrow \mathbf{V}^{\mathbb{C}}$ definido por

$$L^{\mathbb{C}}(\mathbf{v} + i\mathbf{w}) := L(\mathbf{v}) + iL(\mathbf{w})$$

Se $\mathbf{V}$ tem dimensão finita e se $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ é uma sua base, então a mesma família de vetores (pensados como vetores $\mathbf{e}_k + i0$ da complexificação) é uma base de $\mathbf{V}^{\mathbb{C}}$. De fato, um vetor arbitrário $\mathbf{v} + i\mathbf{w} \in \mathbf{V}^{\mathbb{C}}$, com $\mathbf{v} = v_1\mathbf{e}_1 + v_2\mathbf{e}_2 + \cdots + v_n\mathbf{e}_n$ e $\mathbf{w} = w_1\mathbf{e}_1 + w_2\mathbf{e}_2 + \cdots + w_n\mathbf{e}_n$ em $\mathbf{V}$, pode ser representado de maneira única como $(v_1 + iw_1)\mathbf{e}_1 + (v_2 + iw_2)\mathbf{e}_2 + \cdots + (v_n + iw_n)\mathbf{e}_n$. Consequentemente, se A é a matriz real $n \times n$ que define L relativamente à esta base, então a mesma matriz, agora pensada como matriz complexa, define o operador $L^{\mathbb{C}}$ relativamente à base $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$.

e.g. A complexificação de $\mathbb{R}^n$ é, naturalmente, $\mathbb{C}^n$, que pode então ser pensado como $\mathbb{R}^n \oplus i\mathbb{R}^n$. Se um operador $L : \mathbb{R}^n \rightarrow \mathbb{R}^n$ é definido, na base canónica, pela matriz quadrada real A , então a sua complexificação é definida, na base canónica de $\mathbb{C}^n$, pela mesma matriz A .

e.g. A complexificação do espaço $\text{Pol}(\mathbb{R})$ dos polinómios $f(t) = a_0 + a_1t + \cdots + a_nt^n$ com coeficientes reais $a_0, a_1, \dots$ é o espaço $\text{Pol}(\mathbb{C})$ dos polinómios com coeficientes complexos.

ex: Os vetores $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_k$ de $\mathbf{V}^{\mathbb{C}}$ são linearmente independentes sse os vetores $\overline{\mathbf{z}}_1, \overline{\mathbf{z}}_2, \dots, \overline{\mathbf{z}}_k$ são linearmente independentes.

14.5 Forma normal de Jordan real

Estrutura dos operadores reais. Seja $L : \mathbf{V} \rightarrow \mathbf{V}$ um operador definido no espaço vetorial real $\mathbf{V}$ de dimensão finita n , e seja $L^{\mathbb{C}} : \mathbf{V}^{\mathbb{C}} \rightarrow \mathbf{V}^{\mathbb{C}}$ a sua complexificação.

Seja $\mathbf{z} = \mathbf{x} + i\mathbf{y}$ é um vetor próprio de $L^{\mathbb{C}}$, com valor próprio λ . Se λ é real, então

$$L^{\mathbb{C}}(\mathbf{x} + i\mathbf{y}) = L\mathbf{x} + iL\mathbf{y} = \lambda(\mathbf{x} + i\mathbf{y}) \quad \Rightarrow \quad L\mathbf{x} = \lambda\mathbf{x} \quad \text{e} \quad L\mathbf{y} = \lambda\mathbf{y}$$

Portanto, $\mathbf{x}$ ou $\mathbf{y}$ (pois pelo menos um dos dois não é nulo) é um vetor próprio de L com valor próprios λ . Por outro lado, se $\lambda = \alpha + i\omega$ não é real, então $\overline{\mathbf{z}} = \mathbf{x} - i\mathbf{y}$ é um vetor próprio com valor próprio $\overline{\lambda}$, pois

$$L^{\mathbb{C}}(\mathbf{x} + i\mathbf{y}) = L\mathbf{x} + iL\mathbf{y} = (\alpha + i\omega)(\mathbf{x} + i\mathbf{y}) \quad \Rightarrow \quad L\mathbf{x} = \alpha\mathbf{x} - \omega\mathbf{y} \quad \text{e} \quad L\mathbf{y} = \omega\mathbf{x} + \alpha\mathbf{y}$$

e portanto

$$L^{\mathbb{C}}(\mathbf{x} - i\mathbf{y}) = L\mathbf{x} - iL\mathbf{y} = \alpha\mathbf{x} - \omega\mathbf{y} - i(\omega\mathbf{x} + \alpha\mathbf{y}) = (\alpha - i\omega)(\mathbf{x} - i\mathbf{y})$$

Como $\lambda \neq \overline{\lambda}$, $\mathbf{z}$ e $\overline{\mathbf{z}}$ são linearmente independentes sobre os complexos, e isto implica que $\mathbf{x}$ e $\mathbf{y}$ são linearmente independentes sobre os reais, pois se $a\mathbf{x} + b\mathbf{y} = 0$ com a e b reais, então

$$0 = a\mathbf{x} + b\mathbf{y} = a\frac{\mathbf{z} + \overline{\mathbf{z}}}{2} + b\frac{\mathbf{z} - \overline{\mathbf{z}}}{2i} = \frac{a + ib}{2}\mathbf{z} + \frac{a - ib}{2}\overline{\mathbf{z}}$$

implica que $a \pm ib = 0$, logo que $a = b = 0$. Consequentemente, a parte real e a parte imaginária de um vetor próprio $\mathbf{z}$ com valor próprio não real $\lambda = \alpha + i\omega$ geram um plano L -invariante $\mathbb{R}\mathbf{x} + \mathbb{R}\mathbf{y} \subset \mathbf{V}$ onde o operador L é definido pela matriz

$$\begin{pmatrix} \alpha & -\omega \\ \omega & \alpha \end{pmatrix}$$

Uma primeira consequência é que

Teorema 14.13. *Um operador definido num espaço linear de dimensão finita admite um subespaço invariante de dimensão um ou dois.*

Um argumento por indução também mostra que

Teorema 14.14. *Seja $L^{\mathbb{C}}$ a complexificação de um operador $L : \mathbf{V} \rightarrow \mathbf{V}$ definido num espaço vetorial real. Então $\mathbf{z} \in \text{Ker}(\lambda - L^{\mathbb{C}})^k$ sse $\overline{\mathbf{z}} \in \text{Ker}(\overline{\lambda} - L^{\mathbb{C}})^k$.*

Demonstração. O caso dos vetores próprios, ou seja, $k = 1$, foi provado antes. Assumimos então o resultado verdadeiro para k , e consideramos um vetor $\mathbf{z} \in \text{Ker}(\lambda - L^{\mathbb{C}})^{k+1}$. Por definição, $0 = (\lambda - L)^{k+1}\mathbf{z} = (\lambda - L)^k(\lambda - L)\mathbf{z}$, e portanto $\mathbf{z}' = (\lambda - L)\mathbf{z} \in \text{Ker}(\lambda - L)^k$. Pela hipótese indutiva, $\overline{\mathbf{z}}' \in \text{Ker}(\overline{\lambda} - L^{\mathbb{C}})^k$. Um cálculo análogo ao cálculo feito para os vetores próprios mostra que $\overline{\mathbf{z}}' = (\overline{\lambda} - L)\overline{\mathbf{z}}$, e consequentemente que $\overline{\mathbf{z}} \in \text{Ker}(\overline{\lambda} - L^{\mathbb{C}})^{k+1}$. $\square$

Então os valores próprios não reais de $L^{\mathbb{C}}$ ocorrem em pares de números complexos conjugados λ e $\overline{\lambda}$, e os espaços próprios generalizados G_{λ} e $G_{\overline{\lambda}}$ de $L^{\mathbb{C}}$ são obtidos um do outro por conjugação. Em particular, se λ não é real, os espaços próprios generalizados G_{λ} e $G_{\overline{\lambda}}$ de $L^{\mathbb{C}}$ têm a mesma dimensão, e dão origem a subespaços L -invariantes de $\mathbf{V}$ de dimensão real igual $2 \dim G_{\lambda}$.

Em particular, o produto (14.2) que define o determinante de $L^{\mathbb{C}}$ é real, e portanto pode ser usado para definir o determinante do operador real L .

Também, temos mais uma prova, independente dos determinantes, do

Teorema 14.15. *Um operador definido num espaço linear real de dimensão finita ímpar admite pelo menos um vetor próprio.*

Finalmente, a forma normal de Jordan do operador complexificado implica o seguinte

Teorema 14.16 (forma normal de Jordan real). *Seja L um operador no espaço vetorial real $\mathbf{V}$. Então o espaço total é uma soma direta*

$$\mathbf{V} = \left(\bigoplus_{\lambda \in \mathbb{R}} E_{\lambda} \right) \oplus \left(\bigoplus_{\lambda \in \mathbb{C} \setminus \mathbb{R}} E_{\lambda, \bar{\lambda}} \right).$$

de subespaços invariantes E_{λ} , associados aos valores próprios reais λ 's, e $E_{\lambda, \bar{\lambda}}$, associados aos pares de valores próprios conjugados não reais $\lambda, \bar{\lambda}$.

Numa base apropriada, a restrição do operador L a cada subespaços invariantes E_{λ} , com $\lambda \in \mathbb{R}$, é uma soma direta de blocos de Jordan da forma (14.5). Numa base apropriada, a restrição do operador L a cada subespaços invariante $E_{\lambda, \bar{\lambda}}$, com $\lambda = \alpha + i\omega \in \mathbb{C} \setminus \mathbb{R}$, é uma soma direta de operadores com matrizes da forma

$$J = \begin{pmatrix} R & I & & \\ & R & I & \\ & & \ddots & \ddots \\ & & & R & I \\ & & & & R \end{pmatrix} \quad (14.8)$$

com

$$R = \begin{pmatrix} \alpha & -\omega \\ \omega & \alpha \end{pmatrix} \quad e \quad I = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix},$$

Consequentemente, numa base apropriada, o operador L é representado por uma matriz diagonal em blocos (14.7), onde cada bloco J_k é da forma (14.5) ou (14.8).

Demonstração. Consideramos a complexificação do operador. De acordo com o teorema 14.12, $\mathbf{V}^{\mathbb{C}}$ é uma soma direta de subespaços $L_{\lambda}^{\mathbb{C}}$ -cíclicos associados aos valores próprios λ 's de $L^{\mathbb{C}}$.

Consideramos um bloco de Jordan com valor próprio real λ . Se $\mathbf{z} = \mathbf{x} + i\mathbf{y}$ é um vetor $L_{\lambda}^{\mathbb{C}}$ -cíclico que gera o espaço cíclico V com valor próprio λ real, então $\mathbf{x}$ ou $\mathbf{y}$ é um vetor L_{λ} -cíclico real. Este vetor gera portanto um espaço cíclico real W , e portanto uma cadeia de Jordan de dimensão real igual à dimensão complexa de V .

Consideramos agora um bloco de Jordan associado a um espaço $L_{\lambda}^{\mathbb{C}}$ -cíclico $V \subset \mathbf{V}^{\mathbb{C}}$ gerado por $\mathbf{z} = \mathbf{x} + i\mathbf{y}$, com valor próprio $\lambda = \alpha + i\omega$ que não é real (ou seja, com $\omega \neq 0$). Então $L^{\mathbb{C}}$ também admite o valor próprio conjugado $\bar{\lambda} = \alpha - i\omega$, e um espaço $L_{\bar{\lambda}}^{\mathbb{C}}$ -cíclico $\bar{V}$ é obtido de V por conjugação, ou seja, gerado por $\bar{\mathbf{z}}$. Procedendo como no caso dos vetores próprios, é então fácil verificar que o plano $\mathbb{R}\mathbf{x} + \mathbb{R}\mathbf{y}$ dá origem a um subespaço L -invariante real $W \subset \mathbf{V}$, de dimensão real igual à dimensão complexa de $V \oplus \bar{V}$, onde o operador é um bloco do género (14.8). $\square$

Raízes de operadores. As involuções, por exemplo as reflexões em retas do plano, são operadores que satisfazem $R^2 = I$, ou seja, são “raízes quadradas” do operador identidade. É também claro que se a dimensão é $n \geq 1$ existem infinitas destas raízes da identidade. Rotações de ângulos particulares são também exemplos de operadores que satisfazem $R^k = I$, ou seja, são raízes da identidade. Naturalmente, a identidade pode ser substituída por um seu múltiplo λI , desde que o escalar λ admita raízes no corpo considerado. A definição natural é que o operador A é uma raiz k -ésima do operador B se $A^k = B$.

É razoável esperar que pequenas perturbações da identidade também admitem raízes, e uma noção algébrica óbvia de “pequeno” é a nilpotência (outra noção consiste em medir o “tamanho”, e depende portanto de uma definição de uma norma natural no espaço linear dos operadores ...). Por exemplo, $I + N$ com N nilpotente deve ser considerada uma pequena perturbação da identidade, pois um operador nilpotente é uma “raiz do operador nulo”. Este é o caso, e muito mais, e é uma consequência interessante da decomposição de Jordan-Chevalley 14.5.

A ideia é procurar raízes de $I + N$, com $N^{k+1} = 0$, entre os polinômios de grau k em N , que são as “séries de Taylor” $a_0 + a_1x + a_2x^2 + \cdots + a_kx^k$ em uma variável que satisfaz $x^{k+1} = 0$.

Começamos por compreender os casos simples.

e.g. Raízes quadradas de zero. O mais simples é uma raiz quadrada de zero, um operador nilpotente N tal que $N^2 = 0$. O cálculo $(I + aN)^2 = I + 2aN$ mostra que se escolhemos $a = 1/2$ então

$$(I + \frac{1}{2}N)^2 = I + N$$

ou seja, encontramos uma raiz quadrada de $I + N$. Também acontece que $(I + aN)^3 = I + 3aN$, e portanto

$$(I + \frac{1}{3}N)^3 = I + N$$

ou seja, encontramos uma raiz cúbica de $I + N$. É claro que, em geral,

$$(I + \frac{1}{n}N)^n = I + N$$

e portanto $I + N$ admite raízes n -ésimas, para todos os $n \geq 1$. É interessante observar que a esquerda temos uma famosa fórmula de Euler para a função exponencial, sem a necessidade de passar ao limite, e a direita o seu polinômio de Taylor de grau um centrado na origem. Portanto, as duas expressões devem ser consideradas equivalentes ao exponencial $\exp(N)$.

e.g. Raízes cúbicas de zero. Passamos a uma raiz cúbica de zero, um operador N tal que $N^3 = 0$. O cálculo

$$(I + aN + bN^2)^2 = I + 2aN + (a^2 + 2b)N^2$$

mostra que se escolhemos $a = 1/2$ e $b = -1/8$, então

$$(I + \frac{1}{2}N - \frac{1}{8}N^2)^2 = I + N$$

Também,

$$(I + aN + bN^2)^3 = I + 3aN + (3a^2 + 3ab)N^2$$

e portanto

$$(I + \frac{1}{3}N - \frac{1}{3}N^2)^3 = I + N$$

...

É claro que as coisas funcionam em geral.

Teorema 14.17. *Todo operador $I + N$ com N nilpotente num espaço linear de dimensão finita admite raízes n -ésimas, para todo $n \geq 1$.*

Demonstração. Seja $n \geq 1$ um inteiro arbitrário, e seja $N^{k+1} = 0$. Se

$$f(z) = 1 + a_1z + a_2z^2 + \cdots + a_kz^k,$$

então $(f(I + N))^n$ é um polinômio

$$\begin{aligned} (f(I + N))^n = & I + na_1N + \left(\binom{n}{2} a_1^2 + na_2 \right) N^2 + \left(\binom{n}{2} a_1a_2 + na_3 \right) N^3 + \\ & + \left(\binom{n}{2} a_1a_3 + \binom{n}{2} a_2^2 + na_4 \right) N^4 + \cdots \end{aligned}$$

de grau $\leq k$ em N , com coeficientes que são polinômios nos $a_1, a_2, \dots, a_k$. É evidente também que o coeficiente de N^i é uma soma de na_i e uma combinação linear de produtos dos a_j com $j < i$. Isto significa que é sempre possível escolher os a_i 's, de forma recursiva a partir do primeiro $a_1 = 1/n$, até anular todos os coeficientes de N^i com $i \geq 2$, e portanto obter $(f(I + N))^n = I + N$ (e, de fato, de uma única maneira). $\square$

Naturalmente, as raízes que o teorema garante existir não são únicas! O mais provável é (basta pensar nos exemplos das reflexões) ter famílias contínuas de raízes ...

Pelo teorema 14.5, todo operador L num espaço vetorial complexo de dimensão finita é uma soma direta de operadores do género $\lambda + N$ com N nilpotente, e é invertível sse os valores próprios λ 's são diferentes de zero. Mas se $\lambda \neq 0$, então

$$\lambda + N = \lambda (I + N')$$

com $N' = \lambda^{-1} N$ que também é nilpotente. Se ζ é uma raiz k -ésimas de λ (que existe sempre nos complexos), e R é uma raiz k -ésima de $I + N'$ (que existe pelo teorema 14.17), então ζR é uma raiz k -ésima de $\lambda + N$. A soma direta de tais raízes ζR , uma para cada espaço próprio generalizado, é uma raiz k -ésima do operador L . Temos portanto

Teorema 14.18. *Todo operador invertível num espaço linear complexo de dimensão finita admite raízes k -ésimas, para todo $k \geq 1$.*

ex: Calcule uma raiz quadrada de

$$\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 9 & 9 & 0 \\ 0 & 9 & 0 \\ 0 & 0 & 4 \end{pmatrix}$$

ex: Dê exemplos de operadores invertíveis, definidos em espaços vetoriais reais, que não admitem raízes, por exemplo quadradas.

Referências

- [Ap69] T.M. Apostol, *Calculus*, John Wiley & Sons, 1969 [*Cálculo*, Editora Reverté, 1999].
- [Ar85] V.I. Arnold, *Equações diferenciais ordinárias*, MIR, 1985.
- [Ar89] V.I. Arnold, *Metodi geometrici della teoria delle equazioni differenziali ordinarie*, Editori Riuniti - MIR, 1989.
- [Ax15] S. Axler, *Linear Algebra Done Right*, Third edition, Springer, 2015.
- [Ba77] F. Banino, *Geometria per fisici*, Feltrinelli, 1977.
- [Bo89] N. Bourbaki, *Elements of Mathematics, Algebra I*, Springer, 1989.
- [BR98] T.S. Blyth and E.F. Robertson, *Basic Linear Algebra*, McGraw Hill, 1998.
- [Di47] P.A.M. Dirac, *The Principles of Quantum Mechanics*, Clarendon Press, 1947
- [Ef17] J. Efferon, *Linear Algebra*, <http://joshua.smcvt.edu/linearalgebra>, 2017.
- [FIS03] S.H. Friedberg, A.J. Insel and L.E. Spence, *Linear Algebra*, Prentice Hall, 2003.
- [Go96] R. Godement, *Cours d'algèbre*, Hermann Éditeurs, 1996.
- [Ha58] P.R. Halmos, *Finite dimensional vector spaces*, Van Nostrand, 1958.
- [KKR62] C. Kittel, W.D. Knight and M.A. Ruderman, *Berkeley Physics*, McGraw-Hill, 1962.
- [La87] S. Lang, *Linear Algebra*, Third Edition, UTM Springer, 1987.
- [La97] S. Lang, *Introduction to Linear Algebra*, Second Edition, UTM Springer, 1997.
- [LL78] L.D. Landau e E.M. Lifshitz, *Mecânica*, MIR, 1978.
- [Li95] E.L. Lima, *Curso de Análise, volume 1*, Projeto Euclides, 1995.
- [Ma90] L.T. Magalhães, *Álgebra Linear como Introdução à Matemática Aplicada*, Texto Editora, 1990.
- [Me00] C.D. Meyer, *Matrix Analysis and Applied Linear Algebra*, SIAM, 2000.
- [MB99] S. MacLane and G. Birkhoff, *Algebra (Third Edition)*, AMS Chelsea Publishing, 1999.
- [Pe05] R. Penrose, *The Road to Reality: A Complete Guide to the Laws of the Universe*, Knopf, 2005.
- [Po82] M.M. Postnikov, *Lectures in Geometry. Semester I - Analytic Geometry and Semester II - Linear Algebra and differential geometry*, MIR, 1982.
- [RHB06] K.F. Riley, M.P. Hobson and S.J. Bence, *Mathematical Methods for Physics and Engineering*, Cambridge University Press, 2006.
- [Se89] E. Sernesi, *Geometria 1*, Bollati Boringhieri, 1989.
- [Sh77] E.G. Shilov, *Linear algebra*, Dover, 1977.
- [St98] G. Strang, *Linear Algebra and its Applications*, Hartcourt Brace Jonovich Publishers, 1998.
- [St09] G. Strang, *Introduction to Linear Algebra*, fourth edition, Wellesley-Cambridge Press and SIAM 2009.
- [Wa91] B.L. van der Waerden, *Algebra*, Springer, 1991 [*Moderne Algebra*, 1930-1931].
- [We52] H. Weyl, *Space Time Matter*, Dover, 1952 [*Raum Zeit Materie*, 1921].

Índice

aceleração, 16
álgebra, 109
ângulo, 25
aniquilador, 86
argumento, 61
 principal, 62
automorfismo, 96
autovalor, 150
autovetor, 150

base, 44, 73
 canônica, 18
 de Jordan, 170
 dual, 81
 ortonormada, 45
bloco de Jordan, 170
bola, 26

cadeia
 de Jordan, 170
 de Markov, 162
caraterística, 104
circunferência
 unitária, 60
co-dimensão, 79
co-quaterniões, 101
co-vetor, 81
coeficiente, 18, 40
combinação linear, 18, 40
complementar (subespaço), 78
complemento algébrico, 142
complexificação, 173
componente, 18, 24, 73
comprimento, 21
comutador, 92, 110
conjugação, 59, 118
convexo, 37
coordenada, 18, 73

delta de Dirac, 85
dependência linear, 42
desigualdade
 de Schwarz, 24
 do triângulo, 25
deslizamento, 113, 149
deslocamento, 114
determinante, 47, 137, 157
diagonal, 109
dimensão, 44, 73
discriminante, 61
distância, 21

eliminação de Gauss-Jordan, 127
endomorfismo, 87

equação
 homogênea, 84
 linear, 84
esfera, 26
 unitária, 23
espaço
 cíclico, 169
 de funções, 71
 dual, 81
 finitamente gerado, 73
 linear/vetorial, 68
 nulo, 83, 89
 próprio, 151
 próprio generalizado, 165
 quociente, 79
espectro, 158
expansão linear, 40, 70
exponencial, 62

fórmula
 de Lagrange, 49
 de Laplace, 142
família
 ortogonal, 45
forma
 alternada, 135
 linear, 81, 135
fórmula
 de de Moivre, 64
 de Euclides, 61
 de Euler, 61

geradores, 33, 41, 70
grau, 65

hiperplano, 84
homomorfismo, 87
homotetia, 18

identidade
 de Diofanto, 60
 de Jacobi, 48, 110
 de Lagrange, 49
 de polarização, 23
 do paralelogramo, 23
 dos quatros quadrados de Euler, 76
imagem, 90
independência linear, 42, 73
inteiro gaussiano, 60
involução, 96, 156
isomorfismo, 96
isomorfismo linear, 70

lista, 14
livre

- conjunto, 42, 73
- métrica euclidiana, 25
- módulo, 21, 59
- matriz, 47, 98
 - anti-simétrica, 105
 - circulante, 161
 - de transição, 163
 - diagonal, 109
 - diagonal superior, 116
 - diagonalizável, 159
 - dos co-fatores, 142
 - elementar, 131
 - em blocos, 149
 - em escada de linhas, 127
 - estocástica, 162
 - identidade, 100
 - inversa, 115
 - invertível/regular/não-singular, 115
 - semelhante, 118
 - simétrica, 105
 - transposta, 105
- multiplicidade, 166
 - algébrica, 158, 167
 - geométrica, 151
- núcleo, 83, 89
- número complexo, 56
- norma
 - de Frobenius, 121
 - euclidiana, 22
- nulidade, 90
- operações elementares, 127
- operador, 87
 - diagonalizável, 155
 - nilpotente, 94
 - semi-simples, 155
- órbita, 148
- ordem, 90
- paralelepípedo, 53, 137
- parte
 - imaginária, 59
- parte real, 59
- pivot, 127
- plano, 33
- polinómio, 61, 65
 - caraterístico, 157
 - minimal, 167
- polinómios de Chebyshev, 64
- princípio de sobreposição, 126
- produto
 - cartesiano, 79
 - escalar, 21
 - linhas por colunas, 99
 - misto, 51
 - por um escalar, 17, 68
 - triplo, 51
 - vetorial, 48
- projeção, 94, 155
 - ortogonal, 24, 113
- quase-polinómio, 149, 170
- quaterniões, 75
- raízes da unidade, 63
- raio espectral, 158
- raiz, 61
- reflexão, 112
- regra de Cramer, 53, 146
- representação polar, 61
- reta, 29
- rotação, 111
- símbolo
 - de Kronecker, 81, 100
 - de Levi-Civita, 136
- segmento, 37
- semi-simples, 155
- simplexo unitário, 38
- sistema
 - homogéneo, 85, 124
 - linear, 124
- sobreposição, 18
- soma
 - de subespaços, 77
 - de vetores, 17, 68
 - direta, 78
- subespaço
 - cíclico, 148
 - invariante, 148
 - linear/vetorial, 40, 69
- sucessões, 72
- teorema
 - da alternativa de Fredholm, 126
 - da ordem-nulidade, 90
 - de Cayley-Hamilton, 168
 - de Gauss, 66
 - de Kronecker-Capelli, 125
 - de Pitágoras, 23
 - fundamental da álgebra, 66
- terno pitagórico, 61
- traço, 120, 157, 172
- trajetória, 16
- transformação
 - dual/transposta, 106
 - inversa, 95
 - linear, 70, 87
- translação, 18
- triângulo, 37
- valor
 - absoluto, 59

próprio, 150
velocidade, 16
vetor, 17, 68
aplicado, 16

direccional, 29
próprio, 150
próprio generalizado , 165
unitário, 23