

M06: LINEAR ALGEBRA II

Thomas Mettler

Spring Semester 2025

Version 16 Jan 2025



FernUni.ch
UniDistance.ch

Contents

Part 1. Linear Algebra I	7
Chapter 1 Fields and complex numbers	9
1.1 Fields	9
1.2 Complex numbers	12
Chapter 2 Matrices	17
2.1 Definitions	17
2.2 Matrix operations	19
2.3 Mappings associated to matrices	23
Chapter 3 Vector spaces and linear maps	27
3.1 Vector spaces	27
3.2 Linear maps	30
3.3 Vector subspaces and isomorphisms	34
3.4 Generating sets	37
3.5 Linear independence and bases	39
3.6 The dimension	43
3.7 Matrix representation of linear maps	51
Chapter 4 Applications of Gaussian elimination	61
4.1 Gaussian elimination	61
4.2 Applications	63
Chapter 5 The determinant	71
5.1 Axiomatic characterisation	71
5.2 Uniqueness of the determinant	74
5.3 Existence of the determinant	76
5.4 Properties of the determinant	80
5.5 Permutations	81
5.6 The Leibniz formula	84
5.7 Cramer's rule	87
Chapter 6 Endomorphisms	91
6.1 Sums, direct sums and complements	91
6.2 Invariants of endomorphisms	94
6.3 Eigenvectors and eigenvalues	97
6.4 The characteristic polynomial	101
6.5 Properties of eigenvalues	104
6.6 Special endomorphisms	108
Chapter 7 Quotient vector spaces	111
7.1 Affine mappings and affine spaces	111
7.2 Quotient vector spaces	112
Part 2. Linear Algebra II	115

Chapter 8	Symmetry and groups	117
8.1	Symmetry	117
8.2	Groups	118
8.3	Group actions	120
Chapter 9	Bilinear forms	125
9.1	Definitions and basic properties	125
9.2	Symmetric bilinear forms	131
Chapter 10	Euclidean spaces	137
10.1	Inner products	137
10.2	The orthogonal projection	141
10.3	Gram–Schmidt orthonormalisation	145
10.4	The orthogonal group	151
10.5	The adjoint mapping	155
10.6	The spectral theorem	158
10.7	Quadratic forms	161
Chapter 11	Unitary spaces	167
11.1	Hermitian inner products	167
11.2	The unitary group	173
11.3	Adjoint and normal endomorphisms	174
Chapter 12	The Jordan normal form	179
12.1	Generalised eigenvectors and eigenspaces	179
12.2	Jordan blocks	184
12.3	Nilpotent endomorphisms	187
12.4	Calculations	190
12.5	Applications	192
Chapter 13	Duality	197
13.1	The dual vector space	197
13.2	The transpose map	200
13.3	Properties of the transpose	201

Acknowledgements

I am grateful to Micha Wasem for undertaking the meticulous task of proofreading the complete set of lecture notes and offering insights for enhancing pedagogy. Micha's contributions also include crafting exercises with solutions, formulating multiple-choice questions, creating figures, and coding animations. I am also grateful to Keegan Flood for contributing insightful multiple-choice questions.

Furthermore, I would like to express my appreciation to a group of diligent students for spotting typos, in particular to Stéphane Billeter, Daniele Bolla, Johanna Bühler and Liborio Costa.

These lectures notes are inspired by the following sources:

- *Algebra* by Michael Artin, Birkhäuser Grundstudium der Mathematik.
- *Linear Algebra Done Right* by Sheldon Axler, Springer Undergraduate Texts in Mathematics.
- *Linear Algebra* by Emmanuel Kowalski, lecture notes available from his home page at ETH Zurich.
- *Introduction to Linear Algebra* by Gilbert Strang, Wellesley-Cambridge Press

HTML Version

These lecture notes are also available in an HTML version and in app form.

<https://apptest.fernuni.ch>

The HTML version contains the lectures notes and additionally animations, solutions to the exercises and multiple choice questions.

Part 1

Linear Algebra I

Fields and complex numbers

WEEK 1

1.1 Fields

A field $\mathbb{K}$ is roughly speaking a number system in which we can add and multiply numbers, so that the expected properties hold. We will only briefly state the basic facts about fields. For a more detailed account, we refer to the algebra module.

Definition 1.1 A *field* consists of a set $\mathbb{K}$ containing distinguished elements $0_{\mathbb{K}} \neq 1_{\mathbb{K}}$, as well as two binary operations, *addition* $+_{\mathbb{K}} : \mathbb{K} \times \mathbb{K} \rightarrow \mathbb{K}$ and *multiplication* $\cdot_{\mathbb{K}} : \mathbb{K} \times \mathbb{K} \rightarrow \mathbb{K}$, so that the following properties hold:

- Commutativity of addition

$$x +_{\mathbb{K}} y = y +_{\mathbb{K}} x \quad \text{for all } x, y \in \mathbb{K}.$$

- Commutativity of multiplication

$$(1.1) \quad x \cdot_{\mathbb{K}} y = y \cdot_{\mathbb{K}} x \quad \text{for all } x, y \in \mathbb{K}.$$

- Associativity of addition

$$(1.2) \quad (x +_{\mathbb{K}} y) +_{\mathbb{K}} z = x +_{\mathbb{K}} (y +_{\mathbb{K}} z) \quad \text{for all } x, y, z \in \mathbb{K}.$$

- Associativity of multiplication

$$(1.3) \quad (x \cdot_{\mathbb{K}} y) \cdot_{\mathbb{K}} z = x \cdot_{\mathbb{K}} (y \cdot_{\mathbb{K}} z) \quad \text{for all } x, y, z \in \mathbb{K}.$$

- $0_{\mathbb{K}}$ is the identity element of addition

$$(1.4) \quad x +_{\mathbb{K}} 0_{\mathbb{K}} = 0_{\mathbb{K}} +_{\mathbb{K}} x = x \quad \text{for all } x \in \mathbb{K}.$$

- $1_{\mathbb{K}}$ is the identity element of multiplication

$$(1.5) \quad x \cdot_{\mathbb{K}} 1_{\mathbb{K}} = 1_{\mathbb{K}} \cdot_{\mathbb{K}} x = x \quad \text{for all } x \in \mathbb{K}.$$

- For any $x \in \mathbb{K}$ there exists a unique element, denoted by $(-x)$ and called the *additive inverse* of x , such that

$$(1.6) \quad x +_{\mathbb{K}} (-x) = (-x) +_{\mathbb{K}} x = 0_{\mathbb{K}}.$$

- For any $x \in \mathbb{K} \setminus \{0_{\mathbb{K}}\}$ there exists a unique element, denoted by x^{-1} or $\frac{1}{x}$ and called the *multiplicative inverse* of x , such that

$$(1.7) \quad x \cdot_{\mathbb{K}} \frac{1}{x} = \frac{1}{x} \cdot_{\mathbb{K}} x = 1_{\mathbb{K}}.$$

- Distributivity of multiplication over addition

$$(1.8) \quad (x +_{\mathbb{K}} y) \cdot_{\mathbb{K}} z = x \cdot_{\mathbb{K}} z +_{\mathbb{K}} y \cdot_{\mathbb{K}} z \quad \text{for all } x, y, z \in \mathbb{K}.$$

Remark 1.2

- It is customary to simply speak of a field $\mathbb{K}$, without explicitly mentioning $0_{\mathbb{K}}$, $1_{\mathbb{K}}$ and $+_{\mathbb{K}}$, $\cdot_{\mathbb{K}}$.

- When $\mathbb{K}$ is clear from the context, we often simply write 0 and 1 instead of $0_{\mathbb{K}}$ and $1_{\mathbb{K}}$. Likewise, it is customary to write $+$ instead of $+\mathbb{K}$ and $\cdot$ instead of $\cdot\mathbb{K}$. Often $\cdot\mathbb{K}$ is omitted entirely so that we write xy instead of $x \cdot\mathbb{K} y$.
- We refer to the elements of a field as *scalars*.
- The set $\mathbb{K} \setminus \{0_{\mathbb{K}}\}$ is usually denoted by $\mathbb{K}^*$.
- For all $x, y \in \mathbb{K}$ we write $x - y = x +_{\mathbb{K}} (-y)$ and for all $x \in \mathbb{K}$ and $y \in \mathbb{K}^*$ we write $\frac{x}{y} = x \cdot\mathbb{K} \frac{1}{y} = x \cdot\mathbb{K} y^{-1}$.
- A field $\mathbb{K}$ containing only finitely many elements is called *finite*. Algorithms in cryptography are typically based on finite fields.

Example 1.3

- The rational numbers or quotients $\mathbb{Q}$, the real numbers $\mathbb{R}$ and the complex numbers $\mathbb{C}$ – that we will study more carefully below – equipped with the usual addition and multiplication are examples of fields.
- The integers $\mathbb{Z}$ (with usual addition and multiplication) are not a field, as only 1 and -1 admit a multiplicative inverse.
- Considering a set $\mathbb{F}_2$ consisting of only two elements that we may denote by 0 and 1, we define $+_{\mathbb{F}_2}$ and $\cdot_{\mathbb{F}_2}$ via the following tables

$+_{\mathbb{F}_2}$	0	1
0	0	1
1	1	0

and

$\cdot_{\mathbb{F}_2}$	0	1
0	0	0
1	1	1

For instance, we have $1 +_{\mathbb{F}_2} 1 = 0$ and $1 \cdot_{\mathbb{F}_2} 1 = 1$. Then, one can check that $\mathbb{F}_2$ equipped with these operations is indeed a field. A way to remember these tables is to think of 0 as representing the even numbers, while 1 represents the odd numbers. So for instance, a sum of two odd numbers is even and a product of two odd numbers is odd. Alternatively, we may think of 0 and 1 representing the boolean values *FALSE* and *TRUE*. In doing so, $+_{\mathbb{F}_2}$ corresponds to the logical *XOR* and $\cdot_{\mathbb{F}_2}$ corresponds to the logical *AND*.

- Considering a set $\mathbb{F}_4$ consisting of four elements, say $\{0, 1, a, b\}$, we define $+_{\mathbb{F}_4}$ and $\cdot_{\mathbb{F}_4}$ via the following tables

$+_{\mathbb{F}_4}$	0	1	a	b
0	0	1	a	b
1	1	0	b	a
a	a	b	0	1
b	b	a	1	0

and

$\cdot_{\mathbb{F}_4}$	0	1	a	b
0	0	0	0	0
1	1	0	1	a
a	a	0	a	b
b	b	0	b	1

Again one can check that $\mathbb{F}_4$ equipped with these operations is indeed a field.

Lemma 1.4 (Field properties) *In a field $\mathbb{K}$ we have the following properties:*

- $0_{\mathbb{K}} \cdot\mathbb{K} x = 0_{\mathbb{K}}$ for all $x \in \mathbb{K}$.
- $-x = (-1_{\mathbb{K}}) \cdot\mathbb{K} x$ for all $x \in \mathbb{K}$.
- For all $x, y \in \mathbb{K}$, if $x \cdot\mathbb{K} y = 0_{\mathbb{K}}$, then $x = 0_{\mathbb{K}}$ or $y = 0_{\mathbb{K}}$.
- $-0_{\mathbb{K}} = 0_{\mathbb{K}}$.
- $(1_{\mathbb{K}})^{-1} = 1_{\mathbb{K}}$.
- $-(-x) = x$ for all $x \in \mathbb{K}$.
- $(-x) \cdot\mathbb{K} y = x \cdot\mathbb{K} (-y) = -(x \cdot\mathbb{K} y)$.
- $(x^{-1})^{-1} = x$ for all $x \in \mathbb{K}^*$.

Proof We will only prove some of the items, the rest are an exercise for the reader.

(i) Using (1.4), we obtain $0_{\mathbb{K}} +_{\mathbb{K}} 0_{\mathbb{K}} = 0_{\mathbb{K}}$. Hence for all $x \in \mathbb{K}$ we have

$$x \cdot_{\mathbb{K}} 0_{\mathbb{K}} = x \cdot_{\mathbb{K}} (0_{\mathbb{K}} + 0_{\mathbb{K}}) = x \cdot_{\mathbb{K}} 0_{\mathbb{K}} +_{\mathbb{K}} x \cdot_{\mathbb{K}} 0_{\mathbb{K}},$$

where the second equality uses (1.8). Adding the additive inverse of $x \cdot_{\mathbb{K}} 0_{\mathbb{K}}$, we get

$$x \cdot_{\mathbb{K}} 0_{\mathbb{K}} - x \cdot_{\mathbb{K}} 0_{\mathbb{K}} = (x \cdot_{\mathbb{K}} 0_{\mathbb{K}} +_{\mathbb{K}} x \cdot_{\mathbb{K}} 0_{\mathbb{K}}) - x \cdot_{\mathbb{K}} 0_{\mathbb{K}}$$

using the associativity of addition (1.2) and (1.6), this last equation is equivalent to

$$0_{\mathbb{K}} = x \cdot_{\mathbb{K}} 0_{\mathbb{K}}$$

as claimed.

(iii) Let $x, y \in \mathbb{K}$ such that $x \cdot_{\mathbb{K}} y = 0_{\mathbb{K}}$. If $x = 0_{\mathbb{K}}$ then we are done, so suppose $x \neq 0_{\mathbb{K}}$. Using (1.7), we have $1_{\mathbb{K}} = x^{-1} \cdot_{\mathbb{K}} x$. Multiplying this equation with y we obtain

$$y = y \cdot_{\mathbb{K}} 1_{\mathbb{K}} = y \cdot_{\mathbb{K}} (x \cdot_{\mathbb{K}} x^{-1}) = (y \cdot_{\mathbb{K}} x) \cdot_{\mathbb{K}} x^{-1} = 0_{\mathbb{K}} \cdot_{\mathbb{K}} x^{-1} = 0_{\mathbb{K}}$$

where we have used (1.5), the commutativity (1.1) and associativity (1.3) of multiplication as well as (i) from above.

(v) By (1.5), we have $1_{\mathbb{K}} \cdot_{\mathbb{K}} 1_{\mathbb{K}} = 1_{\mathbb{K}}$, hence $1_{\mathbb{K}}$ is the multiplicative inverse of $1_{\mathbb{K}}$ and since the multiplicative inverse is unique, it follows that $(1_{\mathbb{K}})^{-1} = 1_{\mathbb{K}}$. $\square$

For a positive integer $n \in \mathbb{N}$ and an element x of a field $\mathbb{K}$, we write

$$nx = \underbrace{x +_{\mathbb{K}} x +_{\mathbb{K}} x +_{\mathbb{K}} \cdots +_{\mathbb{K}} x}_{n \text{ summands}}.$$

The field $\mathbb{F}_2$ has the property that $2x = 0$ for all $x \in \mathbb{F}_2$. In this case we say the $\mathbb{F}_2$ has characteristic 2. More generally, the smallest positive integer p such that $px = 0_{\mathbb{K}}$ for all $x \in \mathbb{K}$ is called the *characteristic of the field*. In the case where no such integer exists the field is said to have characteristic 0. So $\mathbb{Q}, \mathbb{R}, \mathbb{C}$ are fields of characteristic 0. It can be shown that the characteristic of any field is either 0 or a prime number.

A subset $\mathbb{F}$ of a field $\mathbb{K}$ that is itself a field, when equipped with the multiplication and addition of $\mathbb{K}$, is called a *subfield* of $\mathbb{K}$.

Example 1.5

- (i) The rational numbers $\mathbb{Q}$ form a subfield of the real numbers $\mathbb{R}$. Furthermore, as we will see below, the real numbers $\mathbb{R}$ can be interpreted as a subfield of the complex numbers $\mathbb{C}$.
- (ii) $\mathbb{F}_2$ may be thought of as the subfield of $\mathbb{F}_4$ consisting of $\{0, 1\}$.

Throughout your studies in mathematics, you will encounter various mappings having names ending in *morphism*, such as *homomorphism*, *isomorphism*, *endomorphism*, *automorphism*. This is quite confusing and to make things worse, the precise meaning of $\star$ -morphism depends on the structure of the set between which the mapping is defined. But don't worry, we will introduce one $\star$ -morphism at a time, starting with *homomorphism*. Broadly speaking, a *homomorphism* between sets $\mathcal{X}$ and $\mathcal{Y}$ that are equipped with some extra structure of the same type is a map $f : \mathcal{X} \rightarrow \mathcal{Y}$ that *respects* the extra structure.

In the case of a field $\mathbb{K}$, the extra structure consists of addition $+_{\mathbb{K}}$, multiplication $\cdot_{\mathbb{K}}$, the identity element of multiplication $1_{\mathbb{K}}$ and the identity element of addition $0_{\mathbb{K}}$. A field homomorphism respects this structure. More precisely:

Definition 1.6 (Field homomorphism) Let $\mathbb{F}$ and $\mathbb{K}$ be fields. A *field homomorphism* is a mapping $\chi : \mathbb{F} \rightarrow \mathbb{K}$ satisfying $\chi(1_{\mathbb{F}}) = 1_{\mathbb{K}}$ as well as

$$\chi(x +_{\mathbb{F}} y) = \chi(x) +_{\mathbb{K}} \chi(y) \quad \text{and} \quad \chi(x \cdot_{\mathbb{F}} y) = \chi(x) \cdot_{\mathbb{K}} \chi(y)$$

for all $x, y \in \mathbb{F}$.

Example 1.7 From the above tables we see that $\chi : \mathbb{F}_2 \rightarrow \mathbb{F}_4$ defined by $\chi(1_{\mathbb{F}_2}) = 1_{\mathbb{F}_4}$ and $\chi(0_{\mathbb{F}_2}) = 0_{\mathbb{F}_4}$ is a field homomorphism.

Remark 1.8

- We certainly also want that a field homomorphism $\chi : \mathbb{F} \rightarrow \mathbb{K}$ satisfies $\chi(0_{\mathbb{F}}) = 0_{\mathbb{K}}$. It turns out that we don't have to ask for this in the definition of a field homomorphism, it is automatically satisfied with [Definition 1.6](#). Indeed, we have

$$\chi(0_{\mathbb{F}}) = \chi(0_{\mathbb{F}} +_{\mathbb{F}} 0_{\mathbb{F}}) = \chi(0_{\mathbb{F}}) +_{\mathbb{K}} \chi(0_{\mathbb{F}}).$$

Adding the additive inverse of $\chi(0_{\mathbb{F}})$ in $\mathbb{K}$, we conclude that $0_{\mathbb{K}} = \chi(0_{\mathbb{F}})$.

- A field homomorphism is injective. Suppose $x, y \in \mathbb{F}$ satisfy $\chi(x) = \chi(y)$ so that $\chi(x - y) = 0_{\mathbb{K}}$. Assume $w = x - y \neq 0_{\mathbb{F}}$, then $\chi(w) \cdot_{\mathbb{K}} \chi(w^{-1}) = \chi(1_{\mathbb{F}}) = 1_{\mathbb{K}}$. Since by assumption $\chi(w) = 0_{\mathbb{K}}$, we thus obtain $0_{\mathbb{K}} \cdot_{\mathbb{K}} \chi(w^{-1}) = 1_{\mathbb{K}}$, contradicting [Lemma 1.4](#) (i). It follows that $x = y$ and hence χ is injective.

1.2 Complex numbers

Video Complex numbers

Historically the complex numbers arose from an interest to make sense of the square root of a negative number. We may picture the rational numbers $\mathbb{Q}$ as elements of an infinite number line with an origin 0. Positive numbers extending to the right of the origin and negative numbers to the left. Mathematicians have observed early on that this line of numbers contains elements, such as π or $\sqrt{2}$, that are not quotients. Phrased differently, the rational numbers do not fill out the whole number line, there are gaps consisting of *irrational numbers*. In a sense to be made precise in the Analysis module, the real numbers may be thought of as the union of the rational numbers and the gaps on the number line, resulting in a gap less line of numbers, known as the *complete field of real numbers*.

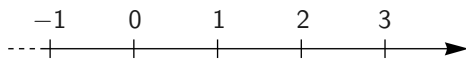


FIGURE 1.1. The real number line.

The square x^2 of a real number x is a non-negative real number, $x^2 \geq 0$, hence if we want to define what the square root of a negative number ought to be, we are in trouble, since there are no numbers left on the line of numbers that we might use. The solution is to consider pairs of real numbers instead. A complex number is an ordered pair (x, y) of real numbers $x, y \in \mathbb{R}$. We denote the set of complex numbers by $\mathbb{C}$. We equip $\mathbb{C}$ with the addition defined by the rule

$$(x_1, y_1) +_{\mathbb{C}} (x_2, y_2) = (x_1 + x_2, y_1 + y_2)$$

for all (x_1, y_1) and $(x_2, y_2) \in \mathbb{C}$ and where $+$ on the right denotes the usual addition $+\mathbb{R}$ of real numbers. Furthermore, we equip $\mathbb{C}$ with the multiplication defined by the rule

$$(1.9) \quad (x_1, y_1) \cdot_{\mathbb{C}} (x_2, y_2) = (x_1 \cdot x_2 - y_1 \cdot y_2, x_1 \cdot y_2 + y_1 \cdot x_2).$$

for all (x_1, y_1) and $(x_2, y_2) \in \mathbb{C}$ and where $\cdot$ on the right denotes the usual multiplication $\cdot_{\mathbb{R}}$ of real numbers.

Definition 1.9 (Complex numbers) The set $\mathbb{C}$ together with the operations $+_{\mathbb{C}}$, $\cdot_{\mathbb{C}}$ and $0_{\mathbb{C}} = (0, 0)$ and $1_{\mathbb{C}} = (1, 0)$ is called the *field of complex numbers*.

The mapping $\chi : \mathbb{R} \rightarrow \mathbb{C}, x \mapsto (x, 0)$ is a field homomorphism. Indeed,

$$\chi(x_1 +_{\mathbb{R}} x_2) = (x_1 +_{\mathbb{R}} x_2, 0) = (x_1, 0) +_{\mathbb{C}} (x_2, 0) = \chi(x_1) +_{\mathbb{C}} \chi(x_2),$$

$$\chi(x_1 \cdot_{\mathbb{R}} x_2) = (x_1 \cdot_{\mathbb{R}} x_2, 0) = (x_1, 0) \cdot_{\mathbb{C}} (x_2, 0) = \chi(x_1) \cdot_{\mathbb{C}} \chi(x_2),$$

for all $x_1, x_2 \in \mathbb{R}$ and $\chi(1) = (1, 0) = 1_{\mathbb{C}}$.

This allows to think of the real numbers $\mathbb{R}$ as the subfield $\{(x, 0) | x \in \mathbb{R}\}$ of the complex numbers $\mathbb{C}$. Because of the injectivity of χ , it is customary to identify x with $\chi(x)$, hence abusing notation, we write $(x, 0) = x$.

Notice that $(0, 1)$ satisfies $(0, 1) \cdot_{\mathbb{C}} (0, 1) = (-1, 0)$ and hence is a square root of the real number $(-1, 0) = -1$. The number $(0, 1)$ is called the *imaginary unit* and usually denoted by i . Sometimes the notation $\sqrt{-1}$ is also used. Every complex number $(x, y) \in \mathbb{C}$ can now be written as

$$(x, y) = (x, 0) +_{\mathbb{C}} (0, y) = (x, 0) +_{\mathbb{C}} i \cdot_{\mathbb{C}} (y, 0) = x + iy,$$

where we follow the usual custom of omitting $\cdot_{\mathbb{C}}$ and writing $+$ instead of $+_{\mathbb{C}}$ on the right hand side. With this convention, complex numbers can be manipulated as real numbers, we just need to keep in mind that i satisfies $i^2 = -1$. For instance, the multiplication of complex numbers $x_1 + iy_1$ and $x_2 + iy_2$ gives

$$(x_1 + iy_1)(x_2 + iy_2) = x_1x_2 + i^2y_1y_2 + i(x_1y_2 + y_1x_2) = x_1x_2 - y_1y_2 + i(x_1y_2 + y_1x_2)$$

in agreement with (1.9). Here we also follow the usual custom of omitting $\cdot_{\mathbb{R}}$ on the right hand side.

Definition 1.10 For a complex number $z = x + iy \in \mathbb{C}$ with $x, y \in \mathbb{R}$ we call

- $\operatorname{Re}(z) = x$ its *real part*;
- $\operatorname{Im}(z) = y$ its *imaginary part*;
- $\bar{z} = x - iy$ the *complex conjugate* of z ;
- $|z| = \sqrt{z\bar{z}} = \sqrt{x^2 + y^2}$ the *absolute value or modulus* of z .

The mapping $z \mapsto \bar{z}$ is called *complex conjugation*.

Remark 1.11

- For $z \in \mathbb{C}$ the following statements are equivalent

$$z \in \mathbb{R} \iff \operatorname{Re}(z) = z \iff \operatorname{Im}(z) = 0 \iff z = \bar{z}.$$

- We have $|z| = 0$ if and only if $z = 0$.

Example 1.12 Let $z = \frac{2+5i}{6-i}$. Then

$$z = \frac{(2+5i)\overline{(6-i)}}{(6-i)\overline{(6-i)}} = \frac{(2+5i)(6+i)}{|6-i|^2} = \frac{1}{37}(7+32i),$$

so that $\operatorname{Re}(z) = \frac{7}{37}$ and $\operatorname{Im}(z) = \frac{32}{37}$. Moreover,

$$|z| = \sqrt{\left(\frac{7}{37}\right)^2 + \left(\frac{32}{37}\right)^2} = \sqrt{\frac{29}{37}}.$$

Remark 1.13

- We may think of a complex number $z = a + ib$ as a point or a vector in the plane $\mathbb{R}^2$ with x -coordinate a and y -coordinate b .
- The real numbers form the horizontal coordinate axis (the real axis) and the *purely imaginary complex numbers* $\{iy | y \in \mathbb{R}\}$ form the vertical coordinate axis (the imaginary axis).
- The point $\bar{z}$ is obtained by reflecting z along the real axis.
- $|z|$ is the distance of z to the origin $0_{\mathbb{C}} = (0, 0) \in \mathbb{C}$.
- The addition of complex numbers corresponds to the usual vector addition.
- For the geometric significance of the multiplication, we refer the reader to the Analysis module.

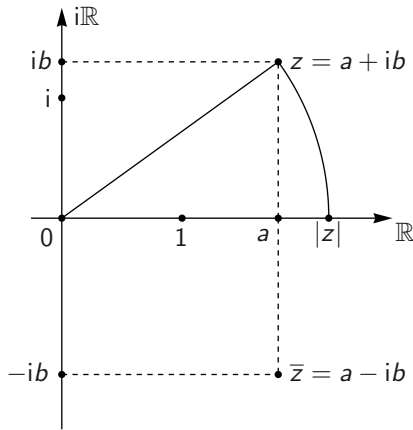


FIGURE 1.2. The complex number plane $\mathbb{C}$

We have the following elementary facts about complex numbers:

Proposition 1.14 For all $z, w \in \mathbb{C}$ we have

- (i) $\operatorname{Re}(z) = \frac{z+\bar{z}}{2}$, $\operatorname{Im}(z) = \frac{z-\bar{z}}{2i}$;
- (ii) $\operatorname{Re}(z+w) = \operatorname{Re}(z) + \operatorname{Re}(w)$, $\operatorname{Im}(z+w) = \operatorname{Im}(z) + \operatorname{Im}(w)$;
- (iii) $\overline{z+w} = \bar{z} + \bar{w}$, $\overline{zw} = \bar{z}\bar{w}$, $\bar{\bar{z}} = z$;
- (iv) $|z|^2 = |\bar{z}|^2 = z\bar{z} = \operatorname{Re}(z)^2 + \operatorname{Im}(z)^2$;
- (v) $|zw| = |z||w|$.

Proof Exercise.

□

Exercises

Exercise 1.15 Check that $\mathbb{C}$ is indeed a field.

Matrices

WEEK 2

2.1 Definitions

A matrix (plural matrices) is simply a rectangular block of numbers. As we will see below, every matrix gives rise to a mapping sending a finite list of numbers to another finite list of numbers. Mappings arising from matrices are called *linear* and linear mappings are among the most fundamental objects in mathematics. In the Linear Algebra modules we develop the theory of linear maps as well as the theory of *vector spaces*, the natural habitat of linear maps. While this theory may come across as quite abstract, it is in fact at the heart of many real world applications, including optics and quantum physics, radio astronomy, MP3 and JPEG compression, X-ray crystallography, MRI scans and machine learning, just to name a few.

Throughout the Linear Algebra modules, $\mathbb{K}$ stands for either the real numbers $\mathbb{R}$ or the complex numbers $\mathbb{C}$, but almost all statements are also valid over arbitrary fields.

We start with some definitions. In this chapter, $m, n, \tilde{m}, \tilde{n}$ denote natural numbers.

Definition 2.1 (Matrix)

- A rectangular block of scalars $A_{ij} \in \mathbb{K}$, $1 \leq i \leq m$, $1 \leq j \leq n$

$$(2.1) \quad \mathbf{A} = \begin{pmatrix} A_{11} & A_{12} & \cdots & A_{1n} \\ A_{21} & A_{22} & \cdots & A_{2n} \\ \vdots & & \ddots & \vdots \\ A_{m1} & A_{m2} & \cdots & A_{mn} \end{pmatrix}$$

is called an $m \times n$ matrix with entries in $\mathbb{K}$.

- We also say that $\mathbf{A}$ is an m -by- n matrix, that $\mathbf{A}$ has size $m \times n$ and that $\mathbf{A}$ has m rows and n columns.
- The entry A_{ij} of $\mathbf{A}$ is said to have *row index* i where $1 \leq i \leq m$, *column index* j where $1 \leq j \leq n$ and will be referred to as the (i, j) -th entry of $\mathbf{A}$.
- A shorthand notation for (2.1) is $\mathbf{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$.
- For matrices $\mathbf{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ and $\mathbf{B} = (B_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ we write $\mathbf{A} = \mathbf{B}$, provided $A_{ij} = B_{ij}$ for all $1 \leq i \leq m$ and all $1 \leq j \leq n$.

Definition 2.2 (Set of matrices)

- The set of m -by- n matrices with entries in $\mathbb{K}$ will be denoted by $M_{m,n}(\mathbb{K})$.
- The elements of the set $M_{m,1}(\mathbb{K})$ are called *column vectors of length m* and the elements of the set $M_{1,n}(\mathbb{K})$ are called *row vectors of length n* .

- We will use the Latin alphabet for column vectors and decorate them with an arrow. For a column vector

$$\vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_m \end{pmatrix} \in M_{m,1}(\mathbb{K})$$

we also use the shorthand notation $\vec{x} = (x_i)_{1 \leq i \leq m}$ and we write $[\vec{x}]_i$ for the i -th entry of $\vec{x}$, so that $[\vec{x}]_i = x_i$ for all $1 \leq i \leq m$.

- We will use the Greek alphabet for row vectors and decorate them with an arrow. For a row vector

$$\vec{\xi} = (\xi_1 \quad \xi_2 \quad \cdots \quad \xi_n) \in M_{1,n}(\mathbb{K})$$

we also use the shorthand notation $\vec{\xi} = (\xi_i)_{1 \leq i \leq n}$ and we write $[\vec{\xi}]_i$ for the i -th entry of $\vec{\xi}$, so that $[\vec{\xi}]_i = \xi_i$ for all $1 \leq i \leq n$.

Remark 2.3 (Notation)

- A matrix is always denoted by a bold capital letter, such as **A**, **B**, **C**, **D**.
- The entries of the matrix are denoted by A_{ij} , B_{ij} , C_{ij} , D_{ij} , respectively.
- We may think of an $m \times n$ matrix as consisting of n column vectors of length m . The column vectors of the matrix are denoted by $\vec{a}_i$, $\vec{b}_i$, $\vec{c}_i$, $\vec{d}_i$, respectively.
- We may think of an $m \times n$ matrix as consisting of m row vectors of length n . The row vectors of the matrix are denoted by $\vec{\alpha}_i$, $\vec{\beta}_i$, $\vec{\gamma}_i$, $\vec{\delta}_i$, respectively.
- For a matrix **A** we also write $[\mathbf{A}]_{ij}$ for the (i, j) -th entry of **A**. So for $\mathbf{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$, we have $[\mathbf{A}]_{ij} = A_{ij}$ for all $1 \leq i \leq m, 1 \leq j \leq n$.

Example 2.4 For

$$\mathbf{A} = \begin{pmatrix} \pi & \sqrt{2} \\ -1 & 5/3 \\ \log 2 & 3 \end{pmatrix} \in M_{3,2}(\mathbb{R}),$$

we have for instance $[\mathbf{A}]_{32} = 3$, $[\mathbf{A}]_{12} = \sqrt{2}$, $[\mathbf{A}]_{21} = -1$ and

$$\vec{a}_1 = \begin{pmatrix} \pi \\ -1 \\ \log 2 \end{pmatrix}, \quad \vec{a}_2 = \begin{pmatrix} \sqrt{2} \\ 5/3 \\ 3 \end{pmatrix}, \quad \vec{\alpha}_2 = (-1 \quad 5/3), \quad \vec{\alpha}_3 = (\log 2 \quad 3).$$

Recall that for sets $\mathcal{X}$ and $\mathcal{Y}$ we write $\mathcal{X} \times \mathcal{Y}$ for the Cartesian product of $\mathcal{X}$ and $\mathcal{Y}$, defined as the set of ordered pairs (x, y) with $x \in \mathcal{X}$ and $y \in \mathcal{Y}$. Moreover, $\mathcal{X} \times \mathcal{X}$ is usually denoted as $\mathcal{X}^2$. Likewise, for a natural number $n \in \mathbb{N}$, we write $\mathcal{X}^n$ for the set of ordered lists consisting of n elements of $\mathcal{X}$. We will also refer to ordered lists consisting of n elements as n -tuples. The elements of $\mathcal{X}^n$ are denoted by $(x_1, x_2, \dots, x_n)$ with $x_i \in \mathcal{X}$ for all $1 \leq i \leq n$. In particular, for all $n \in \mathbb{N}$ we have a bijective map from $\mathbb{K}^n$ to $M_{n,1}(\mathbb{K})$ given by

$$(2.2) \quad (x_1, \dots, x_n) \mapsto \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}.$$

For this reason, we also write $\mathbb{K}^n$ for the set of column vectors of length n with entries in $\mathbb{K}$. The set of row vectors of length n with entries in $\mathbb{K}$ will be denoted by $\mathbb{K}_n$.

Definition 2.5 (Special matrices and vectors)

- The zero matrix $\mathbf{0}_{m,n}$ is the $m \times n$ matrix whose entries are all zero. We will also write $\mathbf{0}_n$ for the $n \times n$ -matrix whose entries are all zero.
- Matrices with equal number n of rows and columns are known as *square matrices*.
- An entry A_{ij} of a square matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ is said to be a *diagonal entry* if $i = j$ and an *off-diagonal entry* otherwise. A matrix whose off-diagonal entries are all zero is said to be *diagonal*.
- We write $\mathbf{1}_n$ for the diagonal $n \times n$ matrix whose diagonal entries are all equal to 1. Using the so-called *Kronecker delta* defined by the rule

$$\delta_{ij} = \begin{cases} 1 & i = j, \\ 0 & i \neq j, \end{cases}$$

we have $[\mathbf{1}_n]_{ij} = \delta_{ij}$ for all $1 \leq i, j \leq n$. The matrix $\mathbf{1}_n$ is called the *unit matrix* or *identity matrix* of size n .

- The *standard basis of $\mathbb{K}^n$* is the set $\{\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n\}$ consisting of the column vectors of the identity matrix $\mathbf{1}_n$ of size n .
- The *standard basis of $\mathbb{K}_n$* is the set $\{\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n\}$ consisting of the row vectors of the identity matrix $\mathbf{1}_n$ of size n .

Example 2.6

(i) Special matrices:

$$\mathbf{0}_{2,3} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad \mathbf{1}_2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}, \quad \mathbf{1}_3 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

(ii) The standard basis of $\mathbb{K}^3$ is $\{\vec{e}_1, \vec{e}_2, \vec{e}_3\}$, where

$$\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \vec{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \quad \text{and} \quad \vec{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

(iii) The standard basis of $\mathbb{K}_3$ is $\{\vec{e}_1, \vec{e}_2, \vec{e}_3\}$, where

$$\vec{e}_1 = (1 \ 0 \ 0), \quad \vec{e}_2 = (0 \ 1 \ 0) \quad \text{and} \quad \vec{e}_3 = (0 \ 0 \ 1).$$

2.2 Matrix operations

We can multiply a matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ with a *scalar* $s \in \mathbb{K}$. This amounts to multiplying each entry of $\mathbf{A}$ with s :

Definition 2.7 Scalar multiplication in $M_{m,n}(\mathbb{K})$ is the map

$$\cdot_{M_{m,n}(\mathbb{K})} : \mathbb{K} \times M_{m,n}(\mathbb{K}) \rightarrow M_{m,n}(\mathbb{K}), \quad (s, \mathbf{A}) \mapsto s \cdot_{M_{m,n}(\mathbb{K})} \mathbf{A}$$

defined by the rule

$$(2.3) \quad s \cdot_{M_{m,n}(\mathbb{K})} \mathbf{A} = (s \cdot_{\mathbb{K}} A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n} \in M_{m,n}(\mathbb{K}),$$

where $s \cdot_{\mathbb{K}} A_{ij}$ denotes the field multiplication of scalars $s, A_{ij} \in \mathbb{K}$.

Remark 2.8 Here we multiply with s from the left. Likewise, we define $\mathbf{A} \cdot_{M_{m,n}(\mathbb{K})} s = (A_{ij} \cdot_{\mathbb{K}} s)_{1 \leq i \leq m, 1 \leq j \leq n}$, that is, we multiply from the right. Of course, since multiplication of scalars is commutative, we have $s \cdot_{M_{m,n}(\mathbb{K})} \mathbf{A} = \mathbf{A} \cdot_{M_{m,n}(\mathbb{K})} s$, that is, left multiplication and right multiplication gives the same matrix. Be aware that this is not true in every number system. An example that you might encounter later on are the so-called *quaternions*, where multiplication fails to be commutative.

The sum of matrices $\mathbf{A}$ and $\mathbf{B}$ of identical size is defined as follows:

Definition 2.9 Addition in $M_{m,n}(\mathbb{K})$ is the map

$$+_{M_{m,n}(\mathbb{K})} : M_{m,n}(\mathbb{K}) \times M_{m,n}(\mathbb{K}) \rightarrow M_{m,n}(\mathbb{K}), \quad (\mathbf{A}, \mathbf{B}) \mapsto \mathbf{A} +_{M_{m,n}(\mathbb{K})} \mathbf{B}$$

defined by the rule

$$(2.4) \quad \mathbf{A} +_{M_{m,n}(\mathbb{K})} \mathbf{B} = (A_{ij} +_{\mathbb{K}} B_{ij})_{1 \leq i \leq m, 1 \leq j \leq n} \in M_{m,n}(\mathbb{K}),$$

where $A_{ij} +_{\mathbb{K}} B_{ij}$ denotes the field addition of scalars $A_{ij}, B_{ij} \in \mathbb{K}$.

Remark 2.10 (Abusing notation)

- Field addition takes two scalars and produces another scalar, thus it is a map $\mathbb{K} \times \mathbb{K} \rightarrow \mathbb{K}$, whereas addition of matrices is a map $M_{m,n}(\mathbb{K}) \times M_{m,n}(\mathbb{K}) \rightarrow M_{m,n}(\mathbb{K})$. For this reason we wrote $+_{M_{m,n}(\mathbb{K})}$ above in order to distinguish matrix addition from field addition of scalars. Of course, it is quite cumbersome to always write $+_{M_{m,n}(\mathbb{K})}$ and $+_{\mathbb{K}}$, so we follow the usual custom of writing $+$, both for field addition of scalars and for matrix addition, trusting that the reader is aware of the difference.
- Likewise, we simply write $\cdot$ instead of $\cdot_{M_{m,n}(\mathbb{K})}$ or omit the dot entirely, so that $s \cdot \mathbf{A} = s\mathbf{A} = s \cdot_{M_{m,n}(\mathbb{K})} \mathbf{A}$ for $s \in \mathbb{K}$ and $\mathbf{A} \in M_{m,n}(\mathbb{K})$.

Example 2.11

(i) Multiplication of a matrix by a scalar:

$$5 \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} 5 = \begin{pmatrix} 5 \cdot 1 & 5 \cdot 2 \\ 5 \cdot 3 & 5 \cdot 4 \end{pmatrix} = \begin{pmatrix} 5 & 10 \\ 15 & 20 \end{pmatrix}.$$

(ii) Addition of matrices:

$$\begin{pmatrix} 3 & -5 \\ -2 & 8 \end{pmatrix} + \begin{pmatrix} -3 & 8 \\ 7 & 10 \end{pmatrix} = \begin{pmatrix} 0 & 3 \\ 5 & 18 \end{pmatrix}.$$

If the number of columns of a matrix $\mathbf{A}$ is equal to the number of rows of a matrix $\mathbf{B}$, we define the matrix product $\mathbf{AB}$ of $\mathbf{A}$ and $\mathbf{B}$ as follows:

Definition 2.12 (Matrix multiplication — Video) Let $\mathbf{A} \in M_{m,n}(\mathbb{K})$ be an m -by- n matrix and $\mathbf{B} \in M_{n,\tilde{m}}(\mathbb{K})$ be an n -by- $\tilde{m}$ matrix. The *matrix product* of $\mathbf{A}$ and $\mathbf{B}$ is the

m -by- $\tilde{m}$ matrix $\mathbf{AB} \in M_{m,\tilde{m}}(\mathbb{K})$ whose entries are defined by the rule

$$[\mathbf{AB}]_{ik} = A_{i1}B_{1k} + A_{i2}B_{2k} + \cdots + A_{in}B_{nk} = \sum_{j=1}^n A_{ij}B_{jk} = \sum_{j=1}^n [\mathbf{A}]_{ij}[\mathbf{B}]_{jk}.$$

for all $1 \leq i \leq m$ and all $1 \leq k \leq \tilde{m}$.

Remark 2.13 (Pairing of row and column vectors) We may define a pairing $\mathbb{K}_n \times \mathbb{K}^n \rightarrow \mathbb{K}$ of a row vector of length n and a column vector of length n by the rule

$$(\vec{\xi}, \vec{x}) \mapsto \vec{\xi}\vec{x} = \xi_1x_1 + \xi_2x_2 + \cdots + \xi_nx_n$$

for all $\vec{\xi} = (\xi_i)_{1 \leq i \leq n} \in \mathbb{K}_n$ and for all $\vec{x} = (x_i)_{1 \leq i \leq n} \in \mathbb{K}^n$. So we multiply the first entry of $\vec{\xi}$ with the first entry of $\vec{x}$, add the product of the second entry of $\vec{\xi}$ and the second entry of $\vec{x}$ and continue in this fashion until the last entry of $\vec{\xi}$ and $\vec{x}$.

The (i, j) -th entry of the matrix product of $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and $\mathbf{B} \in M_{n,\tilde{m}}(\mathbb{K})$ is then given by the pairing

$$[\mathbf{AB}]_{ij} = \vec{\alpha}_i \vec{b}_j$$

of the i -th row vector $\vec{\alpha}_i$ of $\mathbf{A}$ and the j -th column vector $\vec{b}_j$ of $\mathbf{B}$.

Remark 2.14 (Matrix multiplication is not commutative — [Video](#)) If $\mathbf{A}$ is a m -by- n matrix and $\mathbf{B}$ a n -by- m matrix, then both $\mathbf{AB}$ and $\mathbf{BA}$ are defined, but in general $\mathbf{AB} \neq \mathbf{BA}$ since $\mathbf{AB}$ is an m -by- m matrix and $\mathbf{BA}$ is an n -by- n matrix. Even when $n = m$ so that both $\mathbf{A}$ and $\mathbf{B}$ are square matrices, it is false in general that $\mathbf{AB} = \mathbf{BA}$.

The matrix operations have the following properties:

Proposition 2.15 (Properties of matrix operations)

- $\mathbf{0}_{m,n} + \mathbf{A} = \mathbf{A}$ for all $\mathbf{A} \in M_{m,n}(\mathbb{K})$;
- $\mathbf{1}_m \mathbf{A} = \mathbf{A}$ and $\mathbf{A} \mathbf{1}_n = \mathbf{A}$ for all $\mathbf{A} \in M_{m,n}(\mathbb{K})$;
- $\mathbf{0}_{\tilde{m},m} \mathbf{A} = \mathbf{0}_{\tilde{m},n}$ and $\mathbf{A} \mathbf{0}_{n,\tilde{m}} = \mathbf{0}_{m,\tilde{m}}$ for all $\mathbf{A} \in M_{m,n}(\mathbb{K})$;
- $\mathbf{A} + \mathbf{B} = \mathbf{B} + \mathbf{A}$ and $(\mathbf{A} + \mathbf{B}) + \mathbf{C} = \mathbf{A} + (\mathbf{B} + \mathbf{C})$ for all $\mathbf{A}, \mathbf{B}, \mathbf{C} \in M_{m,n}(\mathbb{K})$;
- $\mathbf{0} \cdot \mathbf{A} = \mathbf{0}_{m,n}$ for all $\mathbf{A} \in M_{m,n}(\mathbb{K})$;
- $(s_1 s_2) \mathbf{A} = s_1 (s_2 \mathbf{A})$ for all $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and all $s_1, s_2 \in \mathbb{K}$;
- $\mathbf{A}(s\mathbf{B}) = s(\mathbf{AB}) = (s\mathbf{A})\mathbf{B}$ for all $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and all $\mathbf{B} \in M_{n,\tilde{m}}(\mathbb{K})$ and all $s \in \mathbb{K}$;
- $s(\mathbf{A} + \mathbf{B}) = s\mathbf{A} + s\mathbf{B}$ for all $\mathbf{A}, \mathbf{B} \in M_{m,n}(\mathbb{K})$ and $s \in \mathbb{K}$;
- $(s_1 + s_2) \mathbf{A} = s_1 \mathbf{A} + s_2 \mathbf{A}$ for all $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and for all $s_1, s_2 \in \mathbb{K}$;
- $(\mathbf{B} + \mathbf{C}) \mathbf{A} = \mathbf{BA} + \mathbf{CA}$ for all $\mathbf{B}, \mathbf{C} \in M_{\tilde{m},m}(\mathbb{K})$ and for all $\mathbf{A} \in M_{m,n}(\mathbb{K})$;
- $\mathbf{A}(\mathbf{B} + \mathbf{C}) = \mathbf{AB} + \mathbf{AC}$ for all $\mathbf{A} \in M_{m,\tilde{m}}(\mathbb{K})$ and for all $\mathbf{B}, \mathbf{C} \in M_{\tilde{m},n}(\mathbb{K})$.

Proof We only show the second and the last property. The proofs of the remaining ones are similar and/or elementary consequences of the properties of addition and multiplication of scalars.

To show the second property consider $\mathbf{A} \in M_{m,n}(\mathbb{K})$. Then, by definition, we have for all $1 \leq k \leq m$ and all $1 \leq j \leq n$

$$[\mathbf{1}_m \mathbf{A}]_{kj} = \sum_{i=1}^m [\mathbf{1}_m]_{ki} [\mathbf{A}]_{ij} = \sum_{i=1}^m \delta_{ki} A_{ij} = A_{kj} = [\mathbf{A}]_{kj},$$

where the second last equality uses that δ_{ki} is 0 unless $i = k$, in which case $\delta_{kk} = 1$. We conclude that $\mathbf{1}_m \mathbf{A} = \mathbf{A}$. Likewise, we obtain for all $1 \leq i \leq m$ and all $1 \leq k \leq n$

$$[\mathbf{A} \mathbf{1}_n]_{ik} = \sum_{j=1}^n [\mathbf{A}]_{ij} [\mathbf{1}_n]_{jk} = \sum_{j=1}^n A_{ij} \delta_{jk} = A_{ik} = [\mathbf{A}]_{ik}$$

so that $\mathbf{A} \mathbf{1}_n = \mathbf{A}$. The identities

$$\sum_{i=1}^m \delta_{ki} A_{ij} = A_{kj} \quad \text{and} \quad \sum_{j=1}^n A_{ij} \delta_{jk} = A_{ik}$$

are used repeatedly in Linear Algebra, so make sure you understand them.

For the last property, applying the definition of matrix multiplication gives

$$\mathbf{AB} = \left(\sum_{i=1}^m A_{ki} B_{ij} \right)_{1 \leq k \leq \tilde{m}, 1 \leq j \leq n} \quad \text{and} \quad \mathbf{AC} = \left(\sum_{i=1}^m A_{ki} C_{ij} \right)_{1 \leq k \leq \tilde{m}, 1 \leq j \leq n},$$

so that

$$\begin{aligned} \mathbf{AB} + \mathbf{AC} &= \left(\sum_{i=1}^m A_{ki} B_{ij} + \sum_{i=1}^m A_{ki} C_{ij} \right)_{1 \leq k \leq \tilde{m}, 1 \leq j \leq n} \\ &= \left(\sum_{i=1}^m A_{ki} (B_{ij} + C_{ij}) \right)_{1 \leq k \leq \tilde{m}, 1 \leq j \leq n} = \mathbf{A}(\mathbf{B} + \mathbf{C}), \end{aligned}$$

where we use that

$$\mathbf{B} + \mathbf{C} = (B_{ij} + C_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}.$$

□

Finally, we may flip a matrix along its “diagonal entries”, that is, we interchange the role of rows and columns. More precisely:

Definition 2.16 (Transpose of a matrix)

- The *transpose* of a matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ is the matrix $\mathbf{A}^T \in M_{n,m}(\mathbb{K})$ satisfying

$$[\mathbf{A}^T]_{ij} = [\mathbf{A}]_{ji}$$

for all $1 \leq i \leq n$ and $1 \leq j \leq m$.

- A square matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ that satisfies $\mathbf{A} = \mathbf{A}^T$ is called *symmetric*.
- A square matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ that satisfies $\mathbf{A} = -\mathbf{A}^T$ is called *anti-symmetric*.

Example 2.17 If

$$\mathbf{A} = \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix}, \quad \text{then} \quad \mathbf{A}^T = \begin{pmatrix} 1 & 3 & 5 \\ 2 & 4 & 6 \end{pmatrix}.$$

Remark 2.18 (Properties of the transpose)

- For $\mathbf{A} \in M_{m,n}(\mathbb{K})$ we have by definition $(\mathbf{A}^T)^T = \mathbf{A}$.
- For $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and $\mathbf{B} \in M_{n,\tilde{m}}(\mathbb{K})$, we have

$$(\mathbf{AB})^T = \mathbf{B}^T \mathbf{A}^T.$$

Indeed, by definition we have for all $1 \leq i \leq \tilde{m}$ and all $1 \leq j \leq m$

$$[(\mathbf{AB})^T]_{ij} = [\mathbf{AB}]_{ji} = \sum_{k=1}^n [\mathbf{A}]_{jk} [\mathbf{B}]_{ki} = \sum_{k=1}^n [\mathbf{B}^T]_{ik} [\mathbf{A}^T]_{kj} = [\mathbf{B}^T \mathbf{A}^T]_{ij}.$$

2.3 Mappings associated to matrices

Definition 2.19 (Mapping associated to a matrix) For an $(m \times n)$ -matrix $\mathbf{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n} \in M_{m,n}(\mathbb{K})$ with column vectors $\vec{a}_1, \dots, \vec{a}_n \in \mathbb{K}^m$ we define a mapping

$$f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m, \quad \vec{x} \mapsto \mathbf{A}\vec{x},$$

where the column vector $\mathbf{A}\vec{x} \in \mathbb{K}^m$ is obtained by matrix multiplication of the matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and the column vector $\vec{x} = (x_i)_{1 \leq i \leq n} \in \mathbb{K}^n$

$$\mathbf{A}\vec{x} = \vec{a}_1 x_1 + \vec{a}_2 x_2 + \dots + \vec{a}_n x_n = \begin{pmatrix} A_{11}x_1 + A_{12}x_2 + \dots + A_{1n}x_n \\ A_{21}x_1 + A_{22}x_2 + \dots + A_{2n}x_n \\ \vdots \\ A_{m1}x_1 + A_{m2}x_2 + \dots + A_{mn}x_n \end{pmatrix}.$$

Recall that if $f : \mathcal{X} \rightarrow \mathcal{Y}$ and $g : \mathcal{X} \rightarrow \mathcal{Y}$ are mappings from a set $\mathcal{X}$ into a set $\mathcal{Y}$, then we write $f = g$ if $f(x) = g(x)$ for all elements $x \in \mathcal{X}$.

The matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ uniquely determines the mapping $f_{\mathbf{A}}$:

Proposition 2.20 Let $\mathbf{A}, \mathbf{B} \in M_{m,n}(\mathbb{K})$. Then $f_{\mathbf{A}} = f_{\mathbf{B}}$ if and only if $\mathbf{A} = \mathbf{B}$.

Proof If $\mathbf{A} = \mathbf{B}$, then $A_{ij} = B_{ij}$ for all $1 \leq i \leq m, 1 \leq j \leq n$, hence we conclude that $f_{\mathbf{A}} = f_{\mathbf{B}}$. In order to show the converse direction we consider the standard basis $\vec{e}_i = (\delta_{ij})_{1 \leq j \leq n}, i = 1, \dots, n$ of $\mathbb{K}^n$. Now by assumption

$$f_{\mathbf{A}}(\vec{e}_i) = \begin{pmatrix} A_{1i} \\ A_{2i} \\ \vdots \\ A_{mi} \end{pmatrix} = f_{\mathbf{B}}(\vec{e}_i) = \begin{pmatrix} B_{1i} \\ B_{2i} \\ \vdots \\ B_{mi} \end{pmatrix}.$$

Since this holds for all $i = 1, \dots, n$, we conclude $A_{ij} = B_{ij}$ for all $j = 1, \dots, m$ and $i = 1, \dots, n$. Therefore, we have $\mathbf{A} = \mathbf{B}$, as claimed. $\square$

Recall that if $f : \mathcal{X} \rightarrow \mathcal{Y}$ is a mapping from a set $\mathcal{X}$ into a set $\mathcal{Y}$ and $g : \mathcal{Y} \rightarrow \mathcal{Z}$ a mapping from $\mathcal{Y}$ into a set $\mathcal{Z}$, we can consider the *composition* of g and f

$$g \circ f : \mathcal{X} \rightarrow \mathcal{Z}, \quad x \mapsto g(f(x)).$$

The motivation for the [Definition 2.12](#) of matrix multiplication is given by the following theorem which states that the mapping $f_{\mathbf{AB}}$ associated to the matrix product $\mathbf{AB}$ is the composition of the mapping $f_{\mathbf{A}}$ associated to the matrix $\mathbf{A}$ and the mapping $f_{\mathbf{B}}$ associated to the matrix $\mathbf{B}$. More precisely:

Theorem 2.21 *Let $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and $\mathbf{B} \in M_{n,\tilde{m}}(\mathbb{K})$ so that $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ and $f_{\mathbf{B}} : \mathbb{K}^{\tilde{m}} \rightarrow \mathbb{K}^n$ and $f_{\mathbf{AB}} : \mathbb{K}^{\tilde{m}} \rightarrow \mathbb{K}^m$. Then $f_{\mathbf{AB}} = f_{\mathbf{A}} \circ f_{\mathbf{B}}$.*

Proof For $\vec{x} = (x_k)_{1 \leq k \leq \tilde{m}} \in \mathbb{K}^{\tilde{m}}$ we write $\vec{y} = f_{\mathbf{B}}(\vec{x})$. Then, by definition, $\vec{y} = \mathbf{B}\vec{x} = (y_j)_{1 \leq j \leq n}$ where

$$(2.5) \quad y_j = B_{j1}x_1 + B_{j2}x_2 + \cdots + B_{j\tilde{m}}x_{\tilde{m}} = \sum_{k=1}^{\tilde{m}} B_{jk}x_k.$$

Hence writing $\vec{z} = f_{\mathbf{A}}(\vec{y}) = \mathbf{A}\vec{y}$, we have $\vec{z} = (z_i)_{1 \leq i \leq m}$, where

$$\begin{aligned} z_i &= A_{i1}y_1 + A_{i2}y_2 + \cdots + A_{in}y_n = \sum_{j=1}^n A_{ij}y_j = \sum_{j=1}^n A_{ij} \sum_{k=1}^{\tilde{m}} B_{jk}x_k \\ &= \sum_{k=1}^{\tilde{m}} \left(\sum_{j=1}^n A_{ij}B_{jk} \right) x_k \end{aligned}$$

and where we have used (2.5). Since $\mathbf{AB} = (C_{ik})_{1 \leq i \leq m, 1 \leq k \leq \tilde{m}}$ with

$$C_{ik} = \sum_{j=1}^n A_{ij}B_{jk},$$

we conclude that $\vec{z} = f_{\mathbf{AB}}(\vec{x})$, as claimed. □

Combining [Theorem 2.21](#) and [Proposition 2.20](#), we also obtain:

Corollary 2.22 *Let $\mathbf{A} \in M_{m,n}(\mathbb{K})$, $\mathbf{B} \in M_{n,\tilde{m}}(\mathbb{K})$ and $\mathbf{C} \in M_{\tilde{m},\tilde{n}}(\mathbb{K})$. Then*

$$(\mathbf{AB})\mathbf{C} = \mathbf{A}(\mathbf{BC}),$$

that is, the matrix product is associative.

Proof Using [Proposition 2.20](#) it is enough to show that

$$f_{\mathbf{AB}} \circ f_{\mathbf{C}} = f_{\mathbf{A}} \circ f_{\mathbf{BC}}.$$

Using [Theorem 2.21](#), we get for all $\vec{x} \in \mathbb{K}^{\tilde{n}}$

$$(f_{\mathbf{AB}} \circ f_{\mathbf{C}})(\vec{x}) = f_{\mathbf{AB}}(f_{\mathbf{C}}(\vec{x})) = f_{\mathbf{A}}(f_{\mathbf{B}}(f_{\mathbf{C}}(\vec{x}))) = f_{\mathbf{A}}(f_{\mathbf{BC}}(\vec{x})) = (f_{\mathbf{A}} \circ f_{\mathbf{BC}})(\vec{x}).$$

□

Remark 2.23 For all $\mathbf{A} \in M_{m,n}(\mathbb{K})$, the mapping $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ satisfies the following two very important properties

$$(2.6) \quad \begin{aligned} f_{\mathbf{A}}(\vec{x} + \vec{y}) &= f_{\mathbf{A}}(\vec{x}) + f_{\mathbf{A}}(\vec{y}), & (\text{additivity}), \\ f_{\mathbf{A}}(s \cdot \vec{x}) &= s \cdot f_{\mathbf{A}}(\vec{x}), & (1\text{-homogeneity}), \end{aligned}$$

for all $\vec{x}, \vec{y} \in \mathbb{K}^n$ and $s \in \mathbb{K}$. Indeed, using [Proposition 2.15](#) we have

$$f_{\mathbf{A}}(\vec{x} + \vec{y}) = \mathbf{A}(\vec{x} + \vec{y}) = \mathbf{A}\vec{x} + \mathbf{A}\vec{y} = f_{\mathbf{A}}(\vec{x}) + f_{\mathbf{A}}(\vec{y})$$

and

$$f_{\mathbf{A}}(s \cdot \vec{x}) = \mathbf{A}(s\vec{x}) = s \cdot (\mathbf{A}\vec{x}) = s \cdot f_{\mathbf{A}}(\vec{x}).$$

Mappings satisfying (2.6) are called *linear*.

Example 2.24 Notice that “most” functions $\mathbb{R} \rightarrow \mathbb{R}$ are neither additive nor 1-homogeneous. As an example, consider a mapping $f : \mathbb{R} \rightarrow \mathbb{R}$ which satisfies the 1-homogeneity property. Let $a = f(1) \in \mathbb{R}$. Then the 1-homogeneity implies that for all $x \in \mathbb{R} = \mathbb{R}^1$ we have

$$f(x) = f(x \cdot 1) = x \cdot f(1) = a \cdot x,$$

showing that the only 1-homogeneous mappings from $\mathbb{R} \rightarrow \mathbb{R}$ are of the form $x \mapsto ax$, where a is a real number. In particular, $\sin$, $\cos$, $\tan$, $\log$, $\exp$, $\sqrt{}$ and all polynomials of degree higher than one are not linear.

Vector spaces and linear maps

3.1 Vector spaces

WEEK 3

We have seen that to every matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ we can associate a mapping $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ which is additive and 1-homogeneous. Another example of a mapping which is additive and 1-homogeneous is the derivative. Consider $P(\mathbb{R})$, the set of polynomial functions in one real variable, which we denote by x , with real coefficients. That is, an element $p \in P(\mathbb{R})$ is a function

$$p : \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto a_n x^n + a_{n-1} x^{n-1} + \cdots + a_1 x + a_0 = \sum_{k=0}^n a_k x^k,$$

where $n \in \mathbb{N}$ and the coefficients $a_k \in \mathbb{R}$ for $k = 0, 1, \dots, n$. The largest $m \in \mathbb{N} \cup \{0\}$ such that $a_m \neq 0$ is called the *degree* of p . Notice that we consider polynomials of arbitrary, but *finite degree*. A *power series* $x \mapsto \sum_{k=0}^{\infty} a_k x^k$, that you encounter in the Analysis module, is not a polynomial, unless only finitely many of its coefficients are different from zero.

Clearly, we can multiply p with a real number $s \in \mathbb{R}$ to obtain a new polynomial $s \cdot_{P(\mathbb{R})} p$

$$(3.1) \quad s \cdot_{P(\mathbb{R})} p : \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto s \cdot p(x)$$

so that $(s \cdot_{P(\mathbb{R})} p)(x) = \sum_{k=0}^n s a_k x^k$ for all $x \in \mathbb{R}$. Here $s \cdot p(x)$ is the usual multiplication of the real numbers s and $p(x)$. If we consider another polynomial

$$q : \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto \sum_{k=0}^n b_k x^k$$

with $b_k \in \mathbb{R}$ for $k = 0, 1, \dots, n$, the sum of the polynomials p and q is the polynomial

$$(3.2) \quad p +_{P(\mathbb{R})} q : \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto p(x) + q(x)$$

so that $(p +_{P(\mathbb{R})} q)(x) = \sum_{k=0}^n (a_k + b_k) x^k$ for all $x \in \mathbb{R}$. Here $p(x) + q(x)$ is the usual addition of the real numbers $p(x)$ and $q(x)$. We will henceforth omit writing $+_{P(\mathbb{R})}$ and $\cdot_{P(\mathbb{R})}$ and simply write $+$ and $\cdot$.

We may think of the derivative with respect to the variable x as a mapping

$$\frac{d}{dx} : P(\mathbb{R}) \rightarrow P(\mathbb{R}).$$

Now recall that the derivative satisfies

$$(3.3) \quad \begin{aligned} \frac{d}{dx}(p + q) &= \frac{d}{dx}(p) + \frac{d}{dx}(q) && \text{(additivity),} \\ \frac{d}{dx}(s \cdot p) &= s \cdot \frac{d}{dx}(p) && \text{(1-homogeneity).} \end{aligned}$$

Comparing (2.6) with (3.3) we notice that the polynomials p, q take the role of the vectors $\vec{x}, \vec{y}$ and the derivative takes the role of the mapping $f_{\mathbf{A}}$. This suggests that the mental image of a vector being an arrow in $\mathbb{K}^n$ is too narrow and that we ought to come up with a generalisation of the space $\mathbb{K}^n$ whose elements are *abstract vectors*.

Video Vector spaces

In order to define the notion of a space of abstract vectors, we may ask what key structure the set of (column) vectors $\mathbb{K}^n$ carries. On $\mathbb{K}^n$, we have two fundamental operations,

$$\begin{aligned} + : \mathbb{K}^n \times \mathbb{K}^n &\rightarrow \mathbb{K}^n & (\vec{x}, \vec{y}) &\mapsto \vec{x} + \vec{y}, & (\text{vector addition}), \\ \cdot : \mathbb{K} \times \mathbb{K}^n &\rightarrow \mathbb{K}^n, & (s, \vec{x}) &\mapsto s \cdot \vec{x}, & (\text{scalar multiplication}). \end{aligned}$$

A *vector space* is roughly speaking a set where these two operations are defined and obey the expected properties. More precisely:

Definition 3.1 (Vector space) A $\mathbb{K}$ -vector space, or *vector space over $\mathbb{K}$* is a set V with a distinguished element 0_V (called the zero vector) and two operations

$$+_V : V \times V \rightarrow V \quad (v_1, v_2) \mapsto v_1 +_V v_2 \quad (\text{vector addition})$$

and

$$\cdot_V : \mathbb{K} \times V \rightarrow V \quad (s, v) \mapsto s \cdot_V v \quad (\text{scalar multiplication}),$$

so that the following properties hold:

- Commutativity of vector addition

$$v_1 +_V v_2 = v_2 +_V v_1 \quad (\text{for all } v_1, v_2 \in V);$$

- Associativity of vector addition

$$v_1 +_V (v_2 +_V v_3) = (v_1 +_V v_2) +_V v_3 \quad (\text{for all } v_1, v_2, v_3 \in V);$$

- Identity element of vector addition

$$(3.4) \quad 0_V +_V v = v +_V 0_V = v \quad (\text{for all } v \in V);$$

- Identity element of scalar multiplication

$$1 \cdot_V v = v \quad (\text{for all } v \in V);$$

- Scalar multiplication by zero

$$(3.5) \quad 0 \cdot_V v = 0_V \quad (\text{for all } v \in V);$$

- Compatibility of scalar multiplication with field multiplication

$$(s_1 s_2) \cdot_V v = s_1 \cdot_V (s_2 \cdot_V v) \quad (\text{for all } s_1, s_2 \in \mathbb{K}, v \in V);$$

- Distributivity of scalar multiplication with respect to vector addition

$$s \cdot_V (v_1 +_V v_2) = s \cdot_V v_1 +_V s \cdot_V v_2 \quad (\text{for all } s \in \mathbb{K}, v_1, v_2 \in V);$$

- Distributivity of scalar multiplication with respect to field addition

$$(s_1 + s_2) \cdot_V v = s_1 \cdot_V v +_V s_2 \cdot_V v \quad (\text{for all } s_1, s_2 \in \mathbb{K}, v \in V).$$

The elements of V are called *vectors*.

Example 3.2 (Field) A field $\mathbb{K}$ is a $\mathbb{K}$ -vector space. We may take $V = \mathbb{K}$, $0_V = 0_{\mathbb{K}}$ and equip V with addition $+_V = +_{\mathbb{K}}$ and scalar multiplication $\cdot_V = \cdot_{\mathbb{K}}$. Then the properties of a field imply that $V = \mathbb{K}$ is a $\mathbb{K}$ -vector space.

Example 3.3 (Vector space of matrices) Let $V = M_{m,n}(\mathbb{K})$ denote the set of $m \times n$ -matrices with entries in $\mathbb{K}$ and $0_V = \mathbf{0}_{m,n}$ denote the zero vector. It follows from [Proposition 2.15](#) that V equipped with addition $+_V : V \times V \rightarrow V$ defined by (2.4) and scalar multiplication $\cdot_V : \mathbb{K} \times V \rightarrow V$ defined by (2.3) is a $\mathbb{K}$ -vector

space. In particular, the set of column vectors $\mathbb{K}^n = M_{n,1}(\mathbb{K})$ is a $\mathbb{K}$ -vector space as well.

Example 3.4 (Vector space of polynomials) The set $P(\mathbb{R})$ of polynomials in one real variable and with real coefficients is an $\mathbb{R}$ -vector space, when equipped with addition and scalar multiplication as defined in (3.1) and (3.2) and when the zero vector $0_{P(\mathbb{R})}$ is defined to be the *zero polynomial* $o : \mathbb{R} \rightarrow \mathbb{R}$, that is, the polynomial satisfying $o(x) = 0$ for all $x \in \mathbb{R}$.

More generally, functions form a vector space:

Example 3.5 (Vector space of functions) We follow the convention of calling a mapping with values in $\mathbb{K}$ a *function*. Let $I \subset \mathbb{R}$ be an interval and let $o : I \rightarrow \mathbb{K}$ denote the *zero function* defined by $o(x) = 0$ for all $x \in I$. We consider $V = F(I, \mathbb{K})$, the set of functions from I to $\mathbb{K}$ with zero vector $0_V = o$ given by the zero function and define addition $+_V : V \times V \rightarrow V$ as in (3.2) and scalar multiplication $\cdot_V : \mathbb{K} \times V \rightarrow V$ as in (3.1). It now is a consequence of the properties of addition and multiplication of scalars that $F(I, \mathbb{K})$ is a $\mathbb{K}$ -vector space. (The reader is invited to check this assertion!)

Example 3.6 (Vector space of sequences) A mapping $x : \mathbb{N} \rightarrow \mathbb{K}$ from the natural numbers into a field $\mathbb{K}$ called a *sequence* in $\mathbb{K}$ (or simply a *sequence*, when $\mathbb{K}$ is clear from the context). It is common to write x_n instead of $x(n)$ for $n \in \mathbb{N}$ and to denote a sequence by $(x_n)_{n \in \mathbb{N}} = (x_1, x_2, x_3, \dots)$. We write $\mathbb{K}^\infty$ for the set of sequences in $\mathbb{K}$. For instance, taking $\mathbb{K} = \mathbb{R}$, we may consider the sequence

$$\left(\frac{1}{n}\right)_{n \in \mathbb{N}} = \left(1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{1}{5}, \dots\right)$$

or the sequence

$$(\sqrt{n})_{n \in \mathbb{N}} = (1, \sqrt{2}, \sqrt{3}, 2, \sqrt{5}, \dots).$$

If we equip $\mathbb{K}^\infty$ with the zero vector given by the zero sequence $(0, 0, 0, 0, 0, \dots)$, addition given by $(x_n)_{n \in \mathbb{N}} + (y_n)_{n \in \mathbb{N}} = (x_n + y_n)_{n \in \mathbb{N}}$ and scalar multiplication given by $s \cdot (x_n)_{n \in \mathbb{N}} = (sx_n)_{n \in \mathbb{N}}$ for $s \in \mathbb{K}$, then $\mathbb{K}^\infty$ is a $\mathbb{K}$ -vector space.

Example 3.7 (Zero vector space) Consider a set $V = \{x\}$ consisting of a single element. We define $0_V = x$, addition by $x +_V x = x$ and scalar multiplication by $s \cdot_V x = x$. Then all the properties of Definition 3.1 are satisfied. We write $V = \{0_V\}$ or simply $V = \{0\}$ and call V the *zero vector space* (over $\mathbb{K}$).

The notion of a vector space is an example of an *abstract space*. Later in your studies you will encounter further examples, like *topological spaces*, *metric spaces* and *manifolds*.

Remark 3.8 (Notation & Definition) Let V be a $\mathbb{K}$ -vector space.

- For $v \in V$ we write $-v = (-1) \cdot_V v$ and for $v_1, v_2 \in V$ we write $v_1 - v_2 = v_1 +_V (-v_2)$. In particular, using the properties from [Definition 3.1](#) we have (check which properties we do use!)

$$v - v = v +_V (-v) = v +_V (-1) \cdot_V v = (1 - 1) \cdot_V v = 0 \cdot_V v = 0_V$$

For this reason we call $-v$ the *additive inverse* of v .

- Again, it is too cumbersome to always write $+_V$, for this reason we often write $v_1 + v_2$ instead of $v_1 +_V v_2$.
- Likewise, we will often write $s \cdot v$ or sv instead of $s \cdot_V v$.
- It is also customary to write 0 instead of 0_V .

Lemma 3.9 (Elementary properties of vector spaces) *Let V be a $\mathbb{K}$ -vector space. Then we have:*

- (i) *The zero vector is unique, that is, if $0'_V$ is another vector such that $0'_V + v = v + 0'_V = v$ for all $v \in V$, then $0'_V = 0_V$.*
- (ii) *The additive inverse of every $v \in V$ is unique, that is, if $w \in V$ satisfies $v + w = 0_V$, then $w = -v$.*
- (iii) *For all $s \in \mathbb{K}$ we have $s0_V = 0_V$.*
- (iv) *For $s \in \mathbb{K}$ and $v \in V$ we have $sv = 0_V$ if and only if either $s = 0$ or $v = 0_V$.*

Proof (The reader is invited to check which property of [Definition 3.1](#) is used in each of the equality signs below)

- (i) We have $0'_V = 0'_V + 0_V = 0_V$.
- (ii) Since $v + w = 0_V$, adding $-v$, we obtain $(-v) + v + w = 0_V + (-v) = -v = w$.
- (iii) We compute $s0_V = s(0_V + 0_V) = s0_V + s0_V$ so that $s0_V - s0_V = 0_V = s0_V$.
- (iv) $\Leftarrow$ If $v = 0_V$, then $sv = 0_V$ by (iii). If $s = 0$, then $sv = 0_V$ by (3.5).
 $\Rightarrow$ Let $s \in \mathbb{K}$ and $v \in V$ such that $sv = 0_V$. It is sufficient to show that if $s \neq 0$, then $v = 0_V$. Since $s \neq 0$ we can multiply $sv = 0_V$ with $1/s$ so that

$$\frac{1}{s}(sv) = \left(\frac{1}{s}s\right)v = v = \frac{1}{s}0_V = 0_V.$$

□

3.2 Linear maps

Throughout this section, V, W denote $\mathbb{K}$ -vector spaces.

Previously we saw that the mapping $f_A : \mathbb{K}^n \rightarrow \mathbb{K}^m$ associated to a matrix $M_{m,n}(\mathbb{K})$ is additive and 1-homogeneous. These notions also make sense for mappings between vector spaces.

Definition 3.10 (Linear map) A mapping $f : V \rightarrow W$ is called *linear* if it is additive and 1-homogeneous, that is, if it satisfies

$$(3.6) \quad f(s_1 v_1 + s_2 v_2) = s_1 f(v_1) + s_2 f(v_2)$$

for all $s_1, s_2 \in \mathbb{K}$ and for all $v_1, v_2 \in V$.

The reader is invited to check that the condition (3.6) is indeed equivalent to f being additive and 1-homogeneous.

Example 3.11 As we have seen in Remark 2.23, the mapping $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ associated to a matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ is linear. In Lemma 3.18 below we will see that in fact any linear map $\mathbb{K}^n \rightarrow \mathbb{K}^m$ is of this form.

Example 3.12 The derivative $\frac{d}{dx} : P(\mathbb{R}) \rightarrow P(\mathbb{R})$ is linear, see (3.3).

Example 3.13 The matrix transpose is a map $M_{m,n}(\mathbb{K}) \rightarrow M_{n,m}(\mathbb{K})$ and this map is linear. Indeed, for all $s, t \in \mathbb{K}$ and $\mathbf{A}, \mathbf{B} \in M_{m,n}(\mathbb{K})$, we have

$$\begin{aligned} (s\mathbf{A} + t\mathbf{B})^T &= (sA_{ji} + tB_{ji})_{1 \leq j \leq n, 1 \leq i \leq m} = s(A_{ji})_{1 \leq j \leq n, 1 \leq i \leq m} + \\ &\quad t(B_{ji})_{1 \leq j \leq n, 1 \leq i \leq m} = s\mathbf{A}^T + t\mathbf{B}^T. \end{aligned}$$

Example 3.14 If $\mathcal{X}$ is set, the mapping $\text{Id}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X}$ which returns its input is called the *identity mapping*. Let V be a $\mathbb{K}$ -vector space and $\text{Id}_V : V \rightarrow V$ the identity mapping so that $\text{Id}_V(v) = v$ for all $v \in V$. The identity mapping is linear since for all $s_1, s_2 \in \mathbb{K}$ and $v_1, v_2 \in V$ we have

$$\text{Id}_V(s_1 v_1 + s_2 v_2) = s_1 v_1 + s_2 v_2 = s_1 \text{Id}_V(v_1) + s_2 \text{Id}_V(v_2).$$

A necessary condition for linearity of a mapping is that it maps the zero vector onto the zero vector:

Lemma 3.15 Let $f : V \rightarrow W$ be a linear map, then $f(0_V) = 0_W$.

Proof Since $f : V \rightarrow W$ is linear, we have

$$f(0_V) = f(0 \cdot 0_V) = 0 \cdot f(0_V) = 0_W.$$

□

Proposition 3.16 Let V_1, V_2, V_3 be $\mathbb{K}$ -vector spaces and $f : V_1 \rightarrow V_2$ and $g : V_2 \rightarrow V_3$ be linear maps. Then the composition $g \circ f : V_1 \rightarrow V_3$ is linear. Furthermore, if $f : V_1 \rightarrow V_2$ is bijective, then the inverse function $f^{-1} : V_2 \rightarrow V_1$ (satisfying $f^{-1} \circ f = f \circ f^{-1} = \text{Id}_{V_1}$) is linear.

Proof Let $s, t \in \mathbb{K}$ and $v, w \in V_1$. Then

$$\begin{aligned} (g \circ f)(sv + tw) &= g(f(sv + tw)) = g(sf(v) + tf(w)) \\ &= sg(f(v)) + tg(f(w)) = s(g \circ f)(v) + t(g \circ f)(w), \end{aligned}$$

where we first use the linearity of f and then the linearity of g . It follows that $g \circ f$ is linear.

Now suppose $f : V_1 \rightarrow V_2$ is bijective with inverse function $f^{-1} : V_2 \rightarrow V_1$. Let $s, t \in \mathbb{K}$ and $v, w \in V_2$. Since f is bijective there exist unique vectors $v', w' \in V_1$ with $f(v') = v$ and $f(w') = w$. Hence we can write

$$\begin{aligned} f^{-1}(sv + tw) &= f^{-1}(sf(v') + tf(w')) = f^{-1}(f(sv' + tw')) \\ &= (f^{-1} \circ f)(sv' + tw') = sv' + tw', \end{aligned}$$

where we use the linearity of f . Since we also have $v' = f^{-1}(v)$ and $w' = f^{-1}(w)$, we obtain

$$f^{-1}(sv + tw) = sf^{-1}(v) + tf^{-1}(w),$$

thus showing that $f^{-1} : V_2 \rightarrow V_1$ is linear. $\square$

We also have:

Proposition 3.17 Let $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ the associated linear map. Then $f_{\mathbf{A}}$ is bijective if and only if there exists a matrix $\mathbf{B} \in M_{n,m}(\mathbb{K})$ satisfying $\mathbf{B}\mathbf{A} = \mathbf{1}_n$ and $\mathbf{A}\mathbf{B} = \mathbf{1}_m$. In this case, the matrix $\mathbf{B}$ is unique and will be denoted by $\mathbf{A}^{-1}$. We refer to $\mathbf{A}^{-1}$ as the inverse of $\mathbf{A}$ and call $\mathbf{A}$ invertible.

In order to prove [Proposition 3.17](#) we need the following lemma:

Lemma 3.18 A mapping $g : \mathbb{K}^m \rightarrow \mathbb{K}^n$ is linear if and only if there exists a matrix $\mathbf{B} \in M_{n,m}(\mathbb{K})$ so that $g = f_{\mathbf{B}}$.

Proof Let $\mathbf{B} \in M_{n,m}(\mathbb{K})$, then $f_{\mathbf{B}}$ is linear by [Remark 2.23](#). Conversely, let $g : \mathbb{K}^m \rightarrow \mathbb{K}^n$ be linear. Let $\{\vec{e}_1, \dots, \vec{e}_m\}$ denote the standard basis of $\mathbb{K}^m$. Write

$$g(\vec{e}_i) = \begin{pmatrix} B_{1i} \\ \vdots \\ B_{ni} \end{pmatrix} \quad \text{for } i = 1, \dots, m$$

and consider the matrix

$$\mathbf{B} = \begin{pmatrix} B_{11} & \cdots & B_{1m} \\ \vdots & \ddots & \vdots \\ B_{n1} & \cdots & B_{nm} \end{pmatrix} \in M_{n,m}(\mathbb{K}).$$

For $i = 1, \dots, m$ we obtain

$$(3.7) \quad f_{\mathbf{B}}(\vec{e}_i) = \mathbf{B}\vec{e}_i = g(\vec{e}_i).$$

Any vector $\vec{v} = (v_i)_{1 \leq i \leq m} \in \mathbb{K}^m$ can be written as

$$\vec{v} = v_1\vec{e}_1 + \cdots + v_m\vec{e}_m$$

for (unique) scalars $v_i, i = 1, \dots, m$. Hence using the linearity of g and $f_{\mathbf{B}}$, we compute

$$\begin{aligned} g(\vec{v}) - f_{\mathbf{B}}(\vec{v}) &= g(v_1\vec{e}_1 + \cdots + v_m\vec{e}_m) - f_{\mathbf{B}}(v_1\vec{e}_1 + \cdots + v_m\vec{e}_m) \\ &= v_1(g(\vec{e}_1) - f_{\mathbf{B}}(\vec{e}_1)) + \cdots + v_m(g(\vec{e}_m) - f_{\mathbf{B}}(\vec{e}_m)) = 0_{\mathbb{K}^n}, \end{aligned}$$

where the last equality uses (3.7). Since the vector $\vec{v}$ is arbitrary, it follows that $g = f_{\mathbf{B}}$, as claimed. $\square$

Proof of Proposition 3.17 First, notice that the mapping $f_{\mathbf{1}_n} : \mathbb{K}^n \rightarrow \mathbb{K}^n$ associated to the unit matrix is the identity mapping on $\mathbb{K}^n$, that is, for all $n \in \mathbb{N}$, we have $f_{\mathbf{1}_n} = \text{Id}_{\mathbb{K}^n}$.

Let $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and suppose that $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ is bijective with inverse function $(f_{\mathbf{A}})^{-1} : \mathbb{K}^m \rightarrow \mathbb{K}^n$. By [Proposition 3.16](#), the mapping $(f_{\mathbf{A}})^{-1}$ is linear and hence of

the form $(f_{\mathbf{A}})^{-1} = f_{\mathbf{B}}$ for some matrix $\mathbf{B} \in M_{n,m}(\mathbb{K})$ by the previous [Lemma 3.18](#). Using [Theorem 2.21](#), we obtain

$$(f_{\mathbf{A}})^{-1} \circ f_{\mathbf{A}} = \text{Id}_{\mathbb{K}^n} = f_{\mathbf{B}} \circ f_{\mathbf{A}} = f_{\mathbf{BA}} = f_{\mathbf{1}_n}$$

hence [Proposition 2.20](#) implies that $\mathbf{BA} = \mathbf{1}_n$. Likewise we have

$$f_{\mathbf{A}} \circ (f_{\mathbf{A}})^{-1} = \text{Id}_{\mathbb{K}^m} = f_{\mathbf{A}} \circ f_{\mathbf{B}} = f_{\mathbf{AB}} = f_{\mathbf{1}_m}$$

so that $\mathbf{AB} = \mathbf{1}_m$.

Conversely, let $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and suppose the matrix $\mathbf{B} \in M_{n,m}(\mathbb{K})$ satisfies $\mathbf{AB} = \mathbf{1}_m$ and $\mathbf{BA} = \mathbf{1}_n$. Then, as before, we have

$$f_{\mathbf{AB}} = f_{\mathbf{1}_m} = \text{Id}_{\mathbb{K}^m} = f_{\mathbf{A}} \circ f_{\mathbf{B}} \quad \text{and} \quad f_{\mathbf{BA}} = f_{\mathbf{1}_n} = \text{Id}_{\mathbb{K}^n} = f_{\mathbf{B}} \circ f_{\mathbf{A}}$$

showing that $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ is bijective with inverse function $f_{\mathbf{B}} : \mathbb{K}^m \rightarrow \mathbb{K}^n$.

Finally, to verify the uniqueness of $\mathbf{B}$, we assume that there exists $\mathbf{B}' \in M_{n,m}(\mathbb{K})$ with $\mathbf{AB}' = \mathbf{1}_m$ and $\mathbf{B}'\mathbf{A} = \mathbf{1}_n$. Then

$$\mathbf{B}' = \mathbf{B}'\mathbf{1}_m = \mathbf{B}'\mathbf{AB} = (\mathbf{B}'\mathbf{A})\mathbf{B} = \mathbf{1}_n\mathbf{B} = \mathbf{B},$$

showing that $\mathbf{B}' = \mathbf{B}$, hence $\mathbf{B}$ is unique. □

Exercises

Exercise 3.19 Let $f : V \rightarrow W$ be a linear map, $k \geq 2$ a natural number and $s_1, \dots, s_k \in \mathbb{K}$ and $v_1, \dots, v_k \in V$. Then $f : V \rightarrow W$ satisfies

$$f(s_1 v_1 + \dots + s_k v_k) = s_1 f(v_1) + \dots + s_k f(v_k)$$

or written with the sum symbol

$$f\left(\sum_{i=1}^k s_i v_i\right) = \sum_{i=1}^k s_i f(v_i).$$

This identity is used frequently in Linear Algebra, so make sure you understand it.

Exercise 3.20 Let $a, b, c, d \in \mathbb{K}$ and

$$\mathbf{A} = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in M_{2,2}(\mathbb{K}).$$

Show that $\mathbf{A}$ has an inverse $\mathbf{A}^{-1}$ if and only if $ad - bc \neq 0$. For $ad - bc \neq 0$, compute the inverse $\mathbf{A}^{-1}$.

3.3 Vector subspaces and isomorphisms

3.3.1 Vector subspaces

A vector subspace of a vector space is a subset that is itself a vector space, more precisely:

Definition 3.21 (Vector subspace) Let V be a $\mathbb{K}$ -vector space. A subset $U \subset V$ is called a *vector subspace* of V if U is non-empty and if

$$(3.8) \quad s_1 \cdot_V v_1 +_V s_2 \cdot_V v_2 \in U \quad \text{for all } s_1, s_2 \in \mathbb{K} \text{ and all } v_1, v_2 \in U.$$

Video Subspaces

Remark 3.22

- (i) Observe that since U is non-empty, it contains an element, say u . Since $0 \cdot_V u = 0_V \in U$ it follows that the zero vector 0_V lies in U . A vector subspace U is itself a vector space when we take $0_U = 0_V$ and borrow vector addition and scalar multiplication from V . Indeed, all of the properties in [Definition 3.1](#) of $+_V$ and $\cdot_V$ hold for all elements of V and all scalars, hence also for all elements of $U \subset V$ and all scalars. We only need to verify that we cannot fall out of U by vector addition and scalar multiplication, but this is precisely what the condition (3.8) states.
- (ii) A vector subspace is also called a *linear subspace* or simply a subspace.

The prototypical example of a vector subspace are lines and planes through the origin in $\mathbb{R}^3$:

Example 3.23 (Lines through the origin) Let $\vec{w} \neq 0_{\mathbb{R}^3}$, then the line

$$U = \{s\vec{w} \mid s \in \mathbb{R}\} \subset \mathbb{R}^3$$

is a vector subspace. Indeed, taking $s = 0$ it follows that $0_{\mathbb{R}^3} \in U$ so that U is non-empty. Let $\vec{u}_1, \vec{u}_2$ be vectors in U so that $\vec{u}_1 = t_1 \vec{w}$ and $\vec{u}_2 = t_2 \vec{w}$ for scalars $t_1, t_2 \in \mathbb{R}$. Let $s_1, s_2 \in \mathbb{R}$, then

$$s_1 \vec{u}_1 + s_2 \vec{u}_2 = s_1 t_1 \vec{w} + s_2 t_2 \vec{w} = (s_1 t_1 + s_2 t_2) \vec{w} \in U$$

so that $U \subset \mathbb{R}^3$ is a subspace.

Example 3.24 (Zero vector space) Let V be a $\mathbb{K}$ -vector space and $U = \{0_V\}$ the zero vector space arising from 0_V . Then, by [Definition 3.21](#) and the properties of [Definition 3.1](#), it follows that U is a vector subspace of V .

Example 3.25 (Periodic functions) Taking $I = \mathbb{R}$ and $\mathbb{K} = \mathbb{R}$ in [Example 3.5](#), we see that the functions $f : \mathbb{R} \rightarrow \mathbb{R}$ form an $\mathbb{R}$ -vector space $V = F(\mathbb{R}, \mathbb{R})$. Consider the subset

$$U = \{f \in F(\mathbb{R}, \mathbb{R}) \mid f \text{ is periodic with period } 2\pi\}$$

consisting of 2π -periodic functions, that is, an element $f \in U$ satisfies $f(x + 2\pi) = f(x)$ for all $x \in \mathbb{R}$. Notice that U is not empty, as $\cos : \mathbb{R} \rightarrow \mathbb{R}$ and $\sin : \mathbb{R} \rightarrow \mathbb{R}$ are elements of U . Suppose $f_1, f_2 \in U$ and $s_1, s_2 \in \mathbb{R}$. Then, we have for all $x \in \mathbb{R}$

$$\begin{aligned}(s_1 f_1 + s_2 f_2)(x + 2\pi) &= s_1 f_1(x + 2\pi) + s_2 f_2(x + 2\pi) = s_1 f_1(x) + s_2 f_2(x) \\ &= (s_1 f_1 + s_2 f_2)(x)\end{aligned}$$

showing that $s_1 f_1 + s_2 f_2$ is periodic with period 2π . By [Definition 3.21](#), it follows that U is a vector subspace of $F(\mathbb{R}, \mathbb{R})$.

Recall, if $\mathcal{X}, \mathcal{W}$ are sets, $\mathcal{Y} \subset \mathcal{X}, \mathcal{Z} \subset \mathcal{W}$ subsets and $f : \mathcal{X} \rightarrow \mathcal{W}$ a mapping, then the *image* of $\mathcal{Y}$ under f is the set

$$f(\mathcal{Y}) = \{w \in \mathcal{W} \mid \text{there exists an element } y \in \mathcal{Y} \text{ with } f(y) = w\}$$

consisting of all the elements in $\mathcal{W}$ which are hit by an element of $\mathcal{Y}$ under the mapping f . In the special case where $\mathcal{Y}$ is all of $\mathcal{X}$, that is, $\mathcal{Y} = \mathcal{X}$, it is also customary to write $\text{Im}(f)$ instead of $f(\mathcal{X})$ and simply speak of the image of f . Similarly, the *preimage* of $\mathcal{Z}$ under f is the set

$$f^{-1}(\mathcal{Z}) = \{x \in \mathcal{X} \mid f(x) \in \mathcal{Z}\}$$

consisting of all the elements in $\mathcal{X}$ which are mapped onto elements of $\mathcal{Z}$ under f . Notice that f is not assumed to be bijective, hence the inverse mapping $f^{-1} : \mathcal{W} \rightarrow \mathcal{X}$ does not need to exist (and in fact the definition of the preimage does not involve the inverse mapping). Nonetheless the notation $f^{-1}(\mathcal{Z})$ is customary.

It is natural to ask how the image and preimage of subspaces look like under a linear map:

Proposition 3.26 *Let V, W be $\mathbb{K}$ -vector spaces, $U \subset V$ and $Z \subset W$ be vector subspaces and $f : V \rightarrow W$ a linear map. Then the image $f(U)$ is a vector subspace of W and the preimage $f^{-1}(Z)$ is a vector subspace of V .*

Proof Since U is a vector subspace, we have $0_V \in U$. By [Lemma 3.15](#), $f(0_V) = 0_W$, hence $0_W \in f(U)$. For all $w_1, w_2 \in f(U)$ there exist $u_1, u_2 \in U$ with $f(u_1) = w_1$ and $f(u_2) = w_2$. Hence for all $s_1, s_2 \in \mathbb{K}$ we obtain

$$s_1 w_1 + s_2 w_2 = s_1 f(u_1) + s_2 f(u_2) = f(s_1 u_1 + s_2 u_2),$$

where we use the linearity of f . Since U is a subspace, $s_1 u_1 + s_2 u_2$ is an element of U as well. It follows that $s_1 w_1 + s_2 w_2 \in f(U)$ and hence applying [Definition 3.21](#) again, we conclude that $f(U)$ is a subspace of W . The second claim is left to the reader as an exercise. $\square$

Vector subspaces are stable under intersection in the following sense:

Proposition 3.27 *Let V be a $\mathbb{K}$ -vector space, $n \geq 2$ a natural number and $U_1, \dots, U_n$ vector subspaces of V . Then the intersection*

$$U' = \bigcap_{j=1}^n U_j = \{v \in V \mid v \in U_j \text{ for all } j = 1, \dots, n\}$$

is a vector subspace of V as well.

Proof Since U_j is a vector subspace, $0_V \in U_j$ for all $j = 1, \dots, n$. Therefore, $0_V \in U'$, hence U' is not empty. Let $u_1, u_2 \in U'$ and $s_1, s_2 \in \mathbb{K}$. By assumption, $u_1, u_2 \in U_j$ for all $j = 1, \dots, n$. Since U_j is a vector subspace for all $j = 1, \dots, n$ it follows that $s_1 u_1 + s_2 u_2 \in U_j$ for all $j = 1, \dots, n$ and hence $s_1 u_1 + s_2 u_2 \in U'$. By [Definition 3.21](#), it follows that U' is a vector subspace of V . $\square$

Remark 3.28 Notice that the union of subspaces need not be a subspace. Let $V = \mathbb{R}^2$, $\{\vec{e}_1, \vec{e}_2\}$ its standard basis and

$$U_1 = \{s\vec{e}_1 \mid s \in \mathbb{R}\} \quad \text{and} \quad U_2 = \{s\vec{e}_2 \mid s \in \mathbb{R}\}.$$

Then $\vec{e}_1 \in U_1 \cup U_2$ and $\vec{e}_2 \in U_1 \cup U_2$, but $\vec{e}_1 + \vec{e}_2 \notin U_1 \cup U_2$.

The kernel of a linear map $f : V \rightarrow W$ consists of those vectors in V that are mapped onto the zero vector of W :

Definition 3.29 (Kernel) The *kernel* of a linear map $f : V \rightarrow W$ is the preimage of $\{0_W\}$ under f , that is,

$$\text{Ker}(f) = \{v \in V \mid f(v) = 0_W\} = f^{-1}(\{0_W\}).$$

Example 3.30 The kernel of the linear map $\frac{d}{dx} : P_n(\mathbb{R}) \rightarrow P_{n-1}(\mathbb{R})$ consists of the constant polynomials satisfying $f(x) = c$ for all $x \in \mathbb{R}$ and where $c \in \mathbb{R}$ is some constant.

We can characterise the injectivity of a linear map $f : V \rightarrow W$ in terms of its kernel:

Lemma 3.31 A linear map $f : V \rightarrow W$ is injective if and only if $\text{Ker}(f) = \{0_V\}$.

Proof Let $f : V \rightarrow W$ be injective. Suppose $f(v) = 0_W$. Since $f(0_V) = 0_W$ by [Lemma 3.15](#), we have $f(v) = f(0_V)$, hence $v = 0_V$ by the injectivity assumption. It follows that $\text{Ker}(f) = \{0_V\}$. Conversely, suppose $\text{Ker}(f) = \{0_V\}$ and let $v_1, v_2 \in V$ be such that $f(v_1) = f(v_2)$. Then by the linearity we have $f(v_1) - f(v_2) = 0_W = f(v_1 - v_2)$. Hence $v_1 - v_2$ is in the kernel of f so that $v_1 - v_2 = 0_V$ or $v_1 = v_2$. $\square$

An immediate consequence of [Proposition 3.26](#) is:

Corollary 3.32 Let $f : V \rightarrow W$ be a linear map, then its image $\text{Im}(f)$ is a vector subspace of W and its kernel $\text{Ker}(f)$ is a vector subspace of V .

3.3.2 Isomorphisms

Definition 3.33 (Vector space isomorphism) A bijective linear map $f : V \rightarrow W$ is called a (vector space) *isomorphism*. If an isomorphism $f : V \rightarrow W$ exists, then the $\mathbb{K}$ -vector spaces V and W are called *isomorphic*.

By the definition of surjectivity, a map $f : V \rightarrow W$ is surjective if and only if $\text{Im}(f) = W$. Combining this with [Lemma 3.31](#) gives:

Proposition 3.34 *A linear map $f : V \rightarrow W$ is an isomorphism if and only if $\text{Ker}(f) = \{0_V\}$ and $\text{Im}(f) = W$.*

3.4 Generating sets

Definition 3.35 (Linear combination) Let V be a $\mathbb{K}$ -vector space, $k \in \mathbb{N}$ and $\{v_1, \dots, v_k\}$ a set of vectors from V . A *linear combination* of the vectors $\{v_1, \dots, v_k\}$ is a vector of the form

$$w = s_1 v_1 + \dots + s_k v_k = \sum_{i=1}^k s_i v_i$$

for some $s_1, \dots, s_k \in \mathbb{K}$.

Example 3.36 For $n \in \mathbb{N}$ with $n \geq 2$ consider $V = P_n(\mathbb{R})$ and the polynomials $p_1, p_2, p_3 \in P_n(\mathbb{R})$ defined by the rules $p_1(x) = 1$, $p_2(x) = x$, $p_3(x) = x^2$ for all $x \in \mathbb{R}$. A linear combination of $\{p_1, p_2, p_3\}$ is a polynomial of the form $p(x) = ax^2 + bx + c$ where $a, b, c \in \mathbb{R}$.

Definition 3.37 (Subspace generated by a set) Let V be a $\mathbb{K}$ -vector space and $S \subset V$ be a non-empty subset. The *subspace generated by S* is the set $\text{span}(S)$ whose elements are linear combinations of finitely many vectors in S . The set $\text{span}(S)$ is called the *span* of S . Formally, we have

$$\text{span}(S) = \left\{ v \in V \mid v = \sum_{i=1}^k s_i v_i, k \in \mathbb{N}, s_1, \dots, s_k \in \mathbb{K}, v_1, \dots, v_k \in S \right\}.$$

Remark 3.38 The notation $\langle S \rangle$ for the span of S is also in use.

Proposition 3.39 *Let V be a $\mathbb{K}$ -vector space and $S \subset V$ be a non-empty subset. Then $\text{span}(S)$ is a vector subspace of V .*

Proof Since S is non-empty it contains some element, say u . Since u itself is a linear combination of $\{u\}$, it follows that $\text{span}(S)$ is non-empty. Let $k \in \mathbb{N}$ and $v_1 = t_1 w_1 + \dots + t_k w_k$ for $t_1, \dots, t_k \in \mathbb{K}$ and $w_1, \dots, w_k \in S$ be a linear combination of vectors in S . Furthermore, let $j \in \mathbb{N}$ and $v_2 = \hat{t}_1 \hat{w}_1 + \dots + \hat{t}_j \hat{w}_j$ for $\hat{t}_1, \dots, \hat{t}_j$ and $\hat{w}_1, \dots, \hat{w}_j \in S$ be another linear combination of vectors in S . By [Definition 3.21](#), it suffices to show that for all $s_1, s_2 \in \mathbb{K}$ the vector $s_1 v_1 + s_2 v_2$ is a linear combination of vectors in S . Since

$$\begin{aligned} s_1 v_1 + s_2 v_2 &= s_1(t_1 w_1 + \dots + t_k w_k) + s_2(\hat{t}_1 \hat{w}_1 + \dots + \hat{t}_j \hat{w}_j) \\ &= s_1 t_1 w_1 + \dots + s_1 t_k w_k + s_2 \hat{t}_1 \hat{w}_1 + \dots + s_2 \hat{t}_j \hat{w}_j \end{aligned}$$

is a linear combination of the vectors $\{w_1, \dots, w_k, \hat{w}_1, \dots, \hat{w}_j\}$ in S , the claim follows. $\square$

Remark 3.40 For a subset $S \subset V$, we may alternatively define $\text{span}(S)$ to be the smallest vector subspace of V that contains S . This has the advantage of S being allowed to be empty, in which case $\text{span}(\emptyset) = \{0_V\}$, that is, the empty set is a generating set for the zero vector space.

Definition 3.41 Let V be a $\mathbb{K}$ -vector space. A subset $S \subset V$ is called a *generating set* if $\text{span}(S) = V$. The vector space V is called *finite dimensional* if V admits a generating set with finitely many elements (also called a finite set). A vector space that is not finite dimensional will be called *infinite dimensional*.

Example 3.42 Thinking of a field $\mathbb{K}$ as a $\mathbb{K}$ -vector space, the set $S = \{1_{\mathbb{K}}\}$ consisting of the identity element of multiplication is a generating set for $V = \mathbb{K}$. Indeed, for every $x \in \mathbb{K}$ we have $x = x \cdot_V 1_{\mathbb{K}}$.

Example 3.43 The standard basis $S = \{\vec{e}_1, \dots, \vec{e}_n\}$ is a generating set for $\mathbb{K}^n$, since for all $\vec{x} = (x_i)_{1 \leq i \leq n} \in \mathbb{K}^n$, we can write $\vec{x} = x_1 \vec{e}_1 + \dots + x_n \vec{e}_n$ so that $\vec{x}$ is a linear combination of elements of S .

Example 3.44 Let $\mathbf{E}_{k,l} \in M_{m,n}(\mathbb{K})$ for $1 \leq k \leq m$ and $1 \leq l \leq n$ denote the m -by- n matrix satisfying $\mathbf{E}_{k,l} = (\delta_{ki} \delta_{lj})_{1 \leq i \leq m, 1 \leq j \leq n}$. For example, for $m = 2$ and $n = 3$ we have

$$\mathbf{E}_{1,1} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad \mathbf{E}_{1,2} = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad \mathbf{E}_{1,3} = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix}$$

and

$$\mathbf{E}_{2,1} = \begin{pmatrix} 0 & 0 & 0 \\ 1 & 0 & 0 \end{pmatrix}, \quad \mathbf{E}_{2,2} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}, \quad \mathbf{E}_{2,3} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Then $S = \{\mathbf{E}_{k,l}\}_{1 \leq k \leq m, 1 \leq l \leq n}$ is a generating set for $M_{m,n}(\mathbb{K})$, since a matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ can be written as

$$\mathbf{A} = \sum_{k=1}^m \sum_{l=1}^n A_{kl} \mathbf{E}_{k,l}$$

so that $\mathbf{A}$ is a linear combination of the elements of S .

Example 3.45 The vector space $P(\mathbb{R})$ of polynomials is infinite dimensional. In order to see this, consider a finite set of polynomials $\{p_1, \dots, p_n\}$, $n \in \mathbb{N}$ and let d_i denote the degree of the polynomial p_i for $i = 1, \dots, n$. We set $D = \max\{d_1, \dots, d_n\}$. Since a linear combination of the polynomials $\{p_1, \dots, p_n\}$ has degree at most D , any polynomial q whose degree is strictly larger than D will satisfy $q \notin \text{span}\{p_1, \dots, p_n\}$. It follows that $P(\mathbb{R})$ cannot be generated by a finite set of polynomials.

Lemma 3.46 Let $f : V \rightarrow W$ be linear and $S \subset V$ a generating set. If f is surjective, then $f(S)$ is a generating set for W . Furthermore, if f is bijective, then V is finite dimensional if and only if W is finite dimensional.

Proof Let $w \in W$. Since f is surjective there exists $v \in V$ such that $f(v) = w$. Since $\text{span}(S) = V$, there exists $k \in \mathbb{N}$, as well as elements $v_1, \dots, v_k \in S$ and scalars $s_1, \dots, s_k$ such that $v = \sum_{i=1}^k s_i v_i$ and hence $w = \sum_{i=1}^k s_i f(v_i)$, where we use the linearity of f . We conclude that $w \in \text{span}(f(S))$ and since w is arbitrary, it follows that $W = \text{span}(f(S))$.

For the second claim suppose V is finite dimensional, hence we have a finite set S with $\text{span}(S) = V$. The set $f(S)$ is finite as well and satisfies $\text{span}(f(S)) = W$ by the previous argument, hence W is finite dimensional as well. Conversely suppose W is finite dimensional with generating set $T \subset W$. Since f is bijective there exists an inverse mapping $f^{-1} : W \rightarrow V$ which is surjective, hence $V = \text{span}(f^{-1}(T))$ so that V is finite dimensional as well. $\square$

3.5 Linear independence and bases

A set of vectors where no vector can be expressed as a linear combination of the other vectors is called linearly independent. More precisely:

Definition 3.47 (Linear independence) Let $S \subset V$ be a non-empty finite subset so that $S = \{v_1, \dots, v_k\}$ for distinct vectors $v_i \in V$, $i = 1, \dots, k$. We say S is *linearly independent* if

$$s_1 v_1 + \dots + s_k v_k = 0_V \iff s_1 = \dots = s_k = 0,$$

where $s_1, \dots, s_k \in \mathbb{K}$. If S is not linearly independent, then S is called *linearly dependent*. Furthermore, we call a subset $S \subset V$ linearly independent if every finite subset of S is linearly independent. We will call distinct vectors $v_1, \dots, v_k$ linearly independent/dependent if the set $\{v_1, \dots, v_k\}$ is linearly independent/dependent.

Remark 3.48 Instead of *distinct*, many authors write *pairwise distinct*, which means that all pairs of vectors v_i, v_j with $i \neq j$ satisfy $v_i \neq v_j$. Of course, this simply means that the list $v_1, \dots, v_k$ of vectors is not allowed to contain a vector more than once.

Notice that if the vectors $v_1, \dots, v_k \in V$ are linearly dependent, then there exist scalars $s_1, \dots, s_k$, not all zero, so that $\sum_{i=1}^k s_i v_i = 0_V$. After possibly changing the numbering of the vectors and scalars, we can assume that $s_1 \neq 0$. Therefore, we can write

$$v_1 = - \sum_{i=2}^k \left(\frac{s_i}{s_1} \right) v_i,$$

so that v_1 is a linear combination of the vectors $v_2, \dots, v_k$.

Also, observe that a subset T of a linearly independent set S is itself linearly independent. (Why?)

Example 3.49 We consider the polynomials $p_1, p_2, p_3 \in P(\mathbb{R})$ defined by the rules $p_1(x) = 1, p_2(x) = x, p_3(x) = x^2$ for all $x \in \mathbb{R}$. Then $\{p_1, p_2, p_3\}$ is linearly independent. In order to see this, consider the condition

$$(3.9) \quad s_1 p_1 + s_2 p_2 + s_3 p_3 = 0_{P(\mathbb{R})} = o$$

where $o : \mathbb{R} \rightarrow \mathbb{R}$ denotes the zero polynomial. Since (3.9) means that

$$s_1 p_1(x) + s_2 p_2(x) + s_3 p_3(x) = o(x),$$

for all $x \in \mathbb{R}$, we can evaluate this condition for any choice of real number x . Taking $x = 0$ gives

$$s_1 p_1(0) + s_2 p_2(0) + s_3 p_3(0) = o(0) = 0 = s_1.$$

Taking $x = 1$ and $x = -1$ gives

$$0 = s_2 p_2(1) + s_3 p_3(1) = s_2 + s_3,$$

$$0 = s_2 p_2(-1) + s_3 p_3(-1) = -s_2 + s_3,$$

so that $s_2 = s_3 = 0$ as well. It follows that $\{p_1, p_2, p_3\}$ is linearly independent.

Remark 3.50 By convention, the empty set is linearly independent.

Definition 3.51 (Basis) A subset $\mathcal{S} \subset V$ which is a generating set of V and also linearly independent is called a *basis* of V .

Video Basis

Example 3.52 Thinking of a field $\mathbb{K}$ as a $\mathbb{K}$ -vector space, the set $\{1_{\mathbb{K}}\}$ is linearly independent, since $1_{\mathbb{K}} \neq 0_{\mathbb{K}}$. Example 3.42 implies that $\{1_{\mathbb{K}}\}$ is a basis of $\mathbb{K}$.

Example 3.53 Clearly, the standard basis $\{\vec{e}_1, \dots, \vec{e}_n\}$ of $\mathbb{K}^n$ is linearly independent since

$$s_1 \vec{e}_1 + \dots + s_n \vec{e}_n = \begin{pmatrix} s_1 \\ \vdots \\ s_n \end{pmatrix} = 0_{\mathbb{K}^n} = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix} \iff s_1 = \dots = s_n = 0.$$

It follows together with Example 3.43 that the standard basis of $\mathbb{K}^n$ is indeed a basis in the sense of Definition 3.51.

Example 3.54 The matrices $\mathbf{E}_{k,l} \in M_{m,n}(\mathbb{K})$ for $1 \leq k \leq m$ and $1 \leq l \leq n$ are linearly independent. Suppose we have scalars $s_{kl} \in \mathbb{K}$ such that

$$\sum_{k=1}^m \sum_{l=1}^n s_{kl} \mathbf{E}_{k,l} = \mathbf{0}_{m,n} = \begin{pmatrix} s_{11} & \dots & s_{1n} \\ \vdots & \ddots & \vdots \\ s_{m1} & \dots & s_{mn} \end{pmatrix} = \begin{pmatrix} 0 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 0 \end{pmatrix}$$

so $s_{kl} = 0$ for all $1 \leq k \leq m$ and all $1 \leq l \leq n$. It follows together with [Example 3.44](#) that $\{\mathbf{E}_{k,l}\}_{1 \leq k \leq m, 1 \leq l \leq n}$ is a basis of $M_{m,n}(\mathbb{K})$. We refer to $\{\mathbf{E}_{k,l}\}_{1 \leq k \leq m, 1 \leq l \leq n}$ as the *standard basis* of $M_{m,n}(\mathbb{K})$.

Example 3.55 Combining [Remark 3.40](#) and [Remark 3.50](#) we conclude that the empty set is a basis for the zero vector space $\{0\}$.

Lemma 3.56 Let $f : V \rightarrow W$ be an injective linear map. Suppose $S \subset V$ is linearly independent, then $f(S) \subset W$ is also linearly independent.

Proof Let $\{w_1, \dots, w_k\} \subset f(S)$ be a finite subset for some $k \in \mathbb{N}$ some and distinct vectors $w_i \in W$, where $1 \leq i \leq k$. Then there exist vectors $v_1, \dots, v_k$ with $f(v_i) = w_i$ for $1 \leq i \leq k$. Suppose there exist scalars $s_1, \dots, s_k$ such that $s_1 w_1 + \dots + s_k w_k = 0_W$. Using the linearity of f , this implies

$$0_W = s_1 w_1 + \dots + s_k w_k = s_1 f(v_1) + \dots + s_k f(v_k) = f(s_1 v_1 + \dots + s_k v_k).$$

Since f is injective we have $\text{Ker}(f) = \{0_V\}$ by [Lemma 3.31](#). Since $s_1 v_1 + \dots + s_k v_k \in \text{Ker } f$ it follows that $s_1 v_1 + \dots + s_k v_k = 0_V$, hence $s_1 = \dots = s_k = 0$ by the linear independence of S . It follows that $f(S)$ is linearly independent as well. $\square$

Exercises

Exercise 3.57 Let $U \subset V$ be a vector subspace and $k \in \mathbb{N}$ with $k \geq 2$. Show that for $u_1, \dots, u_k \in U$ and $s_1, \dots, s_k \in \mathbb{K}$, we have $s_1 u_1 + \dots + s_k u_k \in U$.

Exercise 3.58 (Planes through the origin) Let $\vec{w}_1, \vec{w}_2 \neq 0_{\mathbb{R}^3}$ and $\vec{w}_1 \neq s\vec{w}_2$ for all $s \in \mathbb{R}$. Show that the plane

$$U = \{s_1 \vec{w}_1 + s_2 \vec{w}_2 \mid s_1, s_2 \in \mathbb{R}\}$$

is a vector subspace of $\mathbb{R}^3$.

Exercise 3.59 (Polynomials) Let $n \in \mathbb{N} \cup \{0\}$ and $P_n(\mathbb{R})$ denote the subset of $P(\mathbb{R})$ consisting of polynomials of degree *at most* n . Show that $P_n(\mathbb{R})$ is a subspace of $P(\mathbb{R})$ for all $n \in \mathbb{N} \cup \{0\}$.

Exercise 3.60 Show that the $\mathbb{K}$ -vector space $\mathbb{K}^n$ of column vectors with n entries is isomorphic to the $\mathbb{K}$ -vector space $\mathbb{K}_n$ of row vectors with n entries.

Exercise 3.61 Show that the $\mathbb{R}$ -vector spaces $P_n(\mathbb{R})$ and $\mathbb{R}^{n+1}$ are isomorphic for all $n \in \mathbb{N} \cup \{0\}$.

Exercise 3.62 Show that for a non-empty subset S of a $\mathbb{K}$ -vector space V , the set $\text{span}(S)$ as defined in [Definition 3.37](#) is the same as the set $\text{span}(S)$ as defined in [Remark 3.40](#). In particular, [Proposition 3.39](#) remains true when removing the assumption that S is non-empty.

Exercise 3.63 Show that a subset $\{v\}$ consisting of a single vector $v \in V$ is linearly independent if and only if $v \neq 0_V$.

3.6 The dimension

3.6.1 Defining the dimension

Intuitively, we might define the dimension of a finite dimensional vector space V to be the number of elements of any basis of V , so that a line is 1-dimensional, a plane is 2-dimensional and so on. Of course, this definition only makes sense if we know that there always exists a basis of V and that the number of elements in the basis is independent of the chosen basis. Perhaps surprisingly, these facts take quite a bit of work to prove.

Theorem 3.64 *Let V be a $\mathbb{K}$ -vector space.*

- (i) *Any subset $S \subset V$ generating V admits a subset $T \subset S$ that is a basis of V .*
- (ii) *Any subset $S \subset V$ that is linearly independent in V is contained in a subset $T \subset V$ that is a basis of V .*
- (iii) *If S_1, S_2 are bases of V , then there exists a bijective map $f : S_1 \rightarrow S_2$.*
- (iv) *If V is finite dimensional, then any basis of V is a finite set and the number of elements in the basis is independent of the choice of the basis.*

Corollary 3.65 *Every $\mathbb{K}$ -vector space V admits at least one basis.*

Proof Since V is a generating set for V , we can apply (i) from [Theorem 3.64](#) to $S = V$ to obtain a basis of V . □

Remark 3.66 Let $\mathcal{X}$ be a set with finitely many elements. We write $\text{Card}(\mathcal{X})$ – for *cardinality* – for the number of elements of $\mathcal{X}$.

Definition 3.67 The dimension of a finite dimensional $\mathbb{K}$ -vector space V , denoted by $\dim(V)$ or $\dim_{\mathbb{K}}(V)$, is the number of elements of any basis of V .

Example 3.68

- (i) The zero vector space $\{0\}$ has the empty set as a basis and hence is 0-dimensional.
- (ii) A field $\mathbb{K}$ – thought of as a $\mathbb{K}$ -vector space – has $\{1_{\mathbb{K}}\}$ as a basis and hence is 1-dimensional.
- (iii) The vector space $\mathbb{K}^n$ has $\{\vec{e}_1, \dots, \vec{e}_n\}$ as a basis and hence is n -dimensional.
- (iv) The vector space $M_{m,n}(\mathbb{K})$ has $\mathbf{E}_{k,l}$ for $1 \leq k \leq m$ and $1 \leq l \leq n$ as a basis, hence it is mn -dimensional.

We will only prove [Theorem 3.64](#) for finite dimensional vector spaces. This will be done with the help of three lemmas.

Lemma 3.69 *Let V be a $\mathbb{K}$ -vector space, $S \subset V$ linearly independent and $v_0 \in V$. Suppose $v_0 \notin \text{span}(S)$, then $S \cup \{v_0\}$ is linearly independent.*

Proof Let $\mathcal{T}$ be a finite subset of $\mathcal{S} \cup \{v_0\}$. If $v_0 \notin \mathcal{T}$, then $\mathcal{T}$ is linearly independent, as $\mathcal{S}$ is linearly independent. So suppose $v_0 \in \mathcal{T}$. There exist distinct elements $v_1, \dots, v_n$ of $\mathcal{S}$ so that $\mathcal{T} = \{v_0, v_1, \dots, v_n\}$. Suppose $s_0 v_0 + s_1 v_1 + \dots + s_n v_n = 0_V$ for some scalars $s_0, s_1, \dots, s_n \in \mathbb{K}$. If $s_0 \neq 0$, then we can write

$$v_0 = - \sum_{i=1}^n \frac{s_i}{s_0} v_i,$$

contradicting the assumption that $v_0 \notin \text{span}(\mathcal{S})$. Hence we must have $s_0 = 0$. Since $s_0 = 0$ it follows that $s_1 v_1 + \dots + s_n v_n = 0_V$ so that $s_1 = \dots = s_n = 0$ by the linear independence of $\mathcal{S}$. We conclude that $\mathcal{S} \cup \{v_0\}$ is linearly independent. $\square$

Lemma 3.70 *Let V be a $\mathbb{K}$ -vector space and $\mathcal{S} \subset V$ a generating set. If $v_0 \in \text{span}(\mathcal{S} \setminus \{v_0\})$, then $\mathcal{S} \setminus \{v_0\}$ is a generating set.*

Proof Since $v_0 \in \text{span}(\mathcal{S} \setminus \{v_0\})$, there exist vectors $v_1, \dots, v_n \in \mathcal{S}$ with $v_i \neq v_0$ and scalars $s_1, \dots, s_n$ so that $v_0 = s_1 v_1 + \dots + s_n v_n$. Suppose $v \in V$. Since $\mathcal{S}$ is a generating set, there exist vectors $w_1, \dots, w_k \in \mathcal{S}$ and scalars $t_1, \dots, t_k$ so that $v = t_1 w_1 + \dots + t_k w_k$. If $\{w_1, \dots, w_k\}$ does not contain v_0 , then $v \in \text{span}(\mathcal{S} \setminus \{v_0\})$, so assume that $v_0 \in \{w_1, \dots, w_k\}$. After possibly relabelling the elements of $\{w_1, \dots, w_k\}$ we can assume that $v_0 = w_1$. Hence we have

$$v = t_1 (s_1 v_1 + \dots + s_n v_n) + t_2 w_2 + \dots + t_k w_k$$

with $v_0 \neq v_i$ for $1 \leq i \leq n$ and $v_0 \neq w_j$ for $2 \leq j \leq k$. It follows that $v \in \text{span}(\mathcal{S} \setminus \{v_0\})$, as claimed. $\square$

Lemma 3.71 *Let V be a finite dimensional $\mathbb{K}$ -vector space and $\mathcal{S} \subset V$ a finite set with n elements which generates V . If $\mathcal{T} \subset V$ has more than n elements, then $\mathcal{T}$ is linearly dependent.*

Proof We show that if $\mathcal{T}$ has exactly $n + 1$ elements, then it is linearly dependent. In the other cases, $\mathcal{T}$ contains a subset with exactly $n + 1$ elements and if this subset is linearly dependent, then so is $\mathcal{T}$.

We prove the claim by induction on $n \geq 0$. Let $\mathcal{A}(n)$ be the following statement: “For any $\mathbb{K}$ -vector space V , if there exists a generating subset $\mathcal{S} \subset V$ with n elements, then all subsets of V with $n + 1$ elements are linearly dependent.”

We first show that $\mathcal{A}(0)$ is true. A subset with zero elements is the empty set $\emptyset$. Hence $V = \text{span}(\emptyset) = \{0_V\}$ is the zero vector space. The only subset of $\{V\}$ with 1 element is $\{0_V\}$. Since $s 0_V = 0_V$ for all $s \in \mathbb{K}$, the set $\{0_V\}$ is linearly dependent, thus showing that $\mathcal{A}(0)$ is correct.

Suppose $n \geq 1$ and that $\mathcal{A}(n - 1)$ is true. We want to argue that $\mathcal{A}(n)$ is true as well. Suppose V is generated by the set $\mathcal{S} = \{v_1, \dots, v_n\}$ with n elements. Let $\mathcal{T} = \{w_1, \dots, w_{n+1}\}$ be a subset with $n + 1$ elements. We need to show that $\mathcal{T}$ is linearly dependent. Since $\mathcal{S}$ is generating, we have scalars $s_{ij} \in \mathbb{K}$ with $1 \leq i \leq n + 1$ and $1 \leq j \leq n$ so that

$$(3.10) \quad w_i = \sum_{j=1}^n s_{ij} v_j$$

for all $1 \leq i \leq n + 1$. We now consider two cases:

Case 1. If $s_{11} = \cdots = s_{n+1,1} = 0$, then (3.10) gives for all $1 \leq i \leq n+1$

$$w_i = \sum_{j=2}^n s_{ij} v_j.$$

Notice that the summation now starts at $j = 2$. This implies that $\mathcal{T} \subset W$, where $W = \text{span}\{v_2, \dots, v_n\}$. We can now apply $\mathcal{A}(n-1)$ to the vector space W , the generating set $\mathcal{S}_1 = \{v_2, \dots, v_n\}$ and the subset with n elements being $\mathcal{T}_1 = \{w_1, \dots, w_n\}$. It follows that $\mathcal{T}_1$ is linearly dependent and hence so is $\mathcal{T}$, as it contains $\mathcal{T}_1$.

Case 2. Suppose there exists i so that $s_{i1} \neq 0$. Then, after possibly relabelling the vectors, we can assume that $s_{11} \neq 0$. For $2 \leq i \leq n+1$ we thus obtain from (3.10)

$$\begin{aligned} w_i - \frac{s_{i1}}{s_{11}} w_1 &= w_i - \frac{s_{i1}}{s_{11}} \left(\sum_{j=1}^n s_{1j} v_j \right) = \sum_{j=1}^n s_{ij} v_j - \frac{s_{i1}}{s_{11}} \left(\sum_{j=1}^n s_{1j} v_j \right) \\ &= \sum_{j=1}^n \left(s_{ij} - \frac{s_{i1}}{s_{11}} s_{1j} \right) v_j \\ &= \underbrace{\left(s_{i1} - \frac{s_{i1}}{s_{11}} s_{11} \right)}_{=0} v_1 + \sum_{j=2}^n \left(s_{ij} - \frac{s_{i1}}{s_{11}} s_{1j} \right) v_j \\ &= \sum_{j=2}^n \left(s_{ij} - \frac{s_{i1}}{s_{11}} s_{1j} \right) v_j. \end{aligned}$$

Hence, setting

$$(3.11) \quad \hat{w}_i = w_i - \frac{s_{i1}}{s_{11}} w_1$$

for $2 \leq i \leq n+1$ and $\hat{s}_{ij} = s_{ij} - \frac{s_{i1}}{s_{11}} s_{1j}$ for $2 \leq i \leq n+1$ and $2 \leq j \leq n$, we obtain the relations

$$\hat{w}_i = \sum_{j=2}^n \hat{s}_{ij} v_j$$

for all $2 \leq i \leq n+1$. Therefore, the set $\hat{\mathcal{T}} = \{\hat{w}_2, \dots, \hat{w}_{n+1}\}$ with n elements is contained in W which is generated by $n-1$ elements. Applying $\mathcal{A}(n-1)$, we conclude that $\hat{\mathcal{T}}$ is linearly dependent. It follows that we have scalars $t_2, \dots, t_{n+1}$ not all zero so that

$$t_2 \hat{w}_2 + \cdots + t_{n+1} \hat{w}_{n+1} = 0_V.$$

Using (3.11), we get

$$\sum_{i=2}^{n+1} t_i \left(w_i - \frac{s_{i1}}{s_{11}} w_1 \right) = - \left(\sum_{i=2}^{n+1} t_i \frac{s_{i1}}{s_{11}} \right) w_1 + t_2 w_2 + \cdots + t_{n+1} w_{n+1} = 0_V.$$

Since not all scalars $t_2, \dots, t_{n+1}$ are zero, it follows that $w_1, \dots, w_{n+1}$ are linearly dependent and hence so is $\mathcal{T}$. $\square$

Proof of Theorem 3.64 We restrict to the case where V is finite dimensional. Hence there exists an integer $n \geq 0$ so that V has a generating set $\mathcal{S}_0$ with n elements.

(i) Let $\mathcal{S} \subset V$ be a subset generating V . We consider the set $\mathcal{X}$ consisting of those integers $d \geq 0$ for which there exists a linearly independent subset $\mathcal{T} \subset \mathcal{S}$ with d elements. Since $\emptyset \subset \mathcal{S}$, we have $0 \in \mathcal{X}$, so $\mathcal{X}$ is non-empty. Furthermore, $\mathcal{X}$ is a finite set, as it cannot contain any integer greater than n by Lemma 3.71. Let $m \in \mathcal{X}$ be the largest integer and $\mathcal{T} \subset \mathcal{S}$ a set with m elements. We want to argue that $\mathcal{T}$ is a basis of V . Suppose $\mathcal{T}$ is not a basis of V . Then there exists an element $v_0 \in \mathcal{S}$ so that $v_0 \notin \text{span}(\mathcal{T})$, since if no such element exists, we have $\mathcal{S} \subset \text{span}(\mathcal{T})$ and hence $V = \text{span}(\mathcal{S}) \subset \text{span}(\mathcal{T})$ contradicting the assumption that $\mathcal{T}$ is not a basis of V . Applying Lemma 3.69, we conclude that

$\hat{\mathcal{T}} = \{v_0\} \cup \mathcal{T} \subset \mathcal{S}$ is linearly independent. Since $\hat{\mathcal{T}}$ has $m + 1$ elements, we have $m + 1 \in \mathcal{X}$, contradicting the fact that m is the largest integer in $\mathcal{X}$. It follows that $\mathcal{T}$ must be a basis of V .

(ii) Let $\mathcal{S} \subset V$ be a subset that is linearly independent in V . Let $\hat{\mathcal{X}}$ denote the set consisting of those integers $d \geq 0$ for which there exists a subset $\mathcal{T} \subset V$ with d elements, which contains $\mathcal{S}$ and which is a generating set of V . Notice that $\mathcal{S} \cup \mathcal{S}_0$ is such a set, hence $\hat{\mathcal{X}}$ is not empty. Let m denote the smallest element of $\hat{\mathcal{X}}$ and $\mathcal{T}$ be a generating subset of V containing $\mathcal{S}$ and with m elements. We want to argue that $\mathcal{T}$ is basis for V . By assumption, $\mathcal{T}$ generates V , hence we need to check that $\mathcal{T}$ is linearly independent in V . Suppose $\mathcal{T}$ is linearly dependent and write $\mathcal{T} = \{v_1, \dots, v_m\}$ for distinct elements of V . Suppose $\mathcal{S} = \{v_1, \dots, v_k\}$ for some $k \leq m$. This holds true since $\mathcal{S} \subset \mathcal{T}$. Since $\mathcal{T}$ is linearly dependent we have scalars $s_1, \dots, s_m$ so that

$$s_1 v_1 + \dots + s_m v_m = 0_V.$$

There must exist a scalar s_i with $i > k$ such that $s_i \neq 0$. Otherwise $\mathcal{S}$ would be linearly dependent. After possibly relabelling the vectors, we can assume that $s_{k+1} \neq 0$ so that

$$(3.12) \quad v_{k+1} = -\frac{1}{s_{k+1}} (s_1 v_1 + \dots + s_k v_k + s_{k+2} v_{k+2} + \dots + s_m v_m).$$

Let $\hat{\mathcal{T}} = \{v_1, \dots, v_k, v_{k+2}, \dots, v_m\}$. Then $\mathcal{S} \subset \hat{\mathcal{T}}$ and (3.12) shows that $v_{k+1} \in \text{span}(\hat{\mathcal{T}})$. Lemma 3.70 shows that $\hat{\mathcal{T}}$ generates V , contains $\mathcal{S}$ and has $m - 1$ elements, contradicting the minimality of m .

(iii) Suppose $\mathcal{S}_1$ is a basis of V with n_1 elements and $\mathcal{S}_2$ is a basis of V with n_2 elements. Since $\mathcal{S}_2$ is linearly independent and $\mathcal{S}_1$ generates V , Lemma 3.71 implies that $n_2 \leq n_1$. Likewise, we conclude that $n_2 \geq n_1$. It follows that $n_1 = n_2$ and hence there exists a bijective mapping from $\mathcal{S}_1$ to $\mathcal{S}_2$ as these are finite sets with the same number of elements.

(iv) is an immediate consequence of (iii). □

3.6.2 Properties of the dimension

Lemma 3.72 *Isomorphic finite dimensional vector spaces have the same dimension.*

Proof Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ an isomorphism. Let $\mathcal{S} \subset V$ be a basis of V , then $f(\mathcal{S}) \subset W$ is a basis of W , by combining Lemma 3.46 and Lemma 3.56. Since $\mathcal{S}$ and $f(\mathcal{S})$ have the same number of elements, we have $\dim(V) = \dim(W)$. □

Lemma 3.73 *A subspace of a finite dimensional $\mathbb{K}$ -vector space is finite dimensional as well.*

Proof Let V be a finite dimensional $\mathbb{K}$ -vector space and $U \subset V$ a subspace. Let $\mathcal{S} = \{v_1, \dots, v_n\}$ be a basis of V . For $1 \leq i \leq n$, we define $U_i = U \cap \text{span}\{v_1, \dots, v_i\}$. By construction, each U_i is a subspace and $U_1 \subset U_2 \subset \dots \subset U_n = U$, since $\mathcal{S}$ is a basis of V .

We will show inductively that all U_i are finite dimensional. Notice that U_1 is a subspace of $\text{span}\{v_1\}$. The only subspaces of $\text{span}\{v_1\}$ are $\{0_V\}$ and $\{tv_1 \mid t \in \mathbb{K}\}$, both are finite dimensional, hence U_1 is finite dimensional.

Assume $i \geq 2$. We will show next that if U_{i-1} is finite dimensional, then so is U_i . Let $\mathcal{T}_{i-1}$ be a basis of U_{i-1} . If $U_i = U_{i-1}$, then U_i is finite dimensional as well, so assume there exists a non-zero vector $w \in U_i \setminus U_{i-1}$. Since $\mathcal{S}$ is a basis of V and since $w \in \text{span}\{v_1, \dots, v_i\}$, there exist scalars $s_1, \dots, s_i$ so that $w = s_1 v_1 + \dots + s_i v_i$. By assumption, $w \notin U_{i-1}$, hence $s_i \neq 0$. Any vector $v \in U_i$ can be written as $v = t_1 v_1 + \dots + t_i v_i$ for scalars $t_1, \dots, t_i$. We now compute

$$\begin{aligned} v - \frac{t_i}{s_i} w &= \sum_{k=1}^i t_k v_k - \frac{t_i}{s_i} \left(\sum_{k=1}^i s_k v_k \right) = \sum_{k=1}^i \left(t_k - \frac{t_i}{s_i} s_k \right) v_k \\ &= \sum_{k=1}^{i-1} \left(t_k - \frac{t_i}{s_i} s_k \right) v_k \end{aligned}$$

so that $v - (t_i/s_i)w$ can be written as a linear combination of the vectors $v_1, \dots, v_{i-1}$, hence is an element of U_{i-1} . Recall that $\mathcal{T}_{i-1}$ is a basis of U_{i-1} , hence $v - (t_i/s_i)w$ is a linear combination of elements of $\mathcal{T}_{i-1}$. It follows that any vector $v \in U_i$ is a linear combination of elements of $\mathcal{T}_{i-1} \cup \{w\}$, that is, $\mathcal{T}_{i-1} \cup \{w\}$ generates U_i . Since $\mathcal{T}_{i-1} \cup \{w\}$ contains finitely many vectors, it follows that U_i is finite dimensional. $\square$

Proposition 3.74 *Let V be a finite dimensional $\mathbb{K}$ -vector space. Then for any subspace $U \subset V$*

$$0 \leq \dim(U) \leq \dim(V).$$

Furthermore $\dim(U) = 0$ if and only if $U = \{0_V\}$ and $\dim(U) = \dim(V)$ if and only if $V = U$.

Proof By Lemma 3.73, U is finite dimensional and hence by Corollary 3.65 admits a basis $\mathcal{S}$. By Theorem 3.64 (ii), there is a basis $\mathcal{T}$ of V which contains $\mathcal{S}$. Therefore

$$0 \leq \dim(U) = \text{Card}(\mathcal{S}) \leq \text{Card}(\mathcal{T}) = \dim(V).$$

Suppose $\dim(V) = \dim(U)$, then $\text{Card}(\mathcal{S}) = \text{Card}(\mathcal{T})$ and hence $\mathcal{S} = \mathcal{T}$ since every element of $\mathcal{S}$ is an element of $\mathcal{T}$ and $\mathcal{S}$ and $\mathcal{T}$ have the same number of elements. Therefore, we get $U = \text{span}(\mathcal{S}) = \text{span}(\mathcal{T}) = V$. Since $\dim U = 0$ if and only if the empty set is a basis for U we have $\dim U = 0$ if and only if $U = \{0_V\}$. $\square$

Definition 3.75 (Rank of a linear map and matrix) Let V, W be $\mathbb{K}$ -vector spaces with W finite dimensional. The *rank* of a linear map $f : V \rightarrow W$ is defined as

$$\text{rank}(f) = \dim \text{Im}(f).$$

If $\mathbf{A} \in M_{m,n}(\mathbb{K})$ is a matrix, then we define

$$\text{rank}(\mathbf{A}) = \text{rank}(f_{\mathbf{A}}).$$

The *nullity* of a linear map $f : V \rightarrow W$ is the dimension of its kernel, $\text{nullity}(f) = \dim \text{Ker}(f)$. The following *important* theorem establishes a relation between the nullity and the rank of a linear map. It states something that is intuitively not surprising, namely that the dimension of the image of a linear map $f : V \rightarrow W$ is the dimension of the vector space V minus the dimension of the subspace of vectors that we “lose”, that is, those that are mapped onto the zero vector of W . More precisely:

Theorem 3.76 (Rank-nullity theorem) *Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. Then we have*

$$\dim(V) = \dim \text{Ker}(f) + \dim \text{Im}(f) = \text{nullity}(f) + \text{rank}(f).$$

Proof Let $d = \dim \text{Ker}(f)$ and $n = \dim V$, so that $d \leq n$ by [Proposition 3.74](#). Let $\{v_1, \dots, v_d\}$ be a basis of $\mathcal{S} = \text{Ker}(f)$. By [Theorem 3.64](#) (ii) we can find linearly independent vectors $\hat{\mathcal{S}} = \{v_{d+1}, \dots, v_n\}$ so that $\mathcal{T} = \mathcal{S} \cup \hat{\mathcal{S}}$ is a basis of V . Now $U = \text{span}(\hat{\mathcal{S}})$ is a subspace of V of dimension $n - d$. We consider the linear map

$$g : U \rightarrow \text{Im}(f), \quad v \mapsto f(v).$$

We want to show that g is an isomorphism, since then $\dim \text{Im}(f) = \dim(U) = n - d$, so that

$$\dim \text{Im}(f) = n - d = \dim(V) - \dim \text{Ker}(f),$$

as claimed.

We first show that g is injective. Assume $g(v) = 0_W$. Since $v \in U$, we can write $v = s_{d+1}v_{d+1} + \dots + s_nv_n$ for scalars $s_{d+1}, \dots, s_n$. Since $g(v) = 0_W$ we have $v \in \text{Ker}(f)$, hence we can also write $v = s_1v_1 + \dots + s_dv_d$ for scalars $s_1, \dots, s_d$, subtracting the two expressions for v , we get

$$0_V = s_1v_1 + \dots + s_dv_d - s_{d+1}v_{d+1} - \dots - s_nv_n.$$

Since $\{v_1, \dots, v_n\}$ is a basis, it follows that all the coefficients s_i vanish, where $1 \leq i \leq n$. Therefore we have $v = 0_V$ and g is injective.

Second, we show that g is surjective. Suppose $w \in \text{Im}(f)$ so that $w = f(v)$ for some vector $v \in V$. We write $v = \sum_{i=1}^n s_i v_i$ for scalars $s_1, \dots, s_n$. Using the linearity of f , we compute

$$w = f(v) = f\left(\sum_{i=1}^n s_i v_i\right) = f\left(\underbrace{\sum_{i=d+1}^n s_i v_i}_{=\hat{v}}\right) = f(\hat{v})$$

where $\hat{v} \in U$. We thus have an element $\hat{v}$ with $g(\hat{v}) = w$. Since w was arbitrary, we conclude that g is surjective. $\square$

Corollary 3.77 *Let V, W be finite dimensional $\mathbb{K}$ -vector spaces with $\dim(V) = \dim(W)$ and $f : V \rightarrow W$ a linear map. Then the following statements are equivalent:*

- (i) f is injective;
- (ii) f is surjective;
- (iii) f is bijective.

Proof (i) $\Rightarrow$ (ii) By [Lemma 3.31](#), the map f is injective if and only if $\text{Ker}(f) = \{0_V\}$ so that $\dim \text{Ker}(f) = 0$ by [Example 3.68](#) (i). [Theorem 3.76](#) implies that $\dim \text{Im}(f) = \dim(V) = \dim(W)$ and hence [Proposition 3.74](#) implies that $\text{Im}(f) = W$, that is, f is surjective.

(ii) $\Rightarrow$ (iii) Since f is surjective $\text{Im}(f) = W$ and hence $\dim \text{Im}(f) = \dim(W) = \dim(V)$. [Theorem 3.76](#) implies that $\dim \text{Ker}(f) = 0$ so that $\text{Ker}(f) = \{0_V\}$ by [Proposition 3.74](#). Applying [Lemma 3.31](#) again shows that f is injective and hence bijective.

(iii) $\Rightarrow$ (i) Since f is bijective, it is also injective. $\square$

Corollary 3.78 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. Then $\text{rank}(f) \leq \min\{\dim(V), \dim(W)\}$ and

$$\text{rank}(f) = \dim(V) \iff f \text{ is injective,}$$

$$\text{rank}(f) = \dim(W) \iff f \text{ is surjective.}$$

Proof For the first claim it is sufficient to show that $\text{rank}(f) \leq \dim(V)$ and $\text{rank}(f) \leq \dim(W)$. By definition, $\text{rank}(f) = \dim \text{Im}(f)$ and since $\text{Im}(f) \subset W$, we have $\text{rank}(f) = \dim \text{Im}(f) \leq \dim(W)$ with equality if and only if f is surjective, by [Proposition 3.74](#).

[Theorem 3.76](#) implies that $\text{rank}(f) = \dim \text{Im}(f) = \dim(V) - \dim \text{Ker}(f) \leq \dim(V)$ with equality if and only if $\dim \text{Ker}(f) = 0$, that is, when f is injective (as we have just seen in the proof of the previous corollary). $\square$

Corollary 3.79 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. Then we have

(i) If $\dim(V) < \dim(W)$, then f is not surjective;

(ii) If $\dim(V) > \dim(W)$, then f is not injective. In particular, there exist non-zero vectors $v \in V$ with $f(v) = 0_W$.

Proof (i) Suppose $\dim(V) < \dim(W)$, then by [Theorem 3.76](#)

$$\text{rank}(f) = \dim(V) - \dim \text{Ker}(f) \leq \dim(V) < \dim(W)$$

and the claim follows from [Corollary 3.78](#).

(ii) Suppose $\dim(V) > \dim(W)$, then

$$\text{rank}(f) \leq \dim(W) < \dim(V)$$

and the claim follows from [Corollary 3.78](#). $\square$

Proposition 3.80 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces. Then there exists an isomorphism $\Theta : V \rightarrow W$ if and only if $\dim(V) = \dim(W)$.

Proof $\Rightarrow$ This was already proved in [Lemma 3.72](#).

$\Leftarrow$ Let $\dim(V) = \dim(W) = n \in \mathbb{N}$. Choose a basis $\mathcal{T} = \{w_1, \dots, w_n\}$ of W and consider the linear map

$$\Theta : \mathbb{K}^n \rightarrow W, \quad \vec{x} \mapsto x_1 w_1 + \dots + x_n w_n,$$

where $\vec{x} = (x_i)_{1 \leq i \leq n}$. Notice that Θ is injective. Indeed, if $\Theta(\vec{x}) = x_1 w_1 + \dots + x_n w_n = 0_W$, then $x_1 = \dots = x_n = 0$, since $\{w_1, \dots, w_n\}$ are linearly independent. We thus conclude $\text{Ker } \Theta = \{0_V\}$ and hence [Lemma 3.31](#) implies that Θ is injective and therefore bijective by [Corollary 3.77](#). The map Θ is linear and bijective, thus an isomorphism. Likewise, for a choice of basis $\mathcal{S} = \{v_1, \dots, v_n\}$ of V , we obtain an isomorphism $\Phi : \mathbb{K}^n \rightarrow V$. Since the composition of bijective maps is again bijective, the map $\Theta \circ \Phi^{-1} : V \rightarrow W$ is bijective and since by [Proposition 3.16](#) the composition of linear maps is again linear, the map $\Theta \circ \Phi^{-1} : V \rightarrow W$ is an isomorphism. $\square$

Corollary 3.81 Suppose $\mathbf{A} \in M_{m,n}(\mathbb{K})$ is invertible with inverse $\mathbf{A}^{-1} \in M_{n,m}(\mathbb{K})$. Then $n = m$, hence $\mathbf{A}$ is a square matrix.

Proof Consider $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$. By [Proposition 3.17](#), $f_{\mathbf{A}}$ is bijective and hence an isomorphism. [Proposition 3.80](#) implies that $n = m$. □

Exercises

Exercise 3.82 Show that $f : \mathcal{X} \rightarrow \mathcal{Y}$ admits a left inverse if and only if f is injective and that $f : \mathcal{X} \rightarrow \mathcal{Y}$ admits a right inverse if and only if f is surjective.

3.7 Matrix representation of linear maps

Notice that [Proposition 3.80](#) implies that every finite dimensional $\mathbb{K}$ -vector space V is isomorphic to $\mathbb{K}^n$, where $n = \dim(V)$. Choosing an isomorphism from V to $\mathbb{K}^n$ allows to uniquely describe each vector of V in terms of n scalars, its *coordinates*.

Definition 3.83 (Linear coordinate system) Let V be a $\mathbb{K}$ -vector space of dimension $n \in \mathbb{N}$. A *linear coordinate system* is an injective linear map $\varphi : V \rightarrow \mathbb{K}^n$. The entries of the vector $\varphi(v) \in \mathbb{K}^n$ are called the *coordinates* of the vector $v \in V$ with respect to the coordinate system φ .

We only request that φ is injective, but the mapping φ is automatically bijective by [Corollary 3.77](#).

Example 3.84 (Standard coordinates) On the vector space $\mathbb{K}^n$ we have a linear coordinate system defined by the identity mapping, that is, we define $\varphi(\vec{v}) = \vec{v}$ for all $\vec{v} \in \mathbb{K}^n$. We call this coordinate system the *standard coordinate system* of $\mathbb{K}^n$.

Example 3.85 (Non-linear coordinates) In Linear Algebra we only consider linear coordinate systems, but in other areas of mathematics *non-linear coordinate systems* are also used. An example are the so-called *polar coordinates*

$$\rho : \mathbb{R}^2 \setminus \{0_{\mathbb{R}^2}\} \rightarrow (0, \infty) \times (-\pi, \pi] \subset \mathbb{R}^2, \quad \vec{x} \mapsto \begin{pmatrix} r \\ \phi \end{pmatrix} = \begin{pmatrix} \sqrt{(x_1)^2 + (x_2)^2} \\ \arg(\vec{x}) \end{pmatrix},$$

where $\arg(\vec{x}) = \arccos(x_1/r)$ for $x_2 \geq 0$ and $\arg(\vec{x}) = -\arccos(x_1/r)$ for $x_2 < 0$. Notice that the polar coordinates are only defined on $\mathbb{R}^2 \setminus \{0_{\mathbb{R}^2}\}$. For further details we refer to the Analysis module.

A convenient way to visualise a linear coordinate system $\varphi : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ is to consider the preimage $\varphi^{-1}(\mathcal{C})$ of the *standard coordinate grid*

$$(3.13) \quad \mathcal{C} = \{s\vec{e}_1 + k\vec{e}_2 \mid s \in \mathbb{R}, k \in \mathbb{Z}\} \cup \{k\vec{e}_1 + s\vec{e}_2 \mid s \in \mathbb{R}, k \in \mathbb{Z}\}$$

under φ . The first set in the union (3.13) of sets are the *horizontal coordinate lines* and the second set the *vertical coordinate lines*.

Example 3.86 (see [Figure 3.1](#)) The vector $\vec{v} = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$ has coordinates $\begin{pmatrix} 2 \\ 1 \end{pmatrix}$ with respect to the standard coordinate system of $\mathbb{R}^2$. The same vector has coordinates $\varphi(\vec{v}) = \begin{pmatrix} 4 \\ -1 \end{pmatrix}$ with respect to the coordinate system $\varphi\left(\begin{pmatrix} v_1 \\ v_2 \end{pmatrix}\right) = \begin{pmatrix} v_1 + 2v_2 \\ -v_1 + v_2 \end{pmatrix}$.

While $\mathbb{K}^n$ is equipped with the standard coordinate system, in an abstract vector space V there is no preferred linear coordinate system and a choice of linear coordinate system amounts to choosing a so-called ordered basis of V .

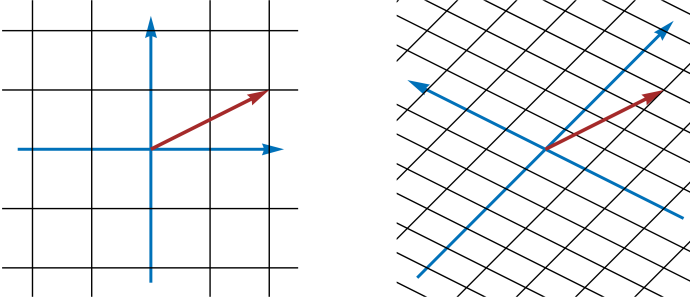


FIGURE 3.1. The coordinates of a vector with respect to different co-ordinate systems.

Definition 3.87 (Ordered basis) Let V be a finite dimensional $\mathbb{K}$ -vector space. An (ordered) n -tuple $\mathbf{b} = (v_1, \dots, v_n)$ of vectors from V is called an *ordered basis* of V if the set $\{v_1, \dots, v_n\}$ is a basis of V .

That there is a bijective correspondence between ordered bases of V and linear coordinate systems on V is a consequence of the following very important lemma which states in particular that two linear maps $f, g : V \rightarrow W$ are the same if and only if they agree on a basis of V .

Lemma 3.88 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces.

- (i) Suppose $f, g : V \rightarrow W$ are linear maps and $\mathbf{b} = (v_1, \dots, v_n)$ is an ordered basis of V . Then $f = g$ if and only if $f(v_i) = g(v_i)$ for all $1 \leq i \leq n$.
- (ii) If $\dim V = \dim W$ and $\mathbf{b} = (v_1, \dots, v_n)$ is an ordered basis of V and $\mathbf{c} = (w_1, \dots, w_n)$ an ordered basis of W , then there exists a unique isomorphism $f : V \rightarrow W$ such that $f(v_i) = w_i$ for all $1 \leq i \leq n$.

Proof (i) $\Rightarrow$ If $f = g$ then $f(v_i) = g(v_i)$ for all $1 \leq i \leq n$. $\Leftarrow$ Let $v \in V$. Since $\mathbf{b}$ is an ordered basis of V there exist unique scalars $s_1, \dots, s_n \in \mathbb{K}$ such that $v = \sum_{i=1}^n s_i v_i$. Using the linearity of f and g , we compute

$$f(v) = f\left(\sum_{i=1}^n s_i v_i\right) = \sum_{i=1}^n s_i f(v_i) = \sum_{i=1}^n s_i g(v_i) = g\left(\sum_{i=1}^n s_i v_i\right) = g(v)$$

so that $f = g$.

(ii) Let $v \in V$. Since $\{v_1, \dots, v_n\}$ is a basis of V there exist unique scalars $s_1, \dots, s_n$ such that $v = \sum_{i=1}^n s_i v_i$. We define $f(v) = \sum_{i=1}^n s_i w_i$, so that in particular $f(v_i) = w_i$ for $1 \leq i \leq n$. Since $\{w_1, \dots, w_n\}$ are linearly independent we have $f(v) = 0_W$ if and only if $s_1 = \dots = s_n = 0$, that is $v = 0_V$. Lemma 3.31 implies that f is injective and hence an isomorphism by Corollary 3.77. The uniqueness of f follows from (i). $\square$

Remark 3.89 Notice that Lemma 3.88 is wrong for maps that are not linear. Consider

$$f : \mathbb{R}^2 \rightarrow \mathbb{R}, \quad \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \mapsto x_1 x_2$$

and

$$g : \mathbb{R}^2 \rightarrow \mathbb{R} \quad \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \mapsto (x_1 - 1)(x_2 - 1).$$

Then $f(\vec{e}_1) = g(\vec{e}_1)$ and $f(\vec{e}_2) = g(\vec{e}_2)$, but $f \neq g$.

Given an ordered basis $\mathbf{b} = (v_1, \dots, v_n)$ of V , the previous lemma implies that there is a unique linear coordinate system $\beta : V \rightarrow \mathbb{K}^n$ such that

$$(3.14) \quad \beta(v_i) = \vec{e}_i$$

for $1 \leq i \leq n$, where $\{\vec{e}_1, \dots, \vec{e}_n\}$ denotes the standard basis of $\mathbb{K}^n$. Conversely, if $\beta : V \rightarrow \mathbb{K}^n$ is a linear coordinate system, we obtain an ordered basis of V

$$\mathbf{b} = (\beta^{-1}(\vec{e}_1), \dots, \beta^{-1}(\vec{e}_n))$$

and these assignments are inverse to each other. Notice that for all $v \in V$ we have

$$\beta(v) = \begin{pmatrix} s_1 \\ \vdots \\ s_n \end{pmatrix} \iff v = s_1 v_1 + \dots + s_n v_n.$$

Remark 3.90 (Notation) We will denote an ordered basis by an upright bold Roman letter, such as $\mathbf{b}$, $\mathbf{c}$, $\mathbf{d}$ or $\mathbf{e}$. We will denote the corresponding linear coordinate system by the corresponding bold Greek letter β, γ, δ or ε , respectively.

Example 3.91 Let $V = \mathbb{K}^3$ and $\mathbf{e} = (\vec{e}_1, \vec{e}_2, \vec{e}_3)$ denote the ordered standard basis. Then for all $\vec{x} = (x_i)_{1 \leq i \leq 3} \in \mathbb{K}^3$ we have

$$\varepsilon(\vec{x}) = \vec{x}.$$

where ε denotes the linear coordinate system corresponding to $\mathbf{e}$. Notice that ε is the standard coordinate system on $\mathbb{K}^n$. Considering instead the ordered basis $\mathbf{b} = (\vec{v}_1, \vec{v}_2, \vec{v}_3) = (\vec{e}_1 + \vec{e}_3, \vec{e}_3, \vec{e}_2 - \vec{e}_1)$, we obtain

$$\beta(\vec{x}) = \begin{pmatrix} x_1 + x_2 \\ x_3 - x_1 - x_2 \\ x_2 \end{pmatrix}$$

since

$$\vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = (x_1 + x_2) \underbrace{\begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}}_{=\vec{v}_1} + (x_3 - x_1 - x_2) \underbrace{\begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}}_{=\vec{v}_2} + x_2 \underbrace{\begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix}}_{=\vec{v}_3}.$$

Fixing linear coordinate systems – or equivalently ordered bases – on finite dimensional vector spaces V, W allows to describe each linear map $g : V \rightarrow W$ in terms of a matrix:

Definition 3.92 (Matrix representation of a linear map — Video) Let V, W be finite dimensional $\mathbb{K}$ -vector spaces, $\mathbf{b}$ an ordered basis of V and $\mathbf{c}$ an ordered basis of W . The matrix representation of a linear map $g : V \rightarrow W$ with respect to the ordered bases $\mathbf{b}$ and $\mathbf{c}$ is the unique matrix $\mathbf{M}(g, \mathbf{b}, \mathbf{c}) \in M_{m,n}(\mathbb{K})$ such that

$$f_{\mathbf{M}(g, \mathbf{b}, \mathbf{c})} = \gamma \circ g \circ \beta^{-1},$$

where β and γ denote the linear coordinate systems corresponding to $\mathbf{b}$ and $\mathbf{c}$, respectively.

The role of the different mappings can be summarised in terms of the following diagram:

$$\begin{array}{ccc} V & \xrightarrow{g} & W \\ \beta^{-1} \uparrow & & \downarrow \gamma \\ \mathbb{K}^n & \xrightarrow{f_{\mathbf{M}(g, \mathbf{b}, \mathbf{c})}} & \mathbb{K}^m \end{array}$$

In practise, we can compute the matrix representation of a linear map as follows:

Proposition 3.93 *Let V, W be finite dimensional $\mathbb{K}$ -vector spaces, $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V , $\mathbf{c} = (w_1, \dots, w_m)$ an ordered basis of W and $g : V \rightarrow W$ a linear map. Then there exist unique scalars $A_{ij} \in \mathbb{K}$, where $1 \leq i \leq m, 1 \leq j \leq n$ such that*

$$(3.15) \quad g(v_j) = \sum_{i=1}^m A_{ij} w_i, \quad 1 \leq j \leq n.$$

Furthermore, the matrix $\mathbf{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ satisfies

$$f_{\mathbf{A}} = \gamma \circ g \circ \beta^{-1}$$

and hence is the matrix representation of g with respect to the ordered bases $\mathbf{b}$ and $\mathbf{c}$.

Remark 3.94 Notice that we sum over the first index of A_{ij} in (3.15).

Proof of Proposition 3.93 For all $1 \leq j \leq n$ the vector $g(v_j)$ is an element of W and hence a linear combination of the vectors $\mathbf{c} = (w_1, \dots, w_m)$, as $\mathbf{c}$ is an ordered basis of W . We thus have scalars $A_{ij} \in \mathbb{K}$ with $1 \leq i \leq m, 1 \leq j \leq n$ such that $g(v_j) = \sum_{i=1}^m A_{ij} w_i$. If $\hat{A}_{ij} \in \mathbb{K}$ with $1 \leq i \leq m, 1 \leq j \leq n$ also satisfy $g(v_j) = \sum_{i=1}^m \hat{A}_{ij} w_i$, then subtracting the two equations gives

$$g(v_j) - g(v_j) = 0_W = \sum_{i=1}^m (A_{ij} - \hat{A}_{ij}) w_i$$

so that $0 = A_{ij} - \hat{A}_{ij}$ for $1 \leq i \leq m, 1 \leq j \leq n$, since the vectors $(w_1, \dots, w_m)$ are linearly independent. It follows that the scalars A_{ij} are unique.

We want to show that $f_{\mathbf{A}} \circ \beta = \gamma \circ g$. Using Lemma 3.88 it is sufficient to show that $(f_{\mathbf{A}} \circ \beta)(v_j) = (\gamma \circ g)(v_j)$ for $1 \leq j \leq n$. Let $\{\vec{e}_1, \dots, \vec{e}_n\}$ denote the standard basis of $\mathbb{K}^n$ so that $\beta(v_j) = \vec{e}_j$ and $\{\vec{d}_1, \dots, \vec{d}_m\}$ the standard basis of $\mathbb{K}^m$ so that $\gamma(w_i) = \vec{d}_i$. We compute

$$\begin{aligned} (f_{\mathbf{A}} \circ \beta)(v_j) &= f_{\mathbf{A}}(\vec{e}_j) = \mathbf{A} \vec{e}_j = \sum_{i=1}^m A_{ij} \vec{d}_i = \sum_{i=1}^m A_{ij} \gamma(w_i) = \gamma \left(\sum_{i=1}^m A_{ij} w_i \right) \\ &= \gamma(g(v_j)) = (\gamma \circ g)(v_j) \end{aligned}$$

where we have used the linearity of γ and (3.15). □

This all translates to a simple recipe for calculating the matrix representation of a linear map, which we now illustrate in some examples.

Example 3.95 Let $V = P_2(\mathbb{R})$ and $W = P_1(\mathbb{R})$ and $g = \frac{d}{dx}$. We consider the ordered basis $\mathbf{b} = (v_1, v_2, v_3) = ((1/2)(3x^2 - 1), x, 1)$ of V and $\mathbf{c} = (w_1, w_2) = (x, 1)$ of W .

(i) Compute the image under g of the elements v_i of the ordered basis $\mathbf{b}$.

$$g\left(\frac{1}{2}(3x^2 - 1)\right) = \frac{d}{dx}\left(\frac{1}{2}(3x^2 - 1)\right) = 3x$$

$$g(x) = \frac{d}{dx}(x) = 1$$

$$g(1) = \frac{d}{dx}(1) = 0.$$

(ii) Write the image vectors as linear combinations of the elements of the ordered basis $\mathbf{c}$.

$$\begin{aligned} 3x &= 3 \cdot w_1 + 0 \cdot w_2 \\ 1 &= 0 \cdot w_1 + 1 \cdot w_2 \\ 0 &= 0 \cdot w_1 + 0 \cdot w_2 \end{aligned} \tag{3.16}$$

(iii) Taking the transpose of the matrix of coefficients appearing in (3.16) gives the matrix representation

$$\mathbf{M}\left(\frac{d}{dx}, \mathbf{b}, \mathbf{c}\right) = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}.$$

of the linear map $g = \frac{d}{dx}$ with respect to the bases $\mathbf{b}, \mathbf{c}$.

Example 3.96 Let $\mathbf{e} = (\vec{e}_1, \dots, \vec{e}_n)$ and $\mathbf{d} = (\vec{d}_1, \dots, \vec{d}_m)$ denote the ordered standard basis of $\mathbb{K}^n$ and $\mathbb{K}^m$, respectively. Then for $\mathbf{A} \in M_{m,n}(\mathbb{K})$, we have

$$\mathbf{A} = \mathbf{M}(f_{\mathbf{A}}, \mathbf{e}, \mathbf{d}),$$

that is, the matrix representation of the mapping $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ with respect to the standard bases is simply the matrix $\mathbf{A}$. Indeed, we have

$$f_{\mathbf{A}}(\vec{e}_j) = \mathbf{A}\vec{e}_j = \begin{pmatrix} A_{1j} \\ \vdots \\ A_{mj} \end{pmatrix} = \sum_{i=1}^m A_{ij} \vec{d}_i.$$

Example 3.97 Let $\mathbf{e} = (\vec{e}_1, \vec{e}_2)$ denote the ordered standard basis of $\mathbb{R}^2$. Consider the matrix

$$\mathbf{A} = \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} = \mathbf{M}(f_{\mathbf{A}}, \mathbf{e}, \mathbf{e}).$$

We want to compute $\text{Mat}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b})$, where $\mathbf{b} = (\vec{v}_1, \vec{v}_2) = (\vec{e}_1 + \vec{e}_2, \vec{e}_2 - \vec{e}_1)$ is *not* the standard basis of $\mathbb{R}^2$. We obtain

$$f_{\mathbf{A}}(\vec{v}_1) = \mathbf{A}\vec{v}_1 = \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 6 \\ 6 \end{pmatrix} = 6 \cdot \vec{v}_1 + 0 \cdot \vec{v}_2$$

$$f_{\mathbf{A}}(\vec{v}_2) = \mathbf{A}\vec{v}_2 = \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} \begin{pmatrix} -1 \\ 1 \end{pmatrix} = \begin{pmatrix} -4 \\ 4 \end{pmatrix} = 0 \cdot \vec{v}_1 + 4 \cdot \vec{v}_2$$

Therefore, we have

$$\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \begin{pmatrix} 6 & 0 \\ 0 & 4 \end{pmatrix}.$$

Proposition 3.98 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces, $\mathbf{b}$ an ordered basis of V with corresponding linear coordinate system β , $\mathbf{c}$ an ordered basis of W with corresponding linear coordinate system γ and $g : V \rightarrow W$ a linear map. Then for all $v \in V$ we have

$$\gamma(g(v)) = \mathbf{M}(g, \mathbf{b}, \mathbf{c})\beta(v).$$

Proof By definition we have for all $\vec{x} \in \mathbb{K}^n$ and $\mathbf{A} \in M_{m,n}(\mathbb{K})$

$$\mathbf{A}\vec{x} = f_{\mathbf{A}}(\vec{x}).$$

Combining this with [Definition 3.92](#), we obtain for all $v \in V$

$$\mathbf{M}(g, \mathbf{b}, \mathbf{c})\beta(v) = f_{\mathbf{M}(g, \mathbf{b}, \mathbf{c})}(\beta(v)) = (\gamma \circ g \circ \beta^{-1})(\beta(v)) = \gamma(g(v)),$$

as claimed. □

Remark 3.99 Explicitly, [Proposition 3.98](#) states the following. Let $\mathbf{A} = \mathbf{M}(g, \mathbf{b}, \mathbf{c})$ and let $v \in V$. Since $\mathbf{b}$ is an ordered basis of V , there exist unique scalars $s_i \in \mathbb{K}$, $1 \leq i \leq n$ such that

$$v = s_1 v_1 + \cdots + s_n v_n.$$

Then we have

$$g(v) = t_1 w_1 + \cdots + t_m w_m,$$

where

$$\begin{pmatrix} t_1 \\ \vdots \\ t_m \end{pmatrix} = \mathbf{A} \begin{pmatrix} s_1 \\ \vdots \\ s_n \end{pmatrix}.$$

Example 3.100 ([Example 3.95](#) continued) With respect to the ordered basis $\mathbf{b} = (\frac{1}{2}(3x^2 - 1), x, 1)$, the polynomial $a_2 x^2 + a_1 x + a_0 \in V = P_2(\mathbb{R})$ is represented by the vector

$$\beta(a_2 x^2 + a_1 x + a_0) = \begin{pmatrix} \frac{2}{3}a_2 \\ a_1 \\ \frac{a_2}{3} + a_0 \end{pmatrix}$$

Indeed

$$a_2 x^2 + a_1 x + a_0 = \frac{2}{3}a_2 \left(\frac{1}{2}(3x^2 - 1) \right) + a_1 x + \left(\frac{a_2}{3} + a_0 \right) 1.$$

Computing $\mathbf{M}(\frac{d}{dx}, \mathbf{b}, \mathbf{c})\beta(a_2 x^2 + a_1 x + a_0)$ gives

$$\begin{pmatrix} 3 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} \frac{2}{3}a_2 \\ a_1 \\ \frac{a_2}{3} + a_0 \end{pmatrix} = \begin{pmatrix} 2a_2 \\ a_1 \end{pmatrix}$$

and this vector represents the polynomial $2a_2 \cdot x + a_1 \cdot 1 = \frac{d}{dx}(a_2 x^2 + a_1 x + a_0)$ with respect to the basis $\mathbf{c} = (x, 1)$ of $P_1(\mathbb{R})$.

As a corollary to [Proposition 3.93](#) we obtain:

Corollary 3.101 Let V_1, V_2, V_3 be finite dimensional $\mathbb{K}$ -vector spaces and $\mathbf{b}_i$ an ordered basis of V_i for $i = 1, 2, 3$. Let $g_1 : V_1 \rightarrow V_2$ and $g_2 : V_2 \rightarrow V_3$ be linear maps.

Then

$$\mathbf{M}(g_2 \circ g_1, \mathbf{b}_1, \mathbf{b}_3) = \mathbf{M}(g_2, \mathbf{b}_2, \mathbf{b}_3)\mathbf{M}(g_1, \mathbf{b}_1, \mathbf{b}_2).$$

Proof Let us write $\mathbf{C} = \mathbf{M}(g_2 \circ g_1, \mathbf{b}_1, \mathbf{b}_3)$ and $\mathbf{A}_1 = \mathbf{M}(g_1, \mathbf{b}_1, \mathbf{b}_2)$ as well as $\mathbf{A}_2 = \mathbf{M}(g_2, \mathbf{b}_2, \mathbf{b}_3)$. Using [Proposition 2.20](#) and [Theorem 2.21](#) it suffices to show that $f_{\mathbf{C}} = f_{\mathbf{A}_2\mathbf{A}_1} = f_{\mathbf{A}_2} \circ f_{\mathbf{A}_1}$. Now [Proposition 3.93](#) gives

$$f_{\mathbf{A}_2} \circ f_{\mathbf{A}_1} = \beta_3 \circ g_2 \circ \beta_2^{-1} \circ \beta_2 \circ g_1 \circ \beta_1^{-1} = \beta_3 \circ g_2 \circ g_1 \circ \beta_1^{-1} = f_{\mathbf{C}}.$$

□

Proposition 3.102 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces, $\mathbf{b}$ an ordered basis of V and $\mathbf{c}$ an ordered basis of W . A linear map $g : V \rightarrow W$ is bijective if and only if $\mathbf{M}(g, \mathbf{b}, \mathbf{c})$ is invertible. Moreover, in the case where g is bijective we have

$$\mathbf{M}(g^{-1}, \mathbf{c}, \mathbf{b}) = (\mathbf{M}(g, \mathbf{b}, \mathbf{c}))^{-1}.$$

Proof Let $n = \dim(V)$ and $m = \dim(W)$.

$\Rightarrow$ Let $g : V \rightarrow W$ be bijective so that g is an isomorphism and hence $n = \dim(V) = \dim(W) = m$ by [Proposition 3.80](#). Then [Corollary 3.101](#) gives

$$\mathbf{M}(g^{-1}, \mathbf{c}, \mathbf{b})\mathbf{M}(g, \mathbf{b}, \mathbf{c}) = \mathbf{M}(g^{-1} \circ g, \mathbf{b}, \mathbf{b}) = \mathbf{M}(\text{Id}_V, \mathbf{b}, \mathbf{b}) = \mathbf{1}_n$$

and

$$\mathbf{M}(g, \mathbf{b}, \mathbf{c})\mathbf{M}(g^{-1}, \mathbf{c}, \mathbf{b}) = \mathbf{M}(g \circ g^{-1}, \mathbf{c}, \mathbf{c}) = \mathbf{M}(\text{Id}_W, \mathbf{c}, \mathbf{c}) = \mathbf{1}_n$$

so that $\mathbf{M}(g, \mathbf{b}, \mathbf{c})$ is invertible with inverse $\mathbf{M}(g^{-1}, \mathbf{c}, \mathbf{b})$.

$\Leftarrow$ Conversely suppose $\mathbf{A} = \mathbf{M}(g, \mathbf{b}, \mathbf{c})$ is invertible with inverse $\mathbf{A}^{-1}$. It follows that $n = m$ by [Corollary 3.81](#). We consider $h = \beta^{-1} \circ f_{\mathbf{A}^{-1}} \circ \gamma : W \rightarrow V$ and since $f_{\mathbf{A}} = \gamma \circ g \circ \beta^{-1}$ by [Proposition 3.93](#), we have

$$g \circ h = \gamma^{-1} \circ f_{\mathbf{A}} \circ \beta \circ \beta^{-1} \circ f_{\mathbf{A}^{-1}} \circ \gamma = \gamma^{-1} \circ f_{\mathbf{A}\mathbf{A}^{-1}} \circ \gamma = \text{Id}_W.$$

Likewise, we have

$$h \circ g = \beta^{-1} \circ f_{\mathbf{A}^{-1}} \circ \gamma \circ \gamma^{-1} \circ f_{\mathbf{A}} \circ \beta = \beta^{-1} \circ f_{\mathbf{A}^{-1}\mathbf{A}} \circ \beta = \text{Id}_V,$$

showing that g admits an inverse mapping $h : W \rightarrow V$ and hence g is bijective. □

Recall that a mapping $f : \mathcal{X} \rightarrow \mathcal{Y}$ between sets $\mathcal{X}, \mathcal{Y}$ is said to admit a *left inverse* if there exists a mapping $g : \mathcal{Y} \rightarrow \mathcal{X}$ such that $g \circ f = \text{Id}_{\mathcal{X}}$. Likewise, a *right inverse* is a mapping $h : \mathcal{Y} \rightarrow \mathcal{X}$ such that $f \circ h = \text{Id}_{\mathcal{Y}}$.

We now have:

Proposition 3.103 Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{K})$ a square matrix. Then the following statements are equivalent:

- (i) The matrix $\mathbf{A}$ admits a left inverse, that is, a matrix $\mathbf{B} \in M_{n,n}(\mathbb{K})$ such that $\mathbf{BA} = \mathbf{1}_n$;
- (ii) The matrix $\mathbf{A}$ admits a right inverse, that is, a matrix $\mathbf{B} \in M_{n,n}(\mathbb{K})$ such that $\mathbf{AB} = \mathbf{1}_n$;
- (iii) The matrix $\mathbf{A}$ is invertible.

Proof By the definition of the invertability of a matrix, (iii) implies both (i) and (ii).

(i) $\Rightarrow$ (iii) Since $\mathbf{BA} = \mathbf{1}_n$ we have $f_{\mathbf{B}} \circ f_{\mathbf{A}} = f_{\mathbf{1}_n} = \text{Id}_{\mathbb{K}^n}$ by [Theorem 2.21](#) and hence $f_{\mathbf{B}}$ is a left inverse for $f_{\mathbf{A}}$. Therefore, by the above exercise, $f_{\mathbf{A}}$ is injective. [Corollary 3.77](#) implies that $f_{\mathbf{A}}$ is also bijective. Denoting the ordered standard basis of $\mathbb{K}^n$ by $\mathbf{e}$, we have $\mathbf{M}(f_{\mathbf{A}}, \mathbf{e}, \mathbf{e}) = \mathbf{A}$ and hence [Proposition 3.102](#) implies that $\mathbf{A}$ is invertible.

(ii) $\Rightarrow$ (iii) is completely analogous to (i) $\Rightarrow$ (iii). □

3.7.1 Change of basis

It is natural to ask how the choice of bases affects the matrix representation of a linear map.

Definition 3.104 (Change of basis matrix) Let V be a finite dimensional $\mathbb{K}$ -vector space and $\mathbf{b}, \mathbf{b}'$ be ordered bases of V with corresponding linear coordinate systems β, β' . The change of basis matrix from $\mathbf{b}$ to $\mathbf{b}'$ is the matrix $\mathbf{C} \in M_{n,n}(\mathbb{K})$ satisfying

$$f_{\mathbf{C}} = \beta' \circ \beta^{-1}$$

We will write $\mathbf{C}(\mathbf{b}, \mathbf{b}')$ for the change of basis matrix from $\mathbf{b}$ to $\mathbf{b}'$.

Remark 3.105 Notice that by definition

$$\mathbf{C}(\mathbf{b}, \mathbf{b}') = \mathbf{M}(\text{Id}_V, \mathbf{b}, \mathbf{b}').$$

Since the identity map $\text{Id}_V : V \rightarrow V$ is bijective with inverse $(\text{Id}_V)^{-1} = \text{Id}_V$, [Proposition 3.102](#) implies that the change of basis matrix $\mathbf{C}(\mathbf{b}, \mathbf{b}')$ is invertible with inverse

$$\mathbf{C}(\mathbf{b}, \mathbf{b}')^{-1} = \mathbf{C}(\mathbf{b}', \mathbf{b}).$$

Example 3.106 Let $V = \mathbb{R}^2$ and $\mathbf{e} = (\vec{e}_1, \vec{e}_2)$ be the ordered standard basis and $\mathbf{b} = (\vec{v}_1, \vec{v}_2) = (\vec{e}_1 + \vec{e}_2, \vec{e}_2 - \vec{e}_1)$ another ordered basis. According to the recipe mentioned in [Example 3.95](#), if we want to compute $\mathbf{C}(\mathbf{e}, \mathbf{b})$ we simply need to write each vector of $\mathbf{e}$ as a linear combination of the elements of $\mathbf{b}$. The transpose of the resulting coefficient matrix is then $\mathbf{C}(\mathbf{e}, \mathbf{b})$. We obtain

$$\begin{aligned}\vec{e}_1 &= \frac{1}{2}\vec{v}_1 - \frac{1}{2}\vec{v}_2, \\ \vec{e}_2 &= \frac{1}{2}\vec{v}_1 + \frac{1}{2}\vec{v}_2,\end{aligned}$$

so that

$$\mathbf{C}(\mathbf{e}, \mathbf{b}) = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ -\frac{1}{2} & \frac{1}{2} \end{pmatrix}.$$

Reversing the role of $\mathbf{e}$ and $\mathbf{b}$ gives $\mathbf{C}(\mathbf{b}, \mathbf{e})$

$$\begin{aligned}\vec{v}_1 &= 1\vec{e}_1 + 1\vec{e}_2, \\ \vec{v}_2 &= -1\vec{e}_1 + 1\vec{e}_2,\end{aligned}$$

so that

$$\mathbf{C}(\mathbf{b}, \mathbf{e}) = \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix}.$$

Notice that indeed we have

$$\mathbf{C}(\mathbf{e}, \mathbf{b})\mathbf{C}(\mathbf{b}, \mathbf{e}) = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ -\frac{1}{2} & \frac{1}{2} \end{pmatrix} \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$$

so that $\mathbf{C}(\mathbf{e}, \mathbf{b})^{-1} = \mathbf{C}(\mathbf{b}, \mathbf{e})$.

Theorem 3.107 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $\mathbf{b}, \mathbf{b}'$ ordered bases of V and $\mathbf{c}, \mathbf{c}'$ ordered bases of W . Let $g : V \rightarrow W$ be a linear map. Then we have

$$\mathbf{M}(g, \mathbf{b}', \mathbf{c}') = \mathbf{C}(\mathbf{c}, \mathbf{c}')\mathbf{M}(g, \mathbf{b}, \mathbf{c})\mathbf{C}(\mathbf{b}', \mathbf{b})$$

In particular, for a linear map $g : V \rightarrow V$ we have

$$\mathbf{M}(g, \mathbf{b}', \mathbf{b}') = \mathbf{C}\mathbf{M}(g, \mathbf{b}, \mathbf{b})\mathbf{C}^{-1},$$

where we write $\mathbf{C} = \mathbf{C}(\mathbf{b}, \mathbf{b}')$.

Proof We write $\mathbf{A} = \mathbf{M}(g, \mathbf{b}, \mathbf{c})$ and $\mathbf{B} = \mathbf{M}(g, \mathbf{b}', \mathbf{c}')$ and $\mathbf{C} = \mathbf{C}(\mathbf{b}, \mathbf{b}')$ and $\mathbf{D} = \mathbf{C}(\mathbf{c}, \mathbf{c}')$. By [Remark 3.105](#) we have $\mathbf{C}^{-1} = \mathbf{C}(\mathbf{b}', \mathbf{b})$, hence applying [Proposition 2.20](#) and [Theorem 2.21](#) and [Corollary 2.22](#), we need to show that

$$f_{\mathbf{B}} = f_{\mathbf{D}} \circ f_{\mathbf{A}} \circ f_{\mathbf{C}^{-1}}.$$

By [Definition 3.92](#) we have

$$\begin{aligned} f_{\mathbf{A}} &= \gamma \circ g \circ \beta^{-1}, \\ f_{\mathbf{B}} &= \gamma' \circ g \circ (\beta')^{-1} \end{aligned}$$

and by [Definition 3.104](#) we have

$$\begin{aligned} f_{\mathbf{C}^{-1}} &= \beta \circ (\beta')^{-1}, \\ f_{\mathbf{D}} &= \gamma' \circ \gamma^{-1}. \end{aligned}$$

Hence we obtain

$$f_{\mathbf{D}} \circ f_{\mathbf{A}} \circ f_{\mathbf{C}^{-1}} = \gamma' \circ \gamma^{-1} \circ \gamma \circ g \circ \beta^{-1} \circ \beta \circ (\beta')^{-1} = \gamma' \circ g \circ (\beta')^{-1} = f_{\mathbf{B}},$$

as claimed. The second statement follows again by applying [Remark 3.105](#). $\square$

Example 3.108 ([Example 3.97](#) and [Example 3.106](#) continued) Let $\mathbf{e} = (\vec{e}_1, \vec{e}_2)$ denote the ordered standard basis of $\mathbb{R}^2$ and

$$\mathbf{A} = \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} = \mathbf{M}(f_{\mathbf{A}}, \mathbf{e}, \mathbf{e}).$$

Let $\mathbf{b} = (\vec{e}_1 + \vec{e}_2, \vec{e}_2 - \vec{e}_1)$. We computed that

$$\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \begin{pmatrix} 6 & 0 \\ 0 & 4 \end{pmatrix}$$

as well as

$$\mathbf{C}(\mathbf{e}, \mathbf{b}) = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ -\frac{1}{2} & \frac{1}{2} \end{pmatrix} \quad \text{and} \quad \mathbf{C}(\mathbf{b}, \mathbf{e}) = \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix}.$$

According to [Theorem 3.107](#) we must have

$$\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \mathbf{C}(\mathbf{e}, \mathbf{b})\mathbf{M}(f_{\mathbf{A}}, \mathbf{e}, \mathbf{e})\mathbf{C}(\mathbf{b}, \mathbf{e})$$

and indeed

$$\begin{pmatrix} 6 & 0 \\ 0 & 4 \end{pmatrix} = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ -\frac{1}{2} & \frac{1}{2} \end{pmatrix} \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix}.$$

Finally, we observe that every invertible matrix can be realised as a change of basis matrix:

Lemma 3.109 *Let V be a finite dimensional $\mathbb{K}$ -vector space, $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V and $\mathbf{C} \in M_{n,n}(\mathbb{K})$ an invertible $n \times n$ -matrix. Define $v'_j = \sum_{i=1}^n C_{ij}v_i$ for $1 \leq j \leq n$. Then $\mathbf{b}' = (v'_1, \dots, v'_n)$ is an ordered basis of V and $\mathbf{C}(\mathbf{b}', \mathbf{b}) = \mathbf{C}$.*

Proof It is sufficient to prove that the vectors $\{v'_1, \dots, v'_n\}$ are linearly independent. Indeed, if they are linearly independent, then they span a subspace U of dimension n and Proposition 3.74 implies that $U = V$, so that $\mathbf{b}'$ is an ordered basis of V . Suppose we have scalars $s_1, \dots, s_n$ such that

$$0_V = \sum_{j=1}^n s_j v'_j = \sum_{j=1}^n \sum_{i=1}^n s_j C_{ij} v_i = \sum_{i=1}^n \left(\sum_{j=1}^n C_{ij} s_j \right) v_i.$$

Since $\{v_1, \dots, v_n\}$ is a basis of V we must have $\sum_{j=1}^n C_{ij} s_j = 0$ for all $i = 1, \dots, n$. In matrix notation this is equivalent to the condition $\mathbf{C}\vec{s} = 0_{\mathbb{K}^n}$, where $\vec{s} = (s_i)_{1 \leq i \leq n}$. Since $\mathbf{C}$ is invertible, we can multiply this last equation from the left with $\mathbf{C}^{-1}$ to obtain $\mathbf{C}^{-1}\mathbf{C}\vec{s} = \mathbf{C}^{-1}0_{\mathbb{K}^n}$ which is equivalent to $\vec{s} = 0_{\mathbb{K}^n}$. It follows that $\mathbf{b}'$ is an ordered basis of V . By definition we have $\mathbf{C}(\mathbf{b}', \mathbf{b}) = \mathbf{C}$. $\square$

Exercises

Exercise 3.110 Let $\text{Id}_V : V \rightarrow V$ denote the identity mapping of the finite dimensional $\mathbb{K}$ -vector space V and let $\mathbf{b} = (v_1, \dots, v_n)$ be any ordered basis of V . Show that $\mathbf{M}(\text{Id}_V, \mathbf{b}, \mathbf{b}) = \mathbf{1}_n$.

Exercise 3.111 Show that $f : \mathcal{X} \rightarrow \mathcal{Y}$ admits a left inverse if and only if f is injective and that $f : \mathcal{X} \rightarrow \mathcal{Y}$ admits a right inverse if and only if f is surjective.

Exercise 3.112 Let V be a finite dimensional $\mathbb{K}$ -vector space and $\mathbf{b}, \mathbf{b}'$ be ordered bases of V . Show that for all $v \in V$ we have

$$\beta'(v) = \mathbf{C}(\mathbf{b}, \mathbf{b}')\beta(v).$$

Applications of Gaussian elimination

4.1 Gaussian elimination

WEEK 7

In the Algorithmics module M01 you learned how to use Gaussian elimination to solve a system of equations of the form

$$(4.1) \quad \mathbf{A}\vec{x} = \vec{b}$$

for some given matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$, vector $\vec{b} \in \mathbb{K}^m$ and unknown $\vec{x} \in \mathbb{K}^n$. Many concrete problems in Linear Algebra lead to systems of the form (4.1). A few sample problems that can be solved with Gaussian elimination are discussed below.

Solving equation of the type (4.1) hinges on the elementary observation that a vector $\vec{x} \in \mathbb{K}^n$ solves $\mathbf{A}\vec{x} = \vec{b}$ if and only if it solves $\mathbf{B}\mathbf{A}\vec{x} = \mathbf{B}\vec{b}$, where $\mathbf{B} \in M_{m,m}(\mathbb{K})$ is any invertible m -by- m matrix.

In the Gaussian elimination algorithm, the matrix $\mathbf{B}$ is chosen among three types of matrices:

Definition 4.1 (Elementary matrices — Video) Let $m \in \mathbb{N}$. The elementary matrices of size m are the square matrices

$$\mathbf{L}_{k,l}(s) = \mathbf{1}_m + s\mathbf{E}_{k,l},$$

$$\mathbf{D}_k(s) = \mathbf{1}_m + (s - 1)\mathbf{E}_{k,k},$$

$$\mathbf{P}_{k,l} = \mathbf{1}_m - \mathbf{E}_{k,k} - \mathbf{E}_{l,l} + \mathbf{E}_{k,l} + \mathbf{E}_{l,k},$$

where $1 \leq k, l \leq m$ with $k \neq l$, $\mathbf{E}_{k,l} \in M_{m,m}(\mathbb{K})$ and $s \in \mathbb{K}$ with $s \neq 0$.

Example 4.2 For $m = 4$ we have for instance

$$\mathbf{L}_{2,3}(s) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & s & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}, \quad \mathbf{D}_4(s) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & s \end{pmatrix}$$

and

$$\mathbf{P}_{2,4} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \end{pmatrix}.$$

As an exercise in matrix multiplication, we compute the effect of left multiplication with elementary matrices.

For $\mathbf{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n} \in M_{m,n}(\mathbb{K})$, we obtain

$$[\mathbf{L}_{k,l}(s)\mathbf{A}]_{ij} = \sum_{r=1}^m (\delta_{ir} + s\delta_{ik}\delta_{lr}) A_{rj} = A_{ij} + s\delta_{ik}A_{lj} = \begin{cases} A_{ij} + sA_{lj} & i = k \\ A_{ij} & i \neq k \end{cases},$$

where we use that $[\mathbf{1}_m]_{ir} = \delta_{ir}$ and $[\mathbf{E}_{k,l}]_{ir} = \delta_{ik}\delta_{lr}$. Therefore, multiplying the matrix $\mathbf{A}$ with $\mathbf{L}_{k,l}(s)$ from the left, adds s times the l -th row of $\mathbf{A}$ to the k -th row of $\mathbf{A}$ and leaves $\mathbf{A}$ unchanged otherwise.

Likewise, we obtain

$$[\mathbf{D}_k(s)\mathbf{A}]_{ij} = \sum_{r=1}^m (\delta_{ir} + (s-1)\delta_{ik}\delta_{kr}) A_{rj} = \begin{cases} sA_{ij} & i = k \\ A_{ij} & i \neq k \end{cases}.$$

Therefore, multiplying the matrix $\mathbf{A}$ with $\mathbf{D}_k(s)$ from the left, multiplies the k -th row of $\mathbf{A}$ with s and leaves $\mathbf{A}$ unchanged otherwise.

Finally,

$$\begin{aligned} [\mathbf{P}_{k,l}\mathbf{A}]_{ij} &= \sum_{r=1}^m (\delta_{ir} - \delta_{ik}\delta_{kr} - \delta_{il}\delta_{lr} + \delta_{ik}\delta_{lr} + \delta_{il}\delta_{rk}) A_{rj} \\ &= A_{ij} - \delta_{ik}A_{kj} - \delta_{il}A_{lj} + \delta_{ik}A_{lj} + \delta_{il}A_{kj} \\ &= A_{ij} + \delta_{ik}(A_{lj} - A_{kj}) + \delta_{il}(A_{kj} - A_{lj}) = \begin{cases} A_{lj} & i = k \\ A_{kj} & i = l \\ A_{ij} & i \neq k, i \neq l \end{cases}. \end{aligned}$$

Therefore, multiplying the matrix $\mathbf{A}$ with $\mathbf{P}_{k,l}$ from the left, swaps the k -th row of $\mathbf{A}$ with the l -th row of $\mathbf{A}$ and leaves $\mathbf{A}$ unchanged otherwise.

These calculations immediately imply:

Proposition 4.3 *The elementary matrices are invertible with*

$$\mathbf{L}_{k,l}(s)^{-1} = \mathbf{L}_{k,l}(-s) \quad \text{and} \quad \mathbf{D}_k(s)^{-1} = \mathbf{D}_k(1/s) \quad \text{and} \quad (\mathbf{P}_{k,l})^{-1} = \mathbf{P}_{k,l}.$$

The sceptical reader may also verify this fact by direct computation with the help of the following lemma:

Lemma 4.4 *Let $m \in \mathbb{N}$. For $1 \leq k, l, p, q \leq m$, we have*

$$\mathbf{E}_{k,l}\mathbf{E}_{p,q} = \begin{cases} \mathbf{E}_{k,q} & p = l \\ \mathbf{0}_{m,m} & p \neq l \end{cases}$$

Proof By definition, we have

$$\mathbf{E}_{k,l}\mathbf{E}_{p,q} = \left(\sum_{r=1}^m \delta_{ik}\delta_{lr}\delta_{rp}\delta_{qj} \right)_{1 \leq i,j \leq m} = \delta_{lp} (\delta_{ik}\delta_{qj})_{1 \leq i,j \leq m} = \begin{cases} \mathbf{E}_{k,q} & p = l \\ \mathbf{0}_{m,m} & p \neq l \end{cases}.$$

□

For each row in a matrix, if the row does not consist of zeros only, then the leftmost nonzero entry is called the leading coefficient of that row.

Definition 4.5 (Row echelon form) A matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ is said to be in *row echelon form* (REF) if

- all rows consisting of only zeros are at the bottom;
- the leading coefficient of a nonzero row is always strictly to the right of the leading coefficient of the row above it.

The matrix $\mathbf{A}$ is said to be in *reduced row echelon form* (rREF) if furthermore

- all of the leading coefficients are equal to 1;
- in every column containing a leading coefficient, all of the other entries in that column are zero.

Gaussian elimination from the Algorithmics module M01 implies the following statement:

Theorem 4.6 (Gauss–Jordan elimination) *Let $\mathbf{A} \in M_{m,n}(\mathbb{K})$ then there exists $N \in \mathbb{N}$ and an N -tuple of elementary matrices $(\mathbf{B}_1, \dots, \mathbf{B}_N)$ such that the matrix $\mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1 \mathbf{A}$ is in reduced row echelon form.*

Proof Applying Gaussian elimination implies the existence of $\hat{N} \in \mathbb{N}$ and elementary matrices $\mathbf{B}_1, \dots, \mathbf{B}_{\hat{N}}$ so that $\mathbf{B}_{\hat{N}} \mathbf{B}_{\hat{N}-1} \cdots \mathbf{B}_2 \mathbf{B}_1 \mathbf{A}$ is REF. After possibly further multiplying this matrix from the left with elementary matrices of the type $\mathbf{D}_k(s)$, we can assume that all leading coefficients are 1. By choosing suitable left multiplications with matrices of the type $\mathbf{L}_{k,l}(s)$, we find a natural number $N \geq \hat{N}$ and elementary matrices $(\mathbf{B}_1, \dots, \mathbf{B}_N)$ so that $\mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1 \mathbf{A}$ is in reduced row echelon form. $\square$

4.2 Applications

4.2.1 Compute the inverse of a matrix

An algorithm using Gaussian elimination for computing the inverse of an invertible matrix relies on the following fact:

Proposition 4.7 *Let $\mathbf{A} \in M_{n,n}(\mathbb{K})$ be a square matrix. Then the following statements are equivalent:*

- (i) $\mathbf{A}$ is invertible;
- (ii) the row vectors of $\mathbf{A}$ are linearly independent;
- (iii) the column vectors of $\mathbf{A}$ are linearly independent.

Proof Part of an exercise sheet. $\square$

Suppose the matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ is invertible. Applying Gauss–Jordan elimination to $\mathbf{A}$, we cannot encounter a zero row, since the occurrence of a zero row corresponds to a non-trivial linear combination of row vectors which gives the zero vector. This is excluded by the above proposition. Having no zero row vectors, the Gauss–Jordan elimination applied to $\mathbf{A}$ must give the identity matrix $\mathbf{1}_n$. Thus we can find a sequence of elementary matrices $\mathbf{B}_1, \dots, \mathbf{B}_N$, $N \in \mathbb{N}$, so that

$$\mathbf{1}_n = \mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1 \mathbf{A}.$$

In other words, $\mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1$ is the inverse of $\mathbf{A}$. This gives the following recipe for computing the inverse of $\mathbf{A}$:

We write the matrix $\mathbf{A}$ and $\mathbf{1}_n$ next to each other, say $\mathbf{A}$ on the left and $\mathbf{1}_n$ on the right. We then perform Gauss–Jordan elimination on $\mathbf{A}$. At each step, we also perform the Gauss–Jordan elimination step to the matrix on the right. Once Gauss–Jordan elimination terminates, we thus obtain $\mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1 \mathbf{A}$ on the left and $\mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1 \mathbf{1}_n$ on the right. But since $\mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1 \mathbf{1}_n = \mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1$ (notice the absence of $\mathbf{1}_n$ after the equality sign), the right hand side is the inverse of $\mathbf{A}$.

Example 4.8 (Inverse of a matrix — [Video](#)) We want to compute the inverse of

$$\mathbf{A} = \begin{pmatrix} 1 & -2 \\ -3 & 4 \end{pmatrix}.$$

Write

$$\left(\begin{array}{cc|cc} 1 & -2 & 1 & 0 \\ -3 & 4 & 0 & 1 \end{array} \right).$$

Adding 3-times the first row to the second row gives

$$\left(\begin{array}{cc|cc} 1 & -2 & 1 & 0 \\ 0 & -2 & 3 & 1 \end{array} \right).$$

Dividing the second row by -2 gives

$$\left(\begin{array}{cc|cc} 1 & -2 & 1 & 0 \\ 0 & 1 & -\frac{3}{2} & -\frac{1}{2} \end{array} \right).$$

Finally, adding the second row twice to the first row gives

$$\left(\begin{array}{cc|cc} 1 & 0 & -2 & -1 \\ 0 & 1 & -\frac{3}{2} & -\frac{1}{2} \end{array} \right),$$

so that

$$\mathbf{A}^{-1} = \begin{pmatrix} -2 & -1 \\ -\frac{3}{2} & -\frac{1}{2} \end{pmatrix}.$$

4.2.2 Compute a basis of a subspace

Gaussian elimination can also be used to compute a basis for a vector subspace U of a finite dimensional $\mathbb{K}$ -vector space V . We assume that $U = \text{span}\{v_1, \dots, v_k\}$ for some vectors $v_i \in V$, $1 \leq i \leq k$. We assume that $\dim U \geq 1$ so that not all vectors are the zero vector.

We first consider the special case where V is the space $\mathbb{K}_n$ of row vectors of length n and with entries in $\mathbb{K}$. Recall that we denote the row vectors by small Greek letters. We write $\mathbb{K}_n^m$ for the m -fold Cartesian product $(\mathbb{K}_n)^m$ of $\mathbb{K}_n$. Clearly, we have a bijective mapping

$$\Omega : \mathbb{K}_n^m \rightarrow M_{m,n}(\mathbb{K}), \quad (\vec{v}_1, \dots, \vec{v}_m) \mapsto \begin{pmatrix} \vec{v}_1 \\ \vdots \\ \vec{v}_m \end{pmatrix}$$

which simply writes the row vectors $(\vec{v}_1, \dots, \vec{v}_m)$ into a matrix with the k -th row vector from the m -tuple of row vectors becoming the k -th row of the matrix.

Example 4.9

$$\Omega((1 \ 2 \ 3), (4 \ 5 \ 6)) = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix}.$$

We have

$$\mathbf{L}_{k,l}(s)\Omega(\vec{v}_1, \dots, \vec{v}_m) = \Omega(\vec{v}_1, \dots, \vec{v}_{k-1}, \vec{v}_k + s\vec{v}_l, \vec{v}_{k+1}, \dots, \vec{v}_m),$$

$$\mathbf{D}_k(s)\Omega(\vec{v}_1, \dots, \vec{v}_m) = \Omega(\vec{v}_1, \dots, \vec{v}_{k-1}, s\vec{v}_k, \vec{v}_{k+1}, \dots, \vec{v}_m),$$

$$\mathbf{P}_{k,l}\Omega(\vec{v}_1, \dots, \vec{v}_m) = \Omega(\vec{v}_1, \dots, \vec{v}_{k-1}, \vec{v}_l, \vec{v}_{k+1}, \dots, \vec{v}_{l-1}, \vec{v}_k, \vec{v}_{l+1}, \dots, \vec{v}_m).$$

Notice that all these operations do not change the span of the vectors $\vec{v}_1, \dots, \vec{v}_m$. More precisely, if $(\vec{v}_1, \dots, \vec{v}_m)$ is an n -tuple of row vectors and if $\Omega(\vec{\omega}_1, \dots, \vec{\omega}_m) = \mathbf{B}\Omega(\vec{v}_1, \dots, \vec{v}_m)$ for some elementary matrix $\mathbf{B}$, then

$$\text{span}\{\vec{v}_1, \dots, \vec{v}_m\} = \text{span}\{\vec{\omega}_1, \dots, \vec{\omega}_m\}.$$

Applying Gaussian elimination to the matrix $\Omega(\vec{v}_1, \dots, \vec{v}_m)$ gives a list of elementary matrices $\mathbf{B}_1, \dots, \mathbf{B}_N$ such that

$$\mathbf{B}_N \mathbf{B}_{N-1} \cdots \mathbf{B}_2 \mathbf{B}_1 \Omega(\vec{v}_1, \dots, \vec{v}_m) = \Omega(\vec{\omega}_1, \dots, \vec{\omega}_r, 0_{\mathbb{K}_n}, \dots, 0_{\mathbb{K}_n})$$

where $1 \leq r \leq m$ and $0_{\mathbb{K}_n}$ denotes the zero vector in $\mathbb{K}_n$. By construction, the matrix $\mathbf{A} = \Omega(\vec{\omega}_1, \dots, \vec{\omega}_r, 0_{\mathbb{K}_n}, \dots, 0_{\mathbb{K}_n})$ is REF. Since the leading coefficient of $\vec{\omega}_i$ is always strictly to the right of the leading coefficient of $\vec{\omega}_{i-1}$, it follows that the vectors $\vec{\omega}_1, \dots, \vec{\omega}_r$ are linearly independent. Therefore, a basis of $\text{span}\{\vec{v}_1, \dots, \vec{v}_m\}$ is given by $\{\vec{\omega}_1, \dots, \vec{\omega}_r\}$.

The general case can be treated with the help of the following facts:

Proposition 4.10 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $\Phi : V \rightarrow W$ an isomorphism. Then $\mathcal{S} \subset V$ is a basis of V if and only if $\Phi(\mathcal{S})$ is a basis of W .

Proof $\Rightarrow$ Since $\mathcal{S}$ is a basis, the set $\mathcal{S}$ is linearly independent and since Φ is injective, so is $\Phi(\mathcal{S})$ by [Lemma 3.56](#). Since $\mathcal{S}$ is a basis, $\mathcal{S}$ is a generating set and since Φ is surjective, the subset $\Phi(\mathcal{S}) \subset W$ is a generating set for W by [Lemma 3.46](#).

$\Leftarrow$ We apply the above implication to $\Phi^{-1} : W \rightarrow V$ and the basis $\Phi(\mathcal{S}) \subset W$. □

Corollary 4.11 Let $\hat{V}, \hat{W}$ be finite dimensional $\mathbb{K}$ -vector spaces, $\Theta : \hat{V} \rightarrow \hat{W}$ an isomorphism and $U \subset \hat{V}$ a vector subspace. Then $\mathcal{S} \subset U$ is a basis of U if and only if $\Theta(\mathcal{S})$ is a basis of $\Theta(U)$.

Proof Apply [Proposition 4.10](#) to the vector space $V = U$, the vector space $W = \Theta(U)$ and the isomorphism $\Phi = \Theta|_U : V \rightarrow W$. □

We now describe a recipe to treat the general case of a subset $U = \text{span}\{v_1, \dots, v_m\}$ of a finite dimensional $\mathbb{K}$ -vector space V :

- (i) Fix an isomorphism $\Phi : V \rightarrow \mathbb{K}_n$ and write $\vec{v}_i = \Phi(v_i)$ for $1 \leq i \leq m$.
- (ii) Apply Gaussian elimination to the matrix $\Omega(\vec{v}_1, \dots, \vec{v}_m)$ to obtain a set of new vectors $(\vec{\omega}_1, \dots, \vec{\omega}_r, 0_{\mathbb{K}_n}, \dots, 0_{\mathbb{K}_n})$ for some $r \in \mathbb{N}$.
- (iii) Apply the inverse isomorphism Φ^{-1} to the obtained list of vectors. This gives the desired basis $\{\Phi^{-1}(\vec{\omega}_1), \dots, \Phi^{-1}(\vec{\omega}_r)\}$ of U .

Example 4.12 (Basis of a subspace — [Video](#)) Let $V = P_3(\mathbb{R})$ so that $\dim(V) = 4$ and

$$U = \text{span}\{x^3 + 2x^2 - x, 4x^3 + 8x^2 - 4x - 3, x^2 + 3x + 4, 2x^3 + 5x + x + 4\}.$$

We want to compute a basis of U . We choose the isomorphism $\Phi : V \rightarrow \mathbb{R}_4$ defined by

$$\Phi(a_3x^3 + a_2x^2 + a_1x + a_0) = (a_3 \ a_2 \ a_1 \ a_0).$$

We thus have $\vec{v}_1 = (1 \ 2 \ -1 \ 0)$, $\vec{v}_2 = (4 \ 8 \ -4 \ -3)$, $\vec{v}_3 = (0 \ 1 \ 3 \ 4)$ and $\vec{v}_4 = (2 \ 5 \ 1 \ 4)$.

Applying Gaussian elimination to the matrix

$$\Omega(\vec{v}_1, \vec{v}_2, \vec{v}_3, \vec{v}_4) = \begin{pmatrix} 1 & 2 & -1 & 0 \\ 4 & 8 & -4 & -3 \\ 0 & 1 & 3 & 4 \\ 2 & 5 & 1 & 4 \end{pmatrix}$$

yields

$$\begin{pmatrix} 1 & 0 & -7 & 0 \\ 0 & 1 & 3 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

Here we applied Gauss-Jordan elimination, but Gaussian elimination is good enough. This gives the vectors $\vec{\omega}_1 = (1 \ 0 \ -7 \ 0)$, $\vec{\omega}_2 = (0 \ 1 \ 3 \ 0)$, $\vec{\omega}_3 = (0 \ 0 \ 0 \ 1)$.

Our basis of U is thus

$$\{\Phi^{-1}(\vec{\omega}_1), \Phi^{-1}(\vec{\omega}_2), \Phi^{-1}(\vec{\omega}_3)\} = \{x^3 - 7x, x^2 + 3x, 1\},$$

where we use that

$$\Phi^{-1}((a_3 \ a_2 \ a_1 \ a_0)) = a_3x^3 + a_2x^2 + a_1x + a_0.$$

4.2.3 Compute the image and rank of a linear map

Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. By computing the image of a linear map f , we mean computing a basis of $\text{Im}(f)$.

In order to compute a basis for $\text{Im}(f)$ we use the following lemma:

Lemma 4.13 *Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. If $\{v_1, \dots, v_n\}$ is a basis of V , then*

$$\text{Im}(f) = \text{span}\{f(v_1), \dots, f(v_n)\}.$$

Proof Let $w \in \text{Im}(f)$ so that $w = f(v)$ for some $v \in V$. We have scalars s_i for $1 \leq i \leq n$ so that $v = \sum_{i=1}^n s_i v_i$. We obtain

$$w = f(v) = f\left(\sum_{i=1}^n s_i v_i\right) = \sum_{i=1}^n s_i f(v_i)$$

so that w is a linear combination of the vectors $\{f(v_1), \dots, f(v_n)\}$. On the other hand, a linear combination of the vectors $f(v_i) \in \text{Im}(f)$ lies in the image of f as well, since $\text{Im}(f)$ is a vector subspace. Hence we have $\text{Im}(f) = \text{span}\{f(v_1), \dots, f(v_n)\}$, as claimed. $\square$

Knowing that $\text{Im}(f) = \text{span}\{f(v_1), \dots, f(v_n)\}$ we can apply the recipe from [Section 4.2.2](#) to $U = \text{span}\{f(v_1), \dots, f(v_n)\}$. By definition, the number of basis vectors for $\text{Im}(f)$ is the rank of f .

Example 4.14 Let

$$\mathbf{A} = \begin{pmatrix} 1 & -2 & 0 & 4 \\ 3 & 1 & 1 & 0 \\ -1 & -5 & -1 & 8 \\ 3 & 8 & 2 & -12 \end{pmatrix}$$

Compute a basis for the image of $f_{\mathbf{A}} : \mathbb{R}^4 \rightarrow \mathbb{R}^4$ and the rank of $f_{\mathbf{A}}$. By [Lemma 4.13](#) we have

$$U = \text{Im}(f_{\mathbf{A}}) = \text{span}\{\mathbf{A}\vec{e}_1, \mathbf{A}\vec{e}_2, \mathbf{A}\vec{e}_3, \mathbf{A}\vec{e}_4\} = \text{span}\{\vec{a}_1, \vec{a}_2, \vec{a}_3, \vec{a}_4\},$$

where $\{\vec{e}_i\}_{1 \leq i \leq 4}$ denotes the standard basis of $\mathbb{R}^4$ and $\{\vec{a}_i\}_{1 \leq i \leq 4}$ the column vectors of $\mathbf{A}$. Comparing with the general setup described above, we are in the case where $V = \mathbb{R}^4$ and $v_i = \mathbf{A}\vec{e}_i$ for $i = 1, 2, 3, 4$.

- (i) For the isomorphism $\Phi : V = \mathbb{R}^4 \rightarrow \mathbb{R}_4$ we usually choose the transpose (but any other isomorphism would work too). We thus have $\vec{v}_1 = (1 \ 3 \ -1 \ 3)$, $\vec{v}_2 = (-2 \ 1 \ -5 \ 8)$, $\vec{v}_3 = (0 \ 1 \ -1 \ 2)$ and $\vec{v}_4 = (4 \ 0 \ 8 \ -12)$.

- (ii) Applying Gaussian elimination to the matrix

$$\Omega(\vec{v}_1, \vec{v}_2, \vec{v}_3, \vec{v}_4) = \mathbf{A}^T = \begin{pmatrix} 1 & 3 & -1 & 3 \\ -2 & 1 & -5 & 8 \\ 0 & 1 & -1 & 2 \\ 4 & 0 & 8 & -12 \end{pmatrix}$$

yields

$$\begin{pmatrix} 1 & 0 & 2 & -3 \\ 0 & 1 & -1 & 2 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

Here again, we applied Gauss-Jordan elimination, but Gaussian elimination is good enough. This gives the vectors $\vec{\omega}_1 = (1 \ 0 \ 2 \ -3)$, $\vec{\omega}_2 = (0 \ 1 \ -1 \ 2)$.

- (iii) Our basis of $\text{Im}(f)$ is thus

$$\{\Phi^{-1}(\vec{\omega}_1), \Phi^{-1}(\vec{\omega}_2)\} = \left\{ \begin{pmatrix} 1 \\ 0 \\ 2 \\ -3 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ -1 \\ 2 \end{pmatrix} \right\},$$

where we use that the transpose is its own inverse. We also conclude that $\text{rank}(f_{\mathbf{A}}) = 2$.

Remark 4.15 In the special case where we want to compute a basis for the image of $f_{\mathbf{A}}$ for some matrix $\mathbf{A}$, the recipe thus reduces to the following steps. Take the transpose of $\mathbf{A}$, perform Gauss elimination, take the transpose again, write down the nonzero column vectors. This gives the desired basis.

4.2.4 Compute the kernel and nullity of a linear map

In order to find a recipe for computing the kernel and nullity of a linear map, we first start with a related problem. Let $\mathbf{A} \in M_{n,m}(\mathbb{K})$ be an $n \times m$ -matrix and

$$U = \left\{ \vec{\xi} \in \mathbb{K}_n \mid \vec{\xi}\mathbf{A} = \mathbf{0}_{\mathbb{K}_m} \right\},$$

where $\vec{\xi}\mathbf{A}$ is defined via matrix multiplication of the row vector $\vec{\xi} \in \mathbb{K}_n = M_{1,n}(\mathbb{K})$ and the matrix $\mathbf{A} \in M_{n,m}(\mathbb{K})$. Notice that $0_{\mathbb{K}_n} \in U$ and if $\vec{\xi}_1, \vec{\xi}_2 \in U$, then $s_1\vec{\xi}_1 + s_2\vec{\xi}_2 \in U$ for all $s_1, s_2 \in \mathbb{K}$. By Definition 3.21, it follows that U is a vector subspace of $\mathbb{K}_n$. We want to compute a basis for U . Applying Gauss elimination to the matrix $\mathbf{A}$, we obtain $r \in \mathbb{N}$ and elementary matrices $\mathbf{B}_1, \dots, \mathbf{B}_N$ so that

$$\mathbf{B}_N \cdots \mathbf{B}_1 \mathbf{A} = \Omega(\vec{\omega}_1, \dots, \vec{\omega}_r, 0_{\mathbb{K}_m}, \dots, 0_{\mathbb{K}_m})$$

for some linearly independent row vectors $(\vec{\omega}_1, \dots, \vec{\omega}_r) \in \mathbb{K}_m$. Since the matrix $\mathbf{B}_N \cdots \mathbf{B}_1$ is invertible, we also obtain a basis $\{\vec{\xi}_1, \dots, \vec{\xi}_n\}$ of $\mathbb{K}_n$ so that

$$\mathbf{B}_N \cdots \mathbf{B}_1 = \Omega(\vec{\xi}_1, \dots, \vec{\xi}_n).$$

We now claim that $S = \{\vec{\xi}_{r+1}, \dots, \vec{\xi}_n\}$ is a basis of U . The set S is linearly independent, hence we only need to show that $\text{span}(S) = U$. Since we have

$$\Omega(\vec{\xi}_1, \dots, \vec{\xi}_n) \mathbf{A} = \Omega(\vec{\omega}_1, \dots, \vec{\omega}_r, 0_{\mathbb{K}_m}, \dots, 0_{\mathbb{K}_m}),$$

the definition of matrix multiplication implies that $\vec{\xi}_i \mathbf{A} = \vec{\omega}_i$ for $1 \leq i \leq r$ and $\vec{\xi}_i \mathbf{A} = 0_{\mathbb{K}_m}$ for $r+1 \leq i \leq n$. Any vector in U can be written as $\vec{v} = \sum_{i=1}^n s_i \vec{\xi}_i$. The condition $\vec{v} \mathbf{A} = 0_{\mathbb{K}_m}$ then implies that $s_1 = \dots = s_r = 0$, hence S is generating.

We can use this observation to compute the kernel and nullity of a linear map $\mathbb{K}^n \rightarrow \mathbb{K}^m$ because of the following lemma whose proof is left as an exercise.

Lemma 4.16 Let $\mathbf{C} \in M_{m,n}(\mathbb{K})$ and $f_{\mathbf{C}} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ be the associated linear map. Then $\vec{x} \in \text{Ker}(f_{\mathbf{C}})$ if and only if $\vec{x}^T \mathbf{C}^T = 0_{\mathbb{K}_m}$.

We simply apply the above procedure to the matrix $\mathbf{A} = \mathbf{C}^T$ and compute the vectors $\{\vec{\xi}_{r+1}, \dots, \vec{\xi}_n\}$. The basis of $\text{Ker}(f_{\mathbf{C}})$ is then given by $\{\vec{\xi}_{r+1}^T, \dots, \vec{\xi}_n^T\}$.

The nullity of $f_{\mathbf{C}}$ is given by the number of basis vectors of $\text{Ker}(f_{\mathbf{C}})$.

Example 4.17 (Kernel of a linear map — Video) Let

$$\mathbf{C} = \begin{pmatrix} 1 & 0 & 1 & 7 \\ -2 & -3 & 1 & 2 \\ 7 & 9 & -2 & 1 \end{pmatrix}$$

In order to compute $\text{Ker}(f_{\mathbf{C}})$ we apply Gaussian elimination to $\mathbf{C}^T$ whilst keeping track of the relevant elementary matrices as in the algorithm for computing the inverse of a matrix. That is, we consider

$$\left(\begin{array}{ccc|cccc} 1 & -2 & 7 & 1 & 0 & 0 & 0 \\ 0 & -3 & 9 & 0 & 1 & 0 & 0 \\ 1 & 1 & -2 & 0 & 0 & 1 & 0 \\ 7 & 2 & 1 & 0 & 0 & 0 & 1 \end{array} \right).$$

Gauss–Jordan elimination (again, Gaussian elimination is enough) gives

$$\left(\begin{array}{ccc|cccc} 1 & 0 & 1 & 0 & 0 & -\frac{2}{5} & \frac{1}{5} \\ 0 & 1 & -3 & 0 & 0 & \frac{7}{5} & -\frac{1}{5} \\ 0 & 0 & 0 & 1 & 0 & \frac{16}{5} & -\frac{3}{5} \\ 0 & 0 & 0 & 0 & 1 & \frac{21}{5} & -\frac{3}{5} \end{array} \right).$$

The vectors $\vec{\xi}_3 = (1 \ 0 \ \frac{16}{5} \ -\frac{3}{5})$ and $\vec{\xi}_4 = (0 \ 1 \ \frac{21}{5} \ -\frac{3}{5})$ thus span the subspace of vectors ξ satisfying $\xi \mathbf{C}^T = 0_{\mathbb{K}_3}$. A basis S for the kernel of $f_{\mathbf{C}}$ is thus given

by

$$\mathcal{S} = \left\{ \begin{pmatrix} 1 \\ 0 \\ \frac{16}{5} \\ -\frac{3}{5} \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ \frac{21}{5} \\ -\frac{3}{5} \end{pmatrix} \right\}$$

and $f_{\mathbf{C}}$ satisfies $\text{nullity}(f_{\mathbf{C}}) = 2$.

Remark 4.18 Section 4.2.3 and Section 4.2.4 can be combined to compute $\text{Ker}(f_{\mathbf{A}})$ and $\text{Im}(f_{\mathbf{A}})$ for $\mathbf{A} \in M_{m,n}(\mathbb{K})$ by a single application of Gaussian elimination.

Remark 4.19 In order to compute the kernel of a linear map $g : V \rightarrow W$ between finite dimensional vector spaces, we can fix an ordered basis $\mathbf{b}$ of V and an ordered basis $\mathbf{c}$ of W , compute $\mathbf{C} = \mathbf{M}(g, \mathbf{b}, \mathbf{c})$, apply the above procedure to the matrix $\mathbf{C}$ in order to obtain a basis $\mathcal{S}$ of $\text{Ker}(f_{\mathbf{C}})$. The desired basis of $\text{Ker}(g)$ is then given by $\beta^{-1}(\mathcal{S})$. While this algorithm can always be carried out, it is computationally quite involved. In many cases it is therefore advisable to compute $\text{Ker}(g)$ by some other technique.

The determinant

WEEK 8

5.1 Axiomatic characterisation

Surprisingly, whether or not a square matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ admits an inverse is captured by a single scalar, called the *determinant of $\mathbf{A}$* and denoted by $\det \mathbf{A}$ or $\det(\mathbf{A})$. That is, the matrix $\mathbf{A}$ admits an inverse if and only if $\det \mathbf{A}$ is nonzero. In practice, however, it is often quicker to use Gauss–Jordan elimination to decide whether the matrix admits an inverse. The determinant is nevertheless a useful tool in linear algebra.

The determinant is an object of *multilinear algebra*, which – for $\ell \in \mathbb{N}$ – considers mappings from the ℓ -fold Cartesian product of a $\mathbb{K}$ -vector space into another $\mathbb{K}$ -vector space. Such a mapping f is required to be linear in each variable. This simply means that if we freeze all variables of f , except for the k -th variable, $1 \leq k \leq \ell$, then the resulting mapping g_k of one variable is required to be linear. More precisely:

Definition 5.1 (Multilinear map — Video) Let V, W be $\mathbb{K}$ -vector spaces and $\ell \in \mathbb{N}$. A mapping $f : V^\ell \rightarrow W$ is called ℓ -*multilinear* (or simply *multilinear*) if the mapping $g_k : V \rightarrow W, v \mapsto f(v_1, \dots, v_{k-1}, v, v_{k+1}, \dots, v_\ell)$ is linear for all $1 \leq k \leq \ell$ and for all ℓ -tuples $(v_1, \dots, v_\ell) \in V^\ell$.

We only need an $(\ell - 1)$ -tuple of vectors to define the map g_k , but the above definition is more convenient to write down.

Two types of multilinear maps are of particular interest:

Definition 5.2 (Symmetric and alternating multilinear maps) Let V, W be $\mathbb{K}$ -vector spaces and $f : V^\ell \rightarrow W$ an ℓ -multilinear map.

- The map f is called *symmetric* if exchanging two arguments does not change the value of f . That is, we have

$$f(v_1, \dots, v_\ell) = f(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_{j-1}, v_i, v_{j+1}, \dots, v_\ell)$$

for all $(v_1, \dots, v_\ell) \in V^\ell$.

- The map f is called *alternating* if $f(v_1, \dots, v_\ell) = 0_W$ whenever at least two arguments agree, that is, there exist $i \neq j$ with $v_i = v_j$. Alternating ℓ -multilinear maps are also called *W -valued ℓ -forms* or simply *ℓ -forms* when $W = \mathbb{K}$.

1-multilinear maps are simply linear maps. 2-multilinear maps are called *bilinear* and 3-multilinear maps are called *trilinear*. Most likely, you are already familiar with two examples of bilinear maps:

Example 5.3 (Bilinear maps)

- (i) The first one is the *scalar product* of two vectors in $\mathbb{R}^3$ (or more generally $\mathbb{R}^n$). So $V = \mathbb{R}^3$ and $W = \mathbb{R}$. Recall that the scalar product is the mapping

$$V^2 = \mathbb{R}^3 \times \mathbb{R}^3 \rightarrow \mathbb{R}, \quad (\vec{x}, \vec{y}) \mapsto \vec{x} \cdot \vec{y} = x_1y_1 + x_2y_2 + x_3y_3,$$

where we write $\vec{x} = (x_i)_{1 \leq i \leq 3}$ and $\vec{y} = (y_i)_{1 \leq i \leq 3}$. Notice that for all $s_1, s_2 \in \mathbb{R}$ and all $\vec{x}_1, \vec{x}_2, \vec{y} \in \mathbb{R}^3$ we have

$$(s_1\vec{x}_1 + s_2\vec{x}_2) \cdot \vec{y} = s_1(\vec{x}_1 \cdot \vec{y}) + s_2(\vec{x}_2 \cdot \vec{y}),$$

so that the scalar product is linear in the first variable. Furthermore, the scalar product is symmetric, $\vec{x} \cdot \vec{y} = \vec{y} \cdot \vec{x}$. It follows that the scalar product is also linear in the second variable, hence it is bilinear or 2-multilinear.

- (ii) The second one is the *cross product* of two vectors in $\mathbb{R}^3$. Here $V = \mathbb{R}^3$ and $W = \mathbb{R}^3$. Recall that the cross product is the mapping

$$V^2 = \mathbb{R}^3 \times \mathbb{R}^3 \rightarrow \mathbb{R}^3, \quad (\vec{x}, \vec{y}) \mapsto \vec{x} \times \vec{y} = \begin{pmatrix} x_2y_3 - x_3y_2 \\ x_3y_1 - x_1y_3 \\ x_1y_2 - x_2y_1 \end{pmatrix}.$$

Notice that for all $s_1, s_2 \in \mathbb{R}$ and all $\vec{x}_1, \vec{x}_2, \vec{y} \in \mathbb{R}^3$ we have

$$(s_1\vec{x}_1 + s_2\vec{x}_2) \times \vec{y} = s_1(\vec{x}_1 \times \vec{y}) + s_2(\vec{x}_2 \times \vec{y}),$$

so that the cross product is linear in the first variable. Likewise, we can check that the cross product is also linear in the second variable, hence it is bilinear or 2-multilinear. Observe that the cross product is alternating.

Example 5.4 (Multilinear map) Let $V = \mathbb{K}$ and consider $f : V^\ell \rightarrow \mathbb{K}, (x_1, \dots, x_\ell) \mapsto x_1x_2 \cdots x_\ell$. Then f is ℓ -multilinear and symmetric.

Example 5.5 Let $\mathbf{A} \in M_{n,n}(\mathbb{R})$ be a symmetric matrix, $\mathbf{A}^T = \mathbf{A}$. Notice that we obtain a symmetric bilinear map

$$f : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}, \quad (x, y) \mapsto \vec{x}^T \mathbf{A} \vec{y},$$

where on the right hand side all products are defined by matrix multiplication.

The [Example 5.5](#) gives us a wealth of symmetric bilinear maps on $\mathbb{R}^n$. As we will see shortly, the situation is quite different if we consider alternating n -multilinear maps on $\mathbb{K}_n$ (notice that we have the same number n of arguments as the dimension of $\mathbb{K}_n$).

Let $\{\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_n\}$ denote the standard basis of $\mathbb{K}_n$ so that $\Omega(\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_n) = \mathbf{1}_n$.

Theorem 5.6 Let $n \in \mathbb{N}$. Then there exists a unique alternating n -multilinear map $f_n : (\mathbb{K}_n)^n \rightarrow \mathbb{K}$ satisfying $f_n(\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_n) = 1$.

Recall that we have bijective mapping $\Omega : (\mathbb{K}_n)^n \rightarrow M_{n,n}(\mathbb{K})$ which forms an $n \times n$ -matrix from n row vectors of length n . For the choice $V = \mathbb{K}_n$, the notion of n -multilinearity thus also makes sense for a mapping $f : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ which takes an $n \times n$ matrix as an input. Here the multilinearity means the the mapping is linear in each row of the matrix. Since $\Omega(\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_n) = \mathbf{1}_n$, we may phrase the above theorem equivalently as:

Theorem 5.7 (Existence and uniqueness of the determinant) *Let $n \in \mathbb{N}$. Then there exists a unique alternating n -multilinear map $f_n : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ satisfying $f_n(\mathbf{1}_n) = 1$.*

Definition 5.8 (Determinant — Video) The mapping $f_n : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ provided by Theorem 5.7 is called the *determinant* and denoted by $\det$. For $\mathbf{A} \in M_{n,n}(\mathbb{K})$ we say $\det(\mathbf{A})$ is the determinant of the matrix $\mathbf{A}$.

Remark 5.9 (Abuse of notation) It would be more precise to write $\det_n$ since the determinant is a family of mappings, one mapping $\det_n : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ for each $n \in \mathbb{N}$. It is however common to simply write $\det$.

Example 5.10 For $n = 1$ the condition that a 1-multilinear (i.e. linear) map $f_1 : M_{1,1}(\mathbb{K}) \rightarrow \mathbb{K}$ is alternating is vacuous. So the Theorem 5.7 states that there is a unique linear map $f_1 : M_{1,1}(\mathbb{K}) \rightarrow \mathbb{K}$ that satisfies $f_1((1)) = 1$. Of course, this is just the map defined by the rule $f_1((a)) = a$, where $(a) \in M_{1,1}(\mathbb{K})$ is any 1-by-1 matrix.

Example 5.11 For $n = 2$ and $a, b, c, d \in \mathbb{K}$ we consider the mapping $f_2 : M_{2,2}(\mathbb{K}) \rightarrow \mathbb{K}$ defined by the rule

$$(5.1) \quad f_2 \left(\begin{pmatrix} a & b \\ c & d \end{pmatrix} \right) = ad - cb.$$

We claim that f_2 is bilinear in the rows and alternating. The condition that f_2 is alternating simplifies to $f(\mathbf{A}) = 0$ whenever the two rows of $\mathbf{A} \in M_{2,2}(\mathbb{K})$ agree. Clearly, f_2 is alternating, since

$$f_2 \left(\begin{pmatrix} a & b \\ a & b \end{pmatrix} \right) = ab - ab = 0.$$

Furthermore, f_2 needs to be linear in each row. The additivity condition applied to the first row gives that we must have

$$f_2 \left(\begin{pmatrix} a_1 + a_2 & b_1 + b_2 \\ c & d \end{pmatrix} \right) = f_2 \left(\begin{pmatrix} a_1 & b_1 \\ c & d \end{pmatrix} \right) + f_2 \left(\begin{pmatrix} a_2 & b_2 \\ c & d \end{pmatrix} \right)$$

for all $a_1, a_2, b_1, b_2, c, d \in \mathbb{K}$. Using the definition (5.1), we obtain

$$\begin{aligned} f_2 \left(\begin{pmatrix} a_1 + a_2 & b_1 + b_2 \\ c & d \end{pmatrix} \right) &= (a_1 + a_2)d - c(b_1 + b_2) \\ &= a_1d - cb_1 + a_2d - cb_2 \\ &= f_2 \left(\begin{pmatrix} a_1 & b_1 \\ c & d \end{pmatrix} \right) + f_2 \left(\begin{pmatrix} a_2 & b_2 \\ c & d \end{pmatrix} \right), \end{aligned}$$

so that f_2 is indeed additive in the first row. The 1-homogeneity condition applied to the first row gives that we must have

$$f_2 \left(\begin{pmatrix} sa & sb \\ c & d \end{pmatrix} \right) = sf_2 \left(\begin{pmatrix} a & b \\ c & d \end{pmatrix} \right)$$

for all $a, b, c, d \in \mathbb{K}$ and $s \in \mathbb{K}$. Indeed, using the definition (5.1), we obtain

$$f_2 \left(\begin{pmatrix} sa & sb \\ c & d \end{pmatrix} \right) = sad - csb = s(ad - cb) = sf_2 \left(\begin{pmatrix} a & b \\ c & d \end{pmatrix} \right),$$

so that f_2 is also 1-homogeneous in the first row. We conclude that f_2 is linear in the first row. Likewise, the reader is invited to check that f_2 is also linear in the second row. Furthermore, we can easily compute that $f_2(\mathbf{1}_2) = 1$. The mapping f_2 thus satisfies all the properties of Theorem 5.7, hence by the uniqueness statement we must have $f_2 = \det$ and we obtain the formula

$$(5.2) \quad \det \left(\begin{pmatrix} a & b \\ c & d \end{pmatrix} \right) = ad - cb$$

for all $a, b, c, d \in \mathbb{K}$.

5.2 Uniqueness of the determinant

So far we have only shown that the determinant exists for $n = 1$ and $n = 2$. However, we need to show the existence and uniqueness part of Theorem 5.7 in general. We first show the uniqueness part. We start by deducing some consequences from the alternating property:

Lemma 5.12 *Let V, W be $\mathbb{K}$ -vector spaces and $\ell \in \mathbb{N}$. An alternating ℓ -multilinear map $f : V^\ell \rightarrow W$ satisfies:*

- (i) *interchanging two arguments of f leads to a minus sign. That is, for $1 \leq i, j \leq \ell$ and $i \neq j$ we obtain*

$$f(v_1, \dots, v_\ell) = -f(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_{j-1}, v_i, v_{j+1}, \dots, v_\ell)$$

for all $(v_1, \dots, v_\ell) \in V^\ell$;

- (ii) *if the vectors $(v_1, \dots, v_\ell) \in V^\ell$ are linearly dependent, then $f(v_1, \dots, v_\ell) = 0_W$;*

- (iii) *for all $1 \leq i \leq \ell$, for all ℓ -tuples of vectors $(v_1, \dots, v_\ell) \in V^\ell$ and scalars $s_1, \dots, s_\ell \in \mathbb{K}$, we have*

$$f(v_1, \dots, v_{i-1}, v_i + w, v_{i+1}, \dots, v_\ell) = f(v_1, \dots, v_\ell)$$

where $w = \sum_{j=1, j \neq i}^\ell s_j v_j$. That is, adding a linear combination of vectors to some argument of f does not change the output, provided the linear combination consists of the remaining arguments.

Proof (i) Since f is alternating, we have for all $(v_1, \dots, v_\ell) \in V^\ell$

$$f(v_1, \dots, v_{i-1}, v_i + v_j, v_{i+1}, \dots, v_{j-1}, v_i + v_j, v_{j+1}, \dots, v_\ell) = 0_W.$$

Using the linearity in the i -th argument, this gives

$$\begin{aligned} 0_W &= f(v_1, \dots, v_{i-1}, v_i, v_{i+1}, \dots, v_{j-1}, v_i + v_j, v_{j+1}, \dots, v_\ell) \\ &\quad + f(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_{j-1}, v_i + v_j, v_{j+1}, \dots, v_\ell). \end{aligned}$$

Using the linearity in the j -th argument, we obtain

$$\begin{aligned} 0_W &= f(v_1, \dots, v_{i-1}, v_i, v_{i+1}, \dots, v_{j-1}, v_i, v_{j+1}, \dots, v_\ell) \\ &\quad + f(v_1, \dots, v_{i-1}, v_i, v_{i+1}, \dots, v_{j-1}, v_j, v_{j+1}, \dots, v_\ell) \\ &\quad + f(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_{j-1}, v_i, v_{j+1}, \dots, v_\ell) \\ &\quad + f(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_{j-1}, v_j, v_{j+1}, \dots, v_\ell). \end{aligned}$$

The first summand has a double occurrence of v_i and hence vanishes by the alternating property. Likewise, the fourth summand has a double occurrence of v_j and hence vanishes as well. Since the second summand equals $f(v_1, \dots, v_\ell)$, we thus obtain

$$f(v_1, \dots, v_\ell) = -f(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_{j-1}, v_i, v_{j+1}, \dots, v_\ell)$$

as claimed.

(ii) Suppose $\{v_1, \dots, v_\ell\}$ are linearly dependent so that we have scalars $s_j \in \mathbb{K}$ not all zero, $1 \leq j \leq \ell$, so that $s_1 v_1 + \dots + s_\ell v_\ell = 0_V$. Suppose $s_i \neq 0$ for some index $1 \leq i \leq \ell$. Then

$$v_i = - \sum_{j=1, j \neq i}^{\ell} \left(\frac{s_j}{s_i} \right) v_j$$

and hence by the linearity in the i -th argument, we obtain

$$\begin{aligned} f \left(v_1, \dots, v_{i-1}, - \sum_{j=1, j \neq i}^{\ell} \left(\frac{s_j}{s_i} \right) v_j, v_{i+1}, \dots, v_\ell \right) \\ = - \sum_{j=1, j \neq i}^{\ell} \left(\frac{s_j}{s_i} \right) f(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_\ell) = 0_W, \end{aligned}$$

where we use that for each $1 \leq j \leq \ell$ with $j \neq i$, the expression

$$f(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_\ell)$$

has a double occurrence of v_j and thus vanishes by the alternating property.

(iii) Let $(v_1, \dots, v_\ell) \in V^\ell$ and $(s_1, \dots, s_\ell) \in \mathbb{K}^\ell$. Then, using the linearity in the i -th argument, we compute

$$\begin{aligned} f(v_1, \dots, v_{i-1}, v_i + \sum_{j=1, j \neq i}^{\ell} s_j v_j, v_{i+1}, \dots, v_\ell) \\ = f(v_1, \dots, v_\ell) + \sum_{j=1, j \neq i}^{\ell} s_j f(v_1, \dots, v_{i-1} v_j, v_{i+1}, \dots, v_\ell) = f(v_1, \dots, v_\ell), \end{aligned}$$

where the last equality follows exactly as in the proof of (ii). $\square$

The alternating property of an n -multilinear map $f_n : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ together with the condition $f_n(\mathbf{1}_n) = 1$ uniquely determines the value of f_n on the elementary matrices:

Lemma 5.13 *Let $n \in \mathbb{N}$ and $f_n : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ an alternating n -multilinear map satisfying $f_n(\mathbf{1}_n) = 1$. Then for all $1 \leq k, l \leq n$ with $k \neq l$ and all $s \in \mathbb{K}$, we have*

$$(5.3) \quad f_n(\mathbf{D}_k(s)) = s, \quad f_n(\mathbf{L}_{k,l}(s)) = 1, \quad f_n(\mathbf{P}_{k,l}) = -1.$$

Moreover, for $\mathbf{A} \in M_{n,n}(\mathbb{K})$ and an elementary matrix $\mathbf{B}$ of size n , we have

$$(5.4) \quad f_n(\mathbf{BA}) = f_n(\mathbf{B})f_n(\mathbf{A}).$$

Proof Recall that $\mathbf{D}_k(s)$ applied to a square matrix $\mathbf{A}$ multiplies the k -th row of $\mathbf{A}$ with s and leaves $\mathbf{A}$ unchanged otherwise. We write $\mathbf{A} \in M_{n,n}(\mathbb{K})$ as $\mathbf{A} = \Omega(\vec{\alpha}_1, \dots, \vec{\alpha}_n)$ for $\vec{\alpha}_i \in \mathbb{K}_n$, $1 \leq i \leq n$. Hence we obtain

$$\mathbf{D}_k(s)\mathbf{A} = \Omega(\vec{\alpha}_1, \dots, \vec{\alpha}_{k-1}, s\vec{\alpha}_k, \vec{\alpha}_{k+1}, \dots, \vec{\alpha}_n).$$

The linearity of f in the k -th row thus gives $f_n(\mathbf{D}_k(s)\mathbf{A}) = sf_n(\mathbf{A})$. In particular, the choice $\mathbf{A} = \mathbf{1}_n$ together with $f_n(\mathbf{1}_n) = 1$ implies that $f_n(\mathbf{D}_k(s)) = f_n(\mathbf{D}_k(s)\mathbf{1}_n) = sf_n(\mathbf{1}_n) = s$.

Therefore, we have

$$f_n(\mathbf{D}_k(s)\mathbf{A}) = f_n(\mathbf{D}_k(s))f_n(\mathbf{A}).$$

Likewise we obtain

$$\mathbf{L}_{k,l}(s)\mathbf{A} = \Omega(\vec{\alpha}_1, \dots, \vec{\alpha}_{k-1}, \vec{\alpha}_k + s\vec{\alpha}_l, \vec{\alpha}_{k+1}, \dots, \vec{\alpha}_n)$$

and we can apply property (iii) of [Lemma 5.12](#) for the choice $w = s\vec{\alpha}_l$ to conclude that $f_n(\mathbf{L}_{k,l}(s)\mathbf{A}) = f_n(\mathbf{A})$. In particular, the choice $\mathbf{A} = \mathbf{1}_n$ together with $f_n(\mathbf{1}_n) = 1$ implies $f_n(\mathbf{L}_{k,l}(s)) = f_n(\mathbf{L}_{k,l}(s)\mathbf{1}_n) = f_n(\mathbf{1}_n) = 1$.

Therefore, we have

$$f_n(\mathbf{L}_{k,l}(s)\mathbf{A}) = f_n(\mathbf{L}_{k,l}(s))f_n(\mathbf{A}).$$

Finally, we have

$$\mathbf{P}_{k,l}\mathbf{A} = \Omega(\vec{\alpha}_1, \dots, \vec{\alpha}_{k-1}, \vec{\alpha}_l, \vec{\alpha}_{k+1}, \dots, \vec{\alpha}_{l-1}, \vec{\alpha}_k, \vec{\alpha}_{l+1}, \dots, \vec{\alpha}_n)$$

so that property (ii) of [Lemma 5.12](#) immediately gives that

$$f_n(\mathbf{P}_{k,l}\mathbf{A}) = -f_n(\mathbf{A}).$$

In particular, the choice $\mathbf{A} = \mathbf{1}_n$ together with $f_n(\mathbf{1}_n) = 1$ implies $f_n(\mathbf{P}_{k,l}) = f_n(\mathbf{P}_{k,l}\mathbf{1}_n) = -f_n(\mathbf{1}_n) = -1$.

Therefore, we have $f_n(\mathbf{P}_{k,l}\mathbf{A}) = f_n(\mathbf{P}_{k,l})f_n(\mathbf{A})$, as claimed. $\square$

We now obtain the uniqueness part of [Theorem 5.7](#).

Proposition 5.14 *Let $n \in \mathbb{N}$ and $f_n, \hat{f}_n : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ be alternating n -multilinear maps satisfying $f_n(\mathbf{1}_n) = \hat{f}_n(\mathbf{1}_n) = 1$. Then $f_n = \hat{f}_n$.*

Proof We need to show that for all $\mathbf{A} \in M_{n,n}(\mathbb{K})$, we have $f_n(\mathbf{A}) = \hat{f}_n(\mathbf{A})$. Suppose first that $\mathbf{A}$ is not invertible. Then, by [Proposition 4.7](#), the row vectors of $\mathbf{A}$ are linearly dependent and hence property (ii) of [Lemma 5.12](#) implies that $f_n(\mathbf{A}) = \hat{f}_n(\mathbf{A}) = 0$.

Now suppose that $\mathbf{A}$ is invertible. Using Gauss–Jordan elimination, we obtain $N \in \mathbb{N}$ and a sequence of elementary matrices $\mathbf{B}_1, \dots, \mathbf{B}_N$ so that $\mathbf{B}_N \cdots \mathbf{B}_1 = \mathbf{A}$. We obtain

$$f_n(\mathbf{A}) = f_n(\mathbf{B}_N \cdots \mathbf{B}_1) = f_n(\mathbf{B}_N)f_n(\mathbf{B}_{N-1} \cdots \mathbf{B}_1) = \hat{f}_n(\mathbf{B}_N)f_n(\mathbf{B}_{N-1} \cdots \mathbf{B}_1),$$

where the second equality uses (5.4) and the third equality uses that (5.3) implies that $\hat{f}_n(\mathbf{B}) = f_n(\mathbf{B})$ for all elementary matrices $\mathbf{B}$. Proceeding in this fashion we get

$$\begin{aligned} f_n(\mathbf{A}) &= \hat{f}_n(\mathbf{B}_N)\hat{f}_n(\mathbf{B}_{N-1}) \cdots \hat{f}_n(\mathbf{B}_1) = \hat{f}_n(\mathbf{B}_N)\hat{f}_n(\mathbf{B}_{N-1}) \cdots \hat{f}_n(\mathbf{B}_2\mathbf{B}_1) = \cdots \\ &= \hat{f}_n(\mathbf{B}_N\mathbf{B}_{N-1} \cdots \mathbf{B}_1) = \hat{f}_n(\mathbf{A}). \end{aligned}$$

$\square$

5.3 Existence of the determinant

It turns out that we can define the determinant recursively in terms of the determinants of certain submatrices. Determinants of submatrices are called *minors*. To this end we first define:

Definition 5.15 Let $n \in \mathbb{N}$. For a square matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ and $1 \leq k, l \leq n$ we denote by $\mathbf{A}^{(k,l)}$ the $(n-1) \times (n-1)$ submatrix obtained by removing the k -th row and l -th column from $\mathbf{A}$.

Example 5.16

$$\mathbf{A} = \begin{pmatrix} a & b \\ c & d \end{pmatrix}, \quad \mathbf{A}^{(1,1)} = (d), \quad \mathbf{A}^{(2,1)} = (b).$$

$$\mathbf{A} = \begin{pmatrix} 1 & -2 & 0 & 4 \\ 3 & 1 & 1 & 0 \\ -1 & -5 & -1 & 8 \\ 3 & 8 & 2 & -12 \end{pmatrix}, \quad \mathbf{A}^{(3,2)} = \begin{pmatrix} 1 & 0 & 4 \\ 3 & 1 & 0 \\ 3 & 2 & -12 \end{pmatrix}.$$

We use induction to prove the existence of the determinant:

Lemma 5.17 Let $n \in \mathbb{N}$ with $n \geq 2$ and $f_{n-1} : M_{n-1,n-1}(\mathbb{K}) \rightarrow \mathbb{K}$ an alternating $(n-1)$ -multilinear mapping satisfying $f_{n-1}(\mathbf{1}_{n-1}) = 1$. Then, for any fixed integer l with $1 \leq l \leq n$, the mapping

$$f_n : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}, \quad \mathbf{A} \mapsto \sum_{k=1}^n (-1)^{l+k} [\mathbf{A}]_{kl} f_{n-1}(\mathbf{A}^{(k,l)})$$

is alternating, n -multilinear and satisfies $f_n(\mathbf{1}_n) = 1$.

Proof of Theorem 5.6 For $n = 1$ we have seen that $f_1 : M_{1,1}(\mathbb{K}) \rightarrow \mathbb{K}$, $(a) \mapsto a$ is 1-multilinear, alternating and satisfies $f_1(\mathbf{1}_1) = 1$. Hence Lemma 5.17 implies that for all $n \in \mathbb{N}$ there exists an n -multilinear and alternating map $f_n : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ satisfying $f_n(\mathbf{1}_n) = 1$. By Proposition 5.14 there is only one such mapping for each $n \in \mathbb{N}$. $\square$

Proof of Lemma 5.17 We take some arbitrary, but then fixed integer l with $1 \leq l \leq n$.

Step 1. We first show that $f_n(\mathbf{1}_n) = 1$. Since $[\mathbf{1}_n]_{kl} = \delta_{kl}$, we obtain

$$f_n(\mathbf{1}_n) = \sum_{k=1}^n (-1)^{l+k} [\mathbf{1}_n]_{kl} f_{n-1}(\mathbf{1}_n^{(k,l)}) = (-1)^{2l} f_{n-1}(\mathbf{1}_n^{(l,l)}) = f_{n-1}(\mathbf{1}_{n-1}) = 1,$$

where we use that $\mathbf{1}_n^{(l,l)} = \mathbf{1}_{n-1}$ and $f_{n-1}(\mathbf{1}_{n-1}) = 1$.

Step 2. We show that f_n is multilinear. Let $\mathbf{A} \in M_{n,n}(\mathbb{K})$ and write $\mathbf{A} = (A_{kj})_{1 \leq k, j \leq n}$. We first show that f_n is 1-homogeneous in each row. Say we multiply the i -th row of $\mathbf{A}$ with s so that we obtain a new matrix $\hat{\mathbf{A}} = (\hat{A}_{kj})_{1 \leq k, j \leq n}$ with

$$\hat{A}_{kj} = \begin{cases} A_{kj}, & k \neq i, \\ sA_{kj}, & k = i. \end{cases}$$

We need to show that $f_n(\hat{\mathbf{A}}) = sf_n(\mathbf{A})$. We compute

$$\begin{aligned} f_n(\hat{\mathbf{A}}) &= \sum_{k=1}^n (-1)^{l+k} \hat{A}_{kl} f_{n-1}(\hat{\mathbf{A}}^{(k,l)}) \\ &= (-1)^{l+i} sA_{il} f_{n-1}(\hat{\mathbf{A}}^{(i,l)}) + \sum_{k=1, k \neq i}^n (-1)^{l+k} A_{kl} f_{n-1}(\hat{\mathbf{A}}^{(k,l)}). \end{aligned}$$

Now notice that $\hat{\mathbf{A}}^{(i,l)} = \mathbf{A}^{(i,l)}$, since $\mathbf{A}$ and $\hat{\mathbf{A}}$ only differ in the i -th row, but this is the row that is removed. Since f_{n-1} is 1-homogeneous in each row, we obtain that $f_{n-1}(\hat{\mathbf{A}}^{(k,l)}) = sf_{n-1}(\mathbf{A}^{(k,l)})$ whenever $k \neq i$. Thus we have

$$\begin{aligned} f_n(\hat{\mathbf{A}}) &= s(-1)^{l+i} A_{il} f_{n-1}(\mathbf{A}^{(i,l)}) + s \sum_{k=1, k \neq i}^n (-1)^{l+k} A_{kl} f_{n-1}(\mathbf{A}^{(k,l)}) \\ &= s \sum_{k=1}^n (-1)^{l+k} A_{kl} f_{n-1}(\mathbf{A}^{(k,l)}) = sf_n(\mathbf{A}). \end{aligned}$$

We now show that f_n is additive in each row. Say the matrix $\mathbf{B} = (B_{kj})_{1 \leq k, j \leq n}$ is identical to the matrix $\mathbf{A}$, except for the i -th row, so that

$$B_{kj} = \begin{cases} A_{kj} & k \neq i \\ B_j & k = i \end{cases}$$

for some scalars B_j with $1 \leq j \leq n$. We need to show that $f_n(\mathbf{C}) = f_n(\mathbf{A}) + f_n(\mathbf{B})$, where $\mathbf{C} = (C_{kj})_{1 \leq k, j \leq n}$ with

$$C_{kj} = \begin{cases} A_{kj} & k \neq i \\ A_{ij} + B_j & k = i \end{cases}$$

We compute

$$f_n(\mathbf{C}) = (-1)^{l+i} (A_{il} + B_l) f_{n-1}(\mathbf{C}^{(i,l)}) + \sum_{k=1, k \neq i}^n (-1)^{l+k} A_{kl} f_{n-1}(\mathbf{C}^{(k,l)}).$$

As before, since $\mathbf{A}$, $\mathbf{B}$, $\mathbf{C}$ only differ in the i -th row, we have $\mathbf{A}^{(i,l)} = \mathbf{B}^{(i,l)} = \mathbf{C}^{(i,l)}$. Using that f_{n-1} is linear in each row, we thus obtain

$$\begin{aligned} f_n(\mathbf{C}) &= (-1)^{l+i} B_l f_{n-1}(\mathbf{B}^{(i,l)}) + \sum_{k=1, k \neq i}^n (-1)^{l+k} A_{kl} f_{n-1}(\mathbf{B}^{(k,l)}) \\ &\quad + (-1)^{l+i} A_{il} f_{n-1}(\mathbf{A}^{(i,l)}) + \sum_{k=1, k \neq i}^n (-1)^{l+k} A_{kl} f_{n-1}(\mathbf{A}^{(k,l)}) = f_n(\mathbf{A}) + f_n(\mathbf{B}). \end{aligned}$$

Step 3. We show that f_n is alternating. Suppose we have $1 \leq i, j \leq n$ with $j > i$ and so that the i -th and j -th row of the matrix $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n}$ are the same. Therefore, unless $k = i$ or $k = j$, the submatrix $\mathbf{A}^{(k,l)}$ also contains two identical rows and since f_{n-1} is alternating, all summands vanish except the one for $k = i$ and $k = j$, this gives

$$\begin{aligned} f_n(\mathbf{A}) &= (-1)^{i+l} A_{il} f_{n-1}(\mathbf{A}^{(i,l)}) + (-1)^{j+l} A_{jl} f_{n-1}(\mathbf{A}^{(j,l)}) \\ &= A_{il} (-1)^l \left((-1)^i f_{n-1}(\mathbf{A}^{(i,l)}) + (-1)^j f_{n-1}(\mathbf{A}^{(j,l)}) \right) \end{aligned}$$

where the second equality sign follows because we have $A_{il} = A_{jl}$ for all $1 \leq l \leq n$ (the i -th and j -th row agree). The mapping f_{n-1} is alternating, hence by the first property of the [Lemma 5.12](#), swapping rows in the matrix $\mathbf{A}^{(j,l)}$ leads to a minus sign in $f_{n-1}(\mathbf{A}^{(j,l)})$. Moving the i -th row of $\mathbf{A}^{(j,l)}$ down by $j - i - 1$ rows (which corresponds to swapping $j - i - 1$ times), we obtain $\mathbf{A}^{(i,l)}$, hence

$$f_{n-1}(\mathbf{A}^{(j,l)}) = (-1)^{j-i-1} f_{n-1}(\mathbf{A}^{(i,l)}).$$

This gives

$$f_n(\mathbf{A}) = A_{il} (-1)^l \left((-1)^i f_{n-1}(\mathbf{A}^{(i,l)}) + (-1)^{2j-i-1} f_{n-1}(\mathbf{A}^{(i,l)}) \right) = 0.$$

□

Remark 5.18 (Laplace expansion — [Video](#)) As a by-product of the proof of [Lemma 5.17](#) we obtain the formula

$$(5.5) \quad \det(\mathbf{A}) = \sum_{k=1}^n (-1)^{l+k} [\mathbf{A}]_{kl} \det(\mathbf{A}^{(k,l)}),$$

known as the *Laplace expansion* of the determinant. The uniqueness statement of [Theorem 5.7](#) thus guarantees that for every $n \times n$ matrix $\mathbf{A}$, the scalar $\sum_{k=1}^n (-1)^{l+k} [\mathbf{A}]_{kl} \det(\mathbf{A}^{(k,l)})$ is independent of the choice of $l \in \mathbb{N}$, $1 \leq l \leq n$. In practice, when computing the determinant, it is thus advisable to choose l such that the corresponding column contains the maximal amount of zeros.

Example 5.19 For $n = 2$ and choosing $l = 1$, we obtain

$$\det \left(\begin{pmatrix} a & b \\ c & d \end{pmatrix} \right) = a \det(\mathbf{A}^{(1,1)}) - c \det(\mathbf{A}^{(2,1)}) = ad - cb,$$

in agreement with [\(5.1\)](#). For $\mathbf{A} = (A_{ij})_{1 \leq i,j \leq 3} \in M_{3,3}(\mathbb{K})$ and choosing $l = 3$ we obtain

$$\begin{aligned} \det \left(\begin{pmatrix} A_{11} & A_{12} & A_{13} \\ A_{21} & A_{22} & A_{23} \\ A_{31} & A_{32} & A_{33} \end{pmatrix} \right) &= A_{13} \det \left(\begin{pmatrix} A_{21} & A_{22} \\ A_{31} & A_{32} \end{pmatrix} \right) \\ &\quad - A_{23} \det \left(\begin{pmatrix} A_{11} & A_{12} \\ A_{31} & A_{32} \end{pmatrix} \right) + A_{33} \det \left(\begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix} \right) \end{aligned}$$

so that

$$\begin{aligned} \det \mathbf{A} &= A_{13}(A_{21}A_{32} - A_{31}A_{22}) - A_{23}(A_{11}A_{32} - A_{31}A_{12}) \\ &\quad + A_{33}(A_{11}A_{22} - A_{21}A_{12}) \\ &= A_{11}A_{22}A_{33} - A_{11}A_{23}A_{32} - A_{12}A_{21}A_{33} \\ &\quad + A_{12}A_{23}A_{31} + A_{13}A_{21}A_{32} - A_{13}A_{22}A_{31}. \end{aligned}$$

Exercises

Exercise 5.20 (Trilinear map) Let $V = \mathbb{R}^3$ and $W = \mathbb{R}$. Show that the map

$$f : V^3 \rightarrow W, \quad (\vec{x}, \vec{y}, \vec{z}) \mapsto (\vec{x} \times \vec{y}) \cdot \vec{z}$$

is alternating and trilinear.

5.4 Properties of the determinant

Proposition 5.21 (Product rule) For matrices $\mathbf{A}, \mathbf{B} \in M_{n,n}(\mathbb{K})$ we have

$$\det(\mathbf{AB}) = \det(\mathbf{A}) \det(\mathbf{B}).$$

Proof We first consider the case where $\mathbf{A}$ is not invertible, then $\det(\mathbf{A}) = 0$ (see the proof of [Proposition 5.14](#)). If $\mathbf{A}$ is not invertible, then neither is $\mathbf{AB}$. Indeed, if $\mathbf{AB}$ were invertible, then there exists a matrix $\mathbf{C}$ such that $(\mathbf{AB})\mathbf{C} = \mathbf{1}_n$. But since, by [Corollary 2.22](#), the matrix product is associative, this also gives $\mathbf{A}(\mathbf{BC}) = \mathbf{1}_n$, so that $\mathbf{BC}$ is the inverse of $\mathbf{A}$, a contradiction. Hence if $\mathbf{A}$ is not invertible, we must also have $\det(\mathbf{AB}) = 0$, which verifies that $\det(\mathbf{AB}) = 0 = \det(\mathbf{A}) \det(\mathbf{B})$ for $\mathbf{A}$ not invertible.

If $\mathbf{A}$ is invertible, we can write it as a product of elementary matrices and applying the second part of [Lemma 5.13](#), we conclude that $\det(\mathbf{AB}) = \det(\mathbf{A}) \det(\mathbf{B})$. $\square$

Corollary 5.22 A matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ is invertible if and only if $\det(\mathbf{A}) \neq 0$. Moreover, in the case where $\mathbf{A}$ is invertible, we have

$$\det(\mathbf{A}^{-1}) = \frac{1}{\det \mathbf{A}}.$$

Proof We have already seen that if $\mathbf{A}$ is not invertible, then $\det(\mathbf{A}) = 0$. This is equivalent to saying that if $\det(\mathbf{A}) \neq 0$, then $\mathbf{A}$ is invertible. It thus remains to show that if $\mathbf{A}$ is invertible, then $\det(\mathbf{A}) \neq 0$. Suppose $\mathbf{A}$ is invertible, then applying [Proposition 5.21](#) gives

$$\det(\mathbf{1}_n) = \det(\mathbf{AA}^{-1}) = \det(\mathbf{A}) \det(\mathbf{A}^{-1}) = 1$$

so that $\det(\mathbf{A}) \neq 0$ and $\det(\mathbf{A}^{-1}) = 1/\det(\mathbf{A})$. $\square$

Remark 5.23 (Product symbol) Recall that for scalars $x_1, \dots, x_n \in \mathbb{K}$, we write

$$\prod_{i=1}^n x_i = x_1 x_2 \cdots x_n.$$

Proposition 5.24 Let $n \in \mathbb{N}$ and $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n} \in M_{n,n}(\mathbb{K})$ be an upper triangular matrix so that $A_{ij} = 0$ for $i > j$. Then

$$(5.6) \quad \det(\mathbf{A}) = \prod_{i=1}^n A_{ii} = A_{11} A_{22} \cdots A_{nn}.$$

Proof We use induction. For $n = 1$ the condition $A_{ij} = 0$ for $i > j$ is vacuous and (5.6) is trivially satisfied, thus the statement is anchored.

Inductive step: Assume $n \in \mathbb{N}$ and $n \geq 2$. We want to show that if (5.6) holds for upper triangular $(n-1) \times (n-1)$ -matrices, then also for upper triangular $n \times n$ -matrices. Let $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n} \in M_{n,n}(\mathbb{K})$ be an upper triangular matrix. Choosing $l = 1$ in the formula

for $\det(\mathbf{A})$, we obtain

$$\begin{aligned}\det(\mathbf{A}) &= \sum_{k=1}^n (-1)^{k+1} A_{k1} \det(\mathbf{A}^{(k,1)}) = A_{11} \det(\mathbf{A}^{(1,1)}) + \sum_{k=2}^n A_{k1} \det(\mathbf{A}^{(k,1)}) \\ &= A_{11} \det(\mathbf{A}^{(1,1)}),\end{aligned}$$

where the last equality uses that $A_{k1} = 0$ for $k > 1$. We have $\mathbf{A}^{(1,1)} = (A_{ij})_{2 \leq i, j \leq n}$ and since $\mathbf{A}$ is an upper triangular matrix, it follows that $\mathbf{A}^{(1,1)}$ is an $(n-1) \times (n-1)$ upper triangular matrix as well. Hence by the induction hypothesis, we obtain

$$\det(\mathbf{A}^{(1,1)}) = \prod_{i=2}^n A_{ii}.$$

We conclude that $\det(\mathbf{A}) = \prod_{i=1}^n A_{ii}$, as claimed. $\square$

5.5 Permutations

A rearrangement of the natural numbers from 1 up to n is called a permutation:

Definition 5.25 (Permutation — Video) Let $n \in \mathbb{N}$ and $\mathcal{X}_n = \{1, 2, 3, \dots, n\}$. A *permutation* is a bijective mapping $\sigma : \mathcal{X}_n \rightarrow \mathcal{X}_n$. The set of all permutations of $\mathcal{X}_n$ is denoted by S_n .

Remark 5.26 If $\tau, \sigma : \mathcal{X}_n \rightarrow \mathcal{X}_n$ are permutations, it is customary to write $\tau\sigma$ or $\tau \cdot \sigma$ instead of $\tau \circ \sigma$. Furthermore, the identity mapping $\text{Id}_{\mathcal{X}_n}$ is often simply denoted by 1. A convenient way to describe a permutation $\sigma \in S_n$ is in terms of a $2 \times n$ matrix

$$\begin{pmatrix} i \\ \sigma(i) \end{pmatrix}_{1 \leq i \leq n}.$$

which we denote by σ . For instance, for $n = 4$, the matrix

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 2 & 3 & 1 & 4 \end{pmatrix}$$

corresponds to the permutation σ satisfying $\sigma(1) = 2, \sigma(2) = 3, \sigma(3) = 1, \sigma(4) = 4$.

Permutations which only swap two natural numbers and leave all remaining numbers unchanged are known as *transpositions*:

Definition 5.27 (Transposition) Let $n \in \mathbb{N}$ and $1 \leq k, l \leq n$ with $k \neq l$. The *transposition* $\tau_{k,l} \in S_n$ is the permutation satisfying

$$\tau_{k,l}(k) = l, \quad \tau_{k,l}(l) = k, \quad \tau_{k,l}(i) = i \text{ if } i \notin \{k, l\}.$$

Every permutation $\sigma \in S_n$ defines a linear map $g : \mathbb{K}^n \rightarrow \mathbb{K}^n$ satisfying $g(\vec{e}_i) = \vec{e}_{\sigma(i)}$, where $\{\vec{e}_1, \dots, \vec{e}_n\}$ denotes the standard basis of $\mathbb{K}^n$. Since g is linear, there exists a unique matrix $\mathbf{P}_\sigma \in M_{n,n}(\mathbb{K})$ so that $g = \mathbf{f}_{\mathbf{P}_\sigma}$. Observe that the column vectors of the matrix $\mathbf{P}_\sigma$ are given by $\vec{e}_{\sigma(1)}, \vec{e}_{\sigma(2)}, \dots, \vec{e}_{\sigma(n)}$.

Definition 5.28 (Permutation matrix) We call $\mathbf{P}_\sigma \in M_{n,n}(\mathbb{K})$ the *permutation matrix* associated to $\sigma \in S_n$.

Example 5.29 Let $n = 4$. For instance, we have

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 2 & 3 & 1 & 4 \end{pmatrix} \quad \mathbf{P}_\sigma = \begin{pmatrix} 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

and

$$\tau_{2,4} = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 1 & 4 & 3 & 2 \end{pmatrix} \quad \mathbf{P}_{\tau_{2,4}} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \end{pmatrix}.$$

Remark 5.30 Notice that $\mathbf{P}_{\tau_{k,l}} = \mathbf{P}_{k,l}$, where $\mathbf{P}_{k,l}$ belongs to the elementary matrices of size n , c.f. [Definition 4.1](#).

Assigning to a permutation its permutation matrix turns composition of permutations into matrix multiplication:

Proposition 5.31 Let $n \in \mathbb{N}$. Then $\mathbf{P}_1 = \mathbf{1}_n$ and for all $\sigma, \pi \in S_n$ we have

$$\mathbf{P}_{\pi \cdot \sigma} = \mathbf{P}_\pi \mathbf{P}_\sigma.$$

In particular, for all $\sigma \in S_n$, the permutation matrix $\mathbf{P}_\sigma$ is invertible with $(\mathbf{P}_\sigma)^{-1} = \mathbf{P}_{\sigma^{-1}}$.

Example 5.32 Considering $n = 3$. For

$$\sigma = \begin{pmatrix} 1 & 2 & 3 \\ 3 & 1 & 2 \end{pmatrix} \quad \text{and} \quad \pi = \begin{pmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{pmatrix} \quad \text{we have} \quad \pi \cdot \sigma = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix},$$

as well as

$$\mathbf{P}_\sigma = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{pmatrix}, \quad \mathbf{P}_\pi = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix} \quad \text{and} \quad \mathbf{P}_{\pi \cdot \sigma} = \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Thus we obtain

$$\mathbf{P}_{\pi \cdot \sigma} = \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{pmatrix} = \mathbf{P}_\pi \mathbf{P}_\sigma,$$

as claimed by [Proposition 5.31](#).

Proof of Proposition 5.31 The matrix $\mathbf{P}_1$ has column vectors given by $\vec{e}_1, \dots, \vec{e}_n$, hence $\mathbf{P}_1 = \mathbf{1}_n$.

Using [Proposition 2.20](#) and [Theorem 2.21](#) it is sufficient to show that for all $\pi, \sigma \in S_n$ we have $f_{\mathbf{P}_{\pi \cdot \sigma}} = f_{\mathbf{P}_\pi} \circ f_{\mathbf{P}_\sigma}$. For all $1 \leq i \leq n$, we obtain

$$f_{\mathbf{P}_\pi}(f_{\mathbf{P}_\sigma}(\vec{e}_i)) = f_{\mathbf{P}_\pi}(\vec{e}_{\sigma(i)}) = \vec{e}_{\pi(\sigma(i))} = \vec{e}_{(\pi \cdot \sigma)(i)} = f_{\mathbf{P}_{\pi \cdot \sigma}}(\vec{e}_i).$$

The two maps thus agree on the ordered basis $\mathbf{e} = (\vec{e}_1, \dots, \vec{e}_n)$ of $\mathbb{K}^n$, so that the second claim follows by applying [Lemma 3.88](#).

We have

$$\mathbf{P}_{\sigma \cdot \sigma^{-1}} = \mathbf{P}_1 = \mathbf{1}_n = \mathbf{P}_\sigma \mathbf{P}_{\sigma^{-1}}$$

showing that $\mathbf{P}_\sigma$ is invertible with inverse $(\mathbf{P}_\sigma)^{-1} = \mathbf{P}_{\sigma^{-1}}$. □

Definition 5.33 (Signature of a permutation) For $\sigma \in S_n$ we call $\text{sgn}(\sigma) = \det(\mathbf{P}_\sigma)$ its *signature*.

Remark 5.34

- Combining [Proposition 5.21](#) and [Proposition 5.31](#), we conclude that

$$\text{sgn}(\pi \cdot \sigma) = \text{sgn}(\pi) \text{sgn}(\sigma)$$

for all $\pi, \sigma \in S_n$.

- Since $\mathbf{P}_{\tau_{k,l}} = \mathbf{P}_{k,l}$ and $\det \mathbf{P}_{k,l} = -1$ by [Lemma 5.13](#), we conclude that

$$\text{sgn}(\tau_{k,l}) = -1$$

for all transpositions $\tau_{k,l} \in S_n$.

Similarly to elementary matrices being the building blocks of invertible matrices, transpositions are the building blocks of permutations:

Proposition 5.35 Let $n \in \mathbb{N}$ and $\sigma \in S_n$. Then there exists $m \geq 0$ and m transpositions $\tau_{k_1, l_1}, \dots, \tau_{k_m, l_m} \in S_n$ such that $\sigma = \tau_{k_m, l_m} \cdots \tau_{k_1, l_1}$, where we use the convention that 0 transpositions corresponds to the identity permutation.

Example 5.36 Let $n = 6$ and σ the permutation defined by the matrix

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 3 & 5 & 2 & 4 & 6 & 1 \end{pmatrix}.$$

To express it as a product of transposition, we write

$$\begin{array}{cccccc|l} 3 & 5 & 2 & 4 & 6 & 1 & \\ 3 & 2 & 5 & 4 & 6 & 1 & \tau_{2,3} \\ 1 & 2 & 5 & 4 & 6 & 3 & \tau_{1,6} \\ 1 & 2 & 5 & 4 & 3 & 6 & \tau_{5,6} \\ 1 & 2 & 3 & 4 & 5 & 6 & \tau_{3,5} \end{array}$$

so that $\sigma = \tau_{3,5} \tau_{5,6} \tau_{1,6} \tau_{2,3}$.

Proof of Proposition 5.35 We use induction. For $n = 1$ we have $\mathcal{X}_n = \{1\}$ and the only permutation is the identity permutation 1, so the statement is trivially true and hence anchored.

Inductive step: Assume $n \in \mathbb{N}$ and $n \geq 2$. We want to show that if the claim holds for S_{n-1} , then also for S_n . Let $\sigma \in S_n$ and define $k = \sigma(n)$. Then the permutation $\sigma_1 = \tau_{n,k}\sigma$ satisfies $\sigma_1(n) = \tau_{n,k}\sigma(n) = \tau_{n,k}(k) = n$ and hence does not permute n . Restricting σ_1 to the first $n - 1$ elements, we obtain a permutation of $\{1, \dots, n - 1\}$. By the induction hypothesis, we thus have $\tilde{m} \in \mathbb{N}$ and $\tau_{k_1, l_1}, \dots, \tau_{k_{\tilde{m}}, l_{\tilde{m}}} \in S_n$ such that

$$\sigma_1 = \tau_{k_{\tilde{m}}, l_{\tilde{m}}} \cdots \tau_{k_1, l_1} = \tau_{n,k}\sigma.$$

Since $\tau_{n,k}^2 = 1$, multiplying from the left with $\tau_{n,k}$ gives $\sigma = \tau_{n,k}\tau_{k_{\tilde{m}}, l_{\tilde{m}}} \cdots \tau_{k_1, l_1}$, the claim follows with $m = \tilde{m} + 1$. $\square$

Combining [Definition 5.33](#), [Remark 5.34](#) and [Proposition 5.35](#), we conclude:

Proposition 5.37 *Let $n \in \mathbb{N}$ and $\sigma \in S_n$. Then $\text{sgn}(\sigma) = \pm 1$. If σ is a product of m transpositions, then $\text{sgn}(\sigma) = (-1)^m$.*

Remark 5.38 Permutations with $\text{sgn}(\sigma) = 1$ are called *even* and permutations with $\text{sgn}(\sigma) = -1$ are called *odd*, since they arise from the composition of an even or odd number of transpositions, respectively.

5.6 The Leibniz formula

Besides the Laplace expansion, there is also a formula for the determinant which relies on permutations. As a warm-up, we first consider the case $n = 2$. Using the linearity of the determinant in the first row, we obtain

$$\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \det \begin{pmatrix} a & 0 \\ c & d \end{pmatrix} + \det \begin{pmatrix} 0 & b \\ c & d \end{pmatrix},$$

where $a, b, c, d \in \mathbb{K}$. Using the linearity of the determinant in the second row, we can further decompose the two above summands

$$\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \underbrace{\det \begin{pmatrix} a & 0 \\ c & 0 \end{pmatrix} + \det \begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix}}_{=\det \begin{pmatrix} a & 0 \\ c & d \end{pmatrix}} + \underbrace{\det \begin{pmatrix} 0 & b \\ c & 0 \end{pmatrix} + \det \begin{pmatrix} 0 & b \\ 0 & d \end{pmatrix}}_{=\det \begin{pmatrix} 0 & b \\ c & d \end{pmatrix}}$$

The first and fourth summand are *always zero* due to the occurrence of a zero column. The second and third summand are *possibly nonzero* (it might still happen that they are zero in the case where some of a, b, c, d are zero). In any case, we get

$$\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \det \begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix} + \det \begin{pmatrix} 0 & b \\ c & 0 \end{pmatrix}.$$

We can proceed analogously in general. Let $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n} \in M_{n,n}(\mathbb{K})$. We denote the rows of $\mathbf{A}$ by $\{\vec{\alpha}_1, \dots, \vec{\alpha}_n\}$. Using the linearity of $\det$ in the first row, we can write

$$\det \mathbf{A} = \det \begin{pmatrix} A_{11} & 0 & 0 & \cdots & 0 \\ & & \vec{\alpha}_2 & & \\ & & \vdots & & \\ & & \vec{\alpha}_n & & \end{pmatrix} + \det \begin{pmatrix} 0 & A_{12} & 0 & \cdots & 0 \\ & & \vec{\alpha}_2 & & \\ & & \vdots & & \\ & & \vec{\alpha}_n & & \end{pmatrix} + \cdots$$

$$\cdots + \det \begin{pmatrix} 0 & 0 & 0 & \cdots & A_{1n} \\ & & \vec{\alpha}_2 & & \\ & & \vdots & & \\ & & \vec{\alpha}_n & & \end{pmatrix}.$$

We can now use the linearity in the second row and proceed in the same fashion with each of the above summands. We continue this procedure until the n -th row. As a result, we can write

$$(5.7) \quad \det \mathbf{A} = \sum_{k=1}^{n^n} \det \mathbf{M}_k$$

where each of the matrices $\mathbf{M}_k$ has exactly one possibly nonzero entry in each row. As above, some of the matrices $\mathbf{M}_k$ will have a zero column so that their determinant vanishes. The matrices $\mathbf{M}_k$ without a zero column must have exactly one possibly nonzero entry in each row and each column. We can thus write the matrices $\mathbf{M}_k$ with possibly nonzero determinant as

$$\mathbf{M}_k = \sum_{i=1}^n A_{\sigma(i)i} \mathbf{E}_{\sigma(i),i}$$

for some permutation $\sigma \in S_n$. Every permutation of $\{1, \dots, n\}$ occurs precisely once in the expansion (5.7), hence we can write

$$\det \mathbf{A} = \sum_{\sigma \in S_n} \det \left(\sum_{i=1}^n A_{\sigma(i)i} \mathbf{E}_{\sigma(i),i} \right),$$

where the notation $\sum_{\sigma \in S_n}$ means that we sum over all possible permutations of $\{1, \dots, n\}$. We will next write the matrix $\sum_{i=1}^n A_{\sigma(i)i} \mathbf{E}_{\sigma(i),i}$ differently. To this end notice that for all $\sigma \in S_n$, the permutation matrix $\mathbf{P}_\sigma$ can be written as $\mathbf{P}_\sigma = \sum_{i=1}^n \mathbf{E}_{\sigma(i),i}$. Furthermore, the diagonal matrix

$$\mathbf{D}_\sigma = \begin{pmatrix} A_{\sigma(1)1} & & & \\ & A_{\sigma(2)2} & & \\ & & \ddots & \\ & & & A_{\sigma(n)n} \end{pmatrix}$$

can be written as $\mathbf{D}_\sigma = \sum_{j=1}^n A_{\sigma(j)j} \mathbf{E}_{j,j}$. Therefore, using Lemma 4.4, we obtain

$$\mathbf{P}_\sigma \mathbf{D}_\sigma = \sum_{i=1}^n \mathbf{E}_{\sigma(i),i} \sum_{j=1}^n A_{\sigma(j)j} \mathbf{E}_{j,j} = \sum_{i=1}^n \sum_{j=1}^n A_{\sigma(j)j} \mathbf{E}_{\sigma(i),i} \mathbf{E}_{j,j} = \sum_{i=1}^n A_{\sigma(i)i} \mathbf{E}_{\sigma(i),i},$$

We thus have the formula

$$\det \mathbf{A} = \sum_{\sigma \in S_n} \det(\mathbf{P}_\sigma \mathbf{D}_\sigma) = \sum_{\sigma \in S_n} \operatorname{sgn}(\sigma) \det(\mathbf{D}_\sigma),$$

where we use the product rule Proposition 5.21 and the definition of the signature of a permutation. By Proposition 5.24, the determinant of a diagonal matrix is the product of its diagonal entries, hence we obtain

$$\det \mathbf{A} = \sum_{\sigma \in S_n} \operatorname{sgn}(\sigma) \prod_{i=1}^n A_{\sigma(i)i}.$$

Finally, writing $\pi = \sigma^{-1}$, we have

$$\prod_{i=1}^n A_{\sigma(i)i} = \prod_{j=1}^n A_{j\pi(j)}.$$

We have thus shown:

Proposition 5.39 (Leibniz formula for the determinant) *Let $n \in \mathbb{N}$ and $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n} \in M_{n,n}(\mathbb{K})$. Then we have*

$$(5.8) \quad \det(\mathbf{A}) = \sum_{\sigma \in S_n} \text{sgn}(\sigma) \prod_{i=1}^n A_{\sigma(i)i} = \sum_{\pi \in S_n} \text{sgn}(\pi) \prod_{j=1}^n A_{j\pi(j)}.$$

Example 5.40 For $n = 3$ we have six permutations

$$\begin{aligned} \sigma_1 &= \begin{pmatrix} 1 & 2 & 3 \\ 1 & 2 & 3 \end{pmatrix}, & \sigma_2 &= \begin{pmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{pmatrix}, & \sigma_3 &= \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix} \\ \sigma_4 &= \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{pmatrix}, & \sigma_5 &= \begin{pmatrix} 1 & 2 & 3 \\ 3 & 1 & 2 \end{pmatrix}, & \sigma_6 &= \begin{pmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{pmatrix}. \end{aligned}$$

For $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq 3} \in M_{3,3}(\mathbb{K})$, the Leibniz formula gives

$$\begin{aligned} \det(\mathbf{A}) &= \text{sgn}(\sigma_1)A_{11}A_{22}A_{33} + \text{sgn}(\sigma_2)A_{11}A_{23}A_{32} + \text{sgn}(\sigma_3)A_{12}A_{21}A_{33} \\ &\quad + \text{sgn}(\sigma_4)A_{12}A_{23}A_{31} + \text{sgn}(\sigma_5)A_{13}A_{21}A_{32} + \text{sgn}(\sigma_6)A_{13}A_{22}A_{31}, \end{aligned}$$

so that in agreement with [Example 5.19](#), we obtain

$$\begin{aligned} \det \mathbf{A} &= A_{11}A_{22}A_{33} - A_{11}A_{23}A_{32} - A_{12}A_{21}A_{33} \\ &\quad + A_{12}A_{23}A_{31} + A_{13}A_{21}A_{32} - A_{13}A_{22}A_{31}. \end{aligned}$$

Remark 5.41 [Exercise 5.49](#) has two important consequences. Since the transpose turns the rows of a matrix into columns and vice versa, we conclude:

- the determinant is also multilinear and alternating, when thought of as a map $(\mathbb{K}^n)^n \rightarrow \mathbb{K}$, that is, when taking n columns vectors as an input. In particular, the determinant is also linear in each column;
- the Laplace expansion is also valid if we expand with respect to a row, that is, for $\mathbf{A} \in M_{n,n}(\mathbb{K})$ and $1 \leq l \leq n$, we have

$$\det(\mathbf{A}) = \sum_{k=1}^n (-1)^{k+l} [\mathbf{A}]_{lk} \det(\mathbf{A}^{(l,k)}).$$

Example 5.42 ($\heartsuit$ – not examinable) For $n \in \mathbb{N}$ and a vector $\vec{x} = (x_i)_{1 \leq i \leq n} \in \mathbb{K}^n$ we can form a matrix $\mathbf{V}_{\vec{x}} = (V_{ij})_{1 \leq i, j \leq n} \in M_{n,n}(\mathbb{K})$ with $V_{ij} = x_i^{j-1}$, that is,

$$\mathbf{V}_{\vec{x}} = \begin{pmatrix} 1 & x_1 & (x_1)^2 & \cdots & (x_1)^{n-1} \\ 1 & x_2 & (x_2)^2 & \cdots & (x_2)^{n-1} \\ 1 & x_3 & (x_3)^2 & \cdots & (x_3)^{n-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_n & (x_n)^2 & \cdots & (x_n)^{n-1} \end{pmatrix}.$$

Such matrices are known as *Vandermonde matrices* and the determinant of a Vandermonde matrix is known as a *Vandermonde determinant*, they satisfy

$$\det(\mathbf{V}_{\vec{x}}) = \prod_{1 \leq i < j \leq n} (x_j - x_i).$$

Sketch of a proof We can define a function $f : \mathbb{K}^n \rightarrow \mathbb{K}$, $\vec{x} \mapsto \det(\mathbf{V}_{\vec{x}})$. By the Leibniz formula, the function f is a polynomial in the variables x_i with integer coefficients. If we freeze all variables of f except the ℓ -th variable, then we obtain a function $g_\ell : \mathbb{K} \rightarrow \mathbb{K}$ of one variable x_ℓ . For $1 \leq i \leq n$ with $i \neq \ell$ we have $g_\ell(x_i) = 0$, since we compute the determinant of a matrix with two identical rows, the ℓ -th row and the i -th row. Factoring the zeros, we can thus write $g_\ell(x_\ell) = q_\ell(x_\ell) \prod_{1 \leq i \leq n, i \neq \ell} (x_\ell - x_i)$ for some polynomial q_ℓ . We can repeat this argument for all ℓ and hence can write $\det(\mathbf{V}_{\vec{x}}) = q(\vec{x}) \prod_{1 \leq i < j \leq n} (x_j - x_i)$ for some polynomial $q(\vec{x})$. The Leibniz formula implies that the sum of the exponents of all the factors x_i in $\det(\mathbf{V}_{\vec{x}})$ must be $\frac{1}{2}n(n-1)$. The same holds true for $\prod_{1 \leq i < j \leq n} (x_j - x_i)$. It follows that q must be a constant. Using the Leibniz formula again, we see that the summand of $\det(\mathbf{V}_{\vec{x}})$ corresponding to the identity permutation is the product of the diagonal entries of $\mathbf{V}_{\vec{x}}$, that is, $x_2(x_3)^2 \cdots (x_n)^{n-1}$. Taking the first term in all factors of $\prod_{1 \leq i < j \leq n} (x_j - x_i)$, we also obtain $x_2(x_3)^2 \cdots (x_n)^{n-1}$, hence $\det(\mathbf{V}_{\vec{x}}) = \prod_{1 \leq i < j \leq n} (x_j - x_i)$, as claimed. $\square$

5.7 Cramer's rule

The determinant can be used to give a formula for the solution of a linear system of equations of the form $\mathbf{A}\vec{x} = \vec{b}$ for an invertible matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$, $\vec{b} \in \mathbb{K}^n$ and unknowns $\vec{x} \in \mathbb{K}^n$. This formula is often referred to as *Cramer's rule*. In order to derive it we start with definitions:

Definition 5.43 (Adjugate matrix — Video) Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{K})$ be a square matrix. The *adjugate matrix* of $\mathbf{A}$ is the $n \times n$ -matrix $\text{Adj}(\mathbf{A})$ whose entries are given by (notice the reverse order of i and j on the right hand side)

$$[\text{Adj}(\mathbf{A})]_{ij} = (-1)^{i+j} \det(\mathbf{A}^{(j,i)}), \quad 1 \leq i, j \leq n.$$

Example 5.44

$$\text{Adj}\left(\begin{pmatrix} a & b \\ c & d \end{pmatrix}\right) = \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}, \quad \text{Adj}\left(\begin{pmatrix} 1 & 1 & 2 \\ 0 & 2 & 1 \\ 1 & 0 & 2 \end{pmatrix}\right) = \begin{pmatrix} 4 & -2 & -3 \\ 1 & 0 & -1 \\ -2 & 1 & 2 \end{pmatrix}$$

The determinant and the adjugate matrix provide a formula for the inverse of a matrix:

Theorem 5.45 Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{K})$. Then we have

$$\text{Adj}(\mathbf{A})\mathbf{A} = \mathbf{A}\text{Adj}(\mathbf{A}) = \det(\mathbf{A})\mathbf{1}_n.$$

In particular, if $\mathbf{A}$ is invertible then

$$\mathbf{A}^{-1} = \frac{1}{\det(\mathbf{A})} \text{Adj}(\mathbf{A}).$$

Proof Let $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n}$. For $1 \leq i \leq n$ we obtain for the i -th diagonal entry

$$[\text{Adj}(\mathbf{A})\mathbf{A}]_{ii} = \sum_{k=1}^n (-1)^{i+k} \det(\mathbf{A}^{(k,i)}) A_{ki} = \det(\mathbf{A}),$$

where we use the Laplace expansion (5.5) of the determinant. The diagonal entries of $\text{Adj}(\mathbf{A})\mathbf{A}$ are thus all equal to $\det \mathbf{A}$. For $1 \leq i, j \leq n$ with $i \neq j$ we have

$$[\text{Adj}(\mathbf{A})\mathbf{A}]_{ij} = \sum_{k=1}^n (-1)^{i+k} \left(\det \mathbf{A}^{(k,i)} \right) A_{kj}.$$

We would like to interpret this last expression as a Laplace expansion. We consider a new matrix $\hat{\mathbf{A}} = (\hat{A}_{ij})_{1 \leq i, j \leq n}$ which is identical to $\mathbf{A}$, except that the i -th column of $\mathbf{A}$ is replaced with the j -th column of $\mathbf{A}$, that is, for $1 \leq k \leq n$, we have

$$(5.9) \quad \hat{A}_{kl} = \begin{cases} A_{kj}, & l = i, \\ A_{kl}, & l \neq i. \end{cases}$$

Then, for all $1 \leq k \leq n$ we have $\hat{\mathbf{A}}^{(k,i)} = \mathbf{A}^{(k,i)}$, since the only column in which $\mathbf{A}$ and $\hat{\mathbf{A}}$ are different is removed in $\mathbf{A}^{(k,i)}$. Using (5.9), the Laplace expansion of $\hat{\mathbf{A}}$ with respect to the i -th column gives

$$\begin{aligned} \det \hat{\mathbf{A}} &= \sum_{k=1}^n (-1)^{(i+k)} \hat{A}_{ki} \det(\hat{\mathbf{A}}^{(k,i)}) = \sum_{k=1}^n (-1)^{i+k} \left(\det \mathbf{A}^{(k,i)} \right) A_{kj} \\ &= [\text{Adj}(\mathbf{A})\mathbf{A}]_{ij} \end{aligned}$$

The matrix $\hat{\mathbf{A}}$ has a double occurrence of the i -th column, hence its column vectors are linearly dependent. Therefore $\hat{\mathbf{A}}$ is not invertible by Proposition 4.7 and so $\det \hat{\mathbf{A}} = [\text{Adj}(\mathbf{A})\mathbf{A}]_{ij} = 0$ by Corollary 5.22. The off-diagonal entries of $\text{Adj}(\mathbf{A})\mathbf{A}$ are thus all zero and we conclude $\text{Adj}(\mathbf{A})\mathbf{A} = \det(\mathbf{A})\mathbf{1}_n$. Using the row version of the Laplace expansion we can conclude analogously that $\mathbf{A} \text{Adj}(\mathbf{A}) = \det(\mathbf{A})\mathbf{1}_n$.

Finally, if $\mathbf{A}$ is invertible, then $\det \mathbf{A} \neq 0$ by Corollary 5.22, so that $\mathbf{A}^{-1} = \text{Adj}(\mathbf{A}) / \det(\mathbf{A})$, as claimed. $\square$

As a corollary we obtain:

Corollary 5.46 Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{K})$ be an invertible upper triangular matrix. Then $\mathbf{A}^{-1}$ is also an upper triangular matrix.

Remark 5.47 Taking the transpose also implies: Let $\mathbf{A} \in M_{n,n}(\mathbb{K})$ be an invertible lower triangular matrix. Then $\mathbf{A}^{-1}$ is also a lower triangular matrix.

Proof of Corollary 5.46 Write $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n}$. Using Theorem 5.45 it suffices to show that $\text{Adj}(\mathbf{A})$ is an upper triangular matrix. If $\mathbf{A}$ is an upper triangular matrix, then $A_{ij} = 0$ for all $i > j$. By definition we have

$$[\text{Adj}(\mathbf{A})]_{ij} = (-1)^{i+j} \det(\mathbf{A}^{(j,i)}), \quad 1 \leq i, j \leq n.$$

Notice that for $i > j$ every element below the diagonal of $\mathbf{A}^{(j,i)}$ is also below the diagonal of $\mathbf{A}$ and hence must be zero. It follows that $\mathbf{A}^{(j,i)}$ is an upper triangular matrix as well. Proposition 5.24 implies that the determinant of $\mathbf{A}^{(j,i)}$ is the product of its diagonal entries. Since $\mathbf{A}^{(j,i)}$ arises from the upper triangular matrix $\mathbf{A}$ by removing a row and a column, at least one of the diagonal entries of $\mathbf{A}^{(j,i)}$ must be zero and thus $\det \mathbf{A}^{(j,i)} = 0$ for $i > j$. We conclude that $\mathbf{A}^{-1}$ is an upper triangular matrix as well. $\square$

We now use [Theorem 5.45](#) to obtain a formula for the solution of the linear system $\mathbf{A}\vec{x} = \vec{b}$ for an invertible matrix $\mathbf{A}$. Multiplying from the left with $\mathbf{A}^{-1}$, we get

$$\vec{x} = \mathbf{A}^{-1}\vec{b} = \frac{1}{\det \mathbf{A}} \operatorname{Adj}(\mathbf{A})\vec{b}.$$

Writing $\vec{x} = (x_i)_{1 \leq i \leq n}$, multiplication with $\det \mathbf{A}$ gives for $1 \leq i \leq n$

$$x_i \det \mathbf{A} = \sum_{k=1}^n [\operatorname{Adj}(\mathbf{A})]_{ik} b_k = \sum_{k=1}^n (-1)^{i+k} \det(\mathbf{A}^{(k,i)}) b_k.$$

We can again interpret the right hand side as a Laplace expansion of the matrix $\hat{\mathbf{A}}_i$ obtained by replacing the i -th column of $\mathbf{A}$ with $\vec{b}$ and leaving $\mathbf{A}$ unchanged otherwise. Hence, we have for all $1 \leq i \leq n$

$$x_i = \frac{\det \hat{\mathbf{A}}_i}{\det \mathbf{A}}.$$

This formula is known as *Cramer's rule*. While this is a neat formula, it is rarely used in computing solutions to linear systems of equations due to the complexity of computing determinants.

Example 5.48 (Cramer's rule) We consider the system $\mathbf{A}\vec{x} = \vec{b}$ for

$$\mathbf{A} = \begin{pmatrix} 2 & 1 & 1 \\ 1 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix} \quad \text{and} \quad \vec{b} = \begin{pmatrix} -2 \\ 2 \\ 4 \end{pmatrix}.$$

Here we obtain

$$\hat{\mathbf{A}}_1 = \begin{pmatrix} -2 & 1 & 1 \\ 2 & 2 & 1 \\ 4 & 1 & 2 \end{pmatrix}, \quad \hat{\mathbf{A}}_2 = \begin{pmatrix} 2 & -2 & 1 \\ 1 & 2 & 1 \\ 1 & 4 & 2 \end{pmatrix}, \quad \hat{\mathbf{A}}_3 = \begin{pmatrix} 2 & 1 & -2 \\ 1 & 2 & 2 \\ 1 & 1 & 4 \end{pmatrix}.$$

We compute $\det \mathbf{A} = 4$, $\det \hat{\mathbf{A}}_1 = -12$, $\det \hat{\mathbf{A}}_2 = 4$ and $\det \hat{\mathbf{A}}_3 = 12$ so that Cramer's rule gives indeed the correct solution

$$\vec{x} = \frac{1}{4} \begin{pmatrix} -12 \\ 4 \\ 12 \end{pmatrix} = \begin{pmatrix} -3 \\ 1 \\ 3 \end{pmatrix}.$$

Exercises

Exercise 5.49 Use the Leibniz formula to show that

$$\det(\mathbf{A}) = \det(\mathbf{A}^T)$$

for all $\mathbf{A} \in M_{n,n}(\mathbb{K})$.

Endomorphisms

WEEK 10

6.1 Sums, direct sums and complements

In this chapter we study linear mappings from a vector space to itself.

Definition 6.1 (Endomorphism — Video) A linear map $g : V \rightarrow V$ from a $\mathbb{K}$ -vector space V to itself is called an *endomorphism*. An endomorphism that is also an isomorphism is called an *automorphism*.

Before we develop the theory of endomorphisms, we introduce some notions for subspaces.

Definition 6.2 (Sum of subspaces — Video) Let V be a $\mathbb{K}$ -vector space, $n \in \mathbb{N}$ and $U_1, \dots, U_n$ be vector subspaces of V . The set

$$\sum_{i=1}^n U_i = U_1 + U_2 + \dots + U_n = \{v \in V \mid v = u_1 + u_2 + \dots + u_n \text{ for } u_i \in U_i\}$$

is called the *sum of the subspaces* U_i .

Recall that by [Proposition 3.27](#), the intersection of two subspaces is again a subspace, whereas the union of two subspaces fails to be a subspace in general. However, subspaces do behave nicely with regards to sums:

Proposition 6.3 The sum of the subspaces $U_i \subset V$, $i = 1, \dots, n$ is again a vector subspace.

Proof The sum $\sum_{i=1}^n U_i$ is non-empty, since it contains the zero vector 0_V . Let v and $v' \in \sum_{i=1}^n U_i$ so that

$$v = v_1 + v_2 + \dots + v_n \quad \text{and} \quad v' = v'_1 + v'_2 + \dots + v'_n$$

for vectors $v_i, v'_i \in U_i$, $i = 1, \dots, n$. Each U_i is a vector subspace of V . Therefore, for all scalars $s, t \in \mathbb{K}$, the vector $sv_i + tv'_i$ is an element of U_i , $i = 1, \dots, n$. Thus

$$sv + tv' = sv_1 + tv'_1 + \dots + sv_n + tv'_n$$

is an element of $U_1 + \dots + U_n$. By [Definition 3.21](#), it follows that $U_1 + \dots + U_n$ is a vector subspace of V . $\square$

Remark 6.4 Notice that $U_1 + \dots + U_n$ is the smallest vector subspace of V containing all vector subspaces U_i , $i = 1, \dots, n$.

If each vector in the sum is in a unique way the sum of vectors from the subspaces we say the subspaces are in direct sum:

Definition 6.5 (Direct sum of subspaces) Let V be a $\mathbb{K}$ -vector space, $n \in \mathbb{N}$ and $U_1, \dots, U_n$ be vector subspaces of V . The subspaces $U_1, \dots, U_n$ are said to be in *direct sum* if each vector $w \in W = U_1 + \dots + U_n$ is in a unique way the sum of vectors $v_i \in U_i$ for $1 \leq i \leq n$. That is, if $w = v_1 + v_2 + \dots + v_n = v'_1 + v'_2 + \dots + v'_n$ for vectors $v_i, v'_i \in U_i$, then $v_i = v'_i$ for all $1 \leq i \leq n$. We write

$$\bigoplus_{i=1}^n U_i$$

in case the subspaces $U_1, \dots, U_n$ are in direct sum.

Example 6.6 Let $n \in \mathbb{N}$ and $V = \mathbb{K}^n$ as well as $U_i = \text{span}\{\vec{e}_i\}$, where $\{\vec{e}_1, \dots, \vec{e}_n\}$ denotes the standard basis of $\mathbb{K}^n$. Then $\mathbb{K}^n = \bigoplus_{i=1}^n U_i$.

Remark 6.7

- Two subspaces U_1, U_2 of V are in direct sum if and only if $U_1 \cap U_2 = \{0_V\}$. Indeed, suppose $U_1 \cap U_2 = \{0_V\}$ and consider $w = v_1 + v_2 = v'_1 + v'_2$ with $v_i, v'_i \in U_i$ for $i = 1, 2$. We then have $v_1 - v'_1 = v'_2 - v_2 \in U_2$, since U_2 is a subspace. Since U_1 is a subspace as well, we also have $v_1 - v'_1 \in U_1$. Since $v_1 - v'_1$ lies both in U_1 and U_2 , we must have $v_1 - v'_1 = 0_V = v'_2 - v_2$. Conversely, suppose U_1, U_2 are in direct sum and let $w \in (U_1 \cap U_2)$. We can write $w = w + 0_V = 0_V + w$, since $w \in U_1$ and $w \in U_2$. Since U_1, U_2 are in direct sum, we must have $w = 0_V$, hence $U_1 \cap U_2 = \{0_V\}$.
- Observe that if the subspaces $U_1, \dots, U_n$ are in direct sum and $v_i \in U_i$ with $v_i \neq 0_V$ for $1 \leq i \leq n$, then the vectors $\{v_1, \dots, v_n\}$ are linearly independent. Indeed, if $s_1, \dots, s_n$ are scalars such that

$$s_1 v_1 + s_2 v_2 + \dots + s_n v_n = 0_V = 0_V + 0_V + \dots + 0_V,$$

then $s_i v_i = 0_V$ for all $1 \leq i \leq n$. By assumption $v_i \neq 0_V$ and hence $s_i = 0$ for all $1 \leq i \leq n$.

Proposition 6.8 Let $n \in \mathbb{N}$, V be a finite dimensional $\mathbb{K}$ -vector space and $U_1, \dots, U_n$ be subspaces of V . Let $\mathbf{b}_i$ be an ordered basis of U_i for $1 \leq i \leq n$. Then we have:

- The tuple of vectors obtained by listing all the vectors of the bases $\mathbf{b}_i$ is a basis of V if and only if $V = \bigoplus_{i=1}^n U_i$.
- $\dim(U_1 + \dots + U_n) \leq \dim(U_1) + \dots + \dim(U_n)$ with equality if and only if the subspaces $U_1, \dots, U_n$ are in direct sum.

Proof Part of an exercise. □

Definition 6.9 (Complement to a subspace) Let V be a $\mathbb{K}$ -vector space and $U \subset V$ a subspace. A subspace U' of V such that $V = U \oplus U'$ is called a *complement* to U .

Example 6.10 Notice that a complement need not be unique. Consider $V = \mathbb{R}^2$ and $U = \text{span}\{\vec{e}_1\}$. Let $v \in V$. Then the subspace $U' = \text{span}\{v\}$ is a complement to U , provided $\vec{e}_1, \vec{v}$ are linearly independent.

Corollary 6.11 (Existence of a complement) *Let U be a subspace of a finite dimensional $\mathbb{K}$ -vector space V . Then there exists a subspace U' so that $V = U \oplus U'$.*

Proof Suppose $(v_1, \dots, v_m)$ is an ordered basis of U . By [Theorem 3.64](#), there exists a basis $\{v_1, \dots, v_m, v_{m+1}, \dots, v_n\}$ of V . Defining $U' = \text{span}\{v_{m+1}, \dots, v_n\}$, [Proposition 6.8](#) implies the claim. $\square$

The dimension of a sum of two subspaces equals the sum of the dimensions of the subspaces minus the dimension of the intersection:

Proposition 6.12 *Let V be a finite dimensional $\mathbb{K}$ -vector space and U_1, U_2 subspaces of V . Then we have*

$$\dim(U_1 + U_2) = \dim(U_1) + \dim(U_2) - \dim(U_1 \cap U_2).$$

Proof Let $r = \dim(U_1 \cap U_2)$ and let $\{u_1, \dots, u_r\}$ be a basis of $U_1 \cap U_2$. These vectors are linearly independent and elements of U_1 , hence by [Theorem 3.64](#), there exist vectors $v_1, \dots, v_{m-r}$ so that $S_1 = \{u_1, \dots, u_r, v_1, \dots, v_{m-r}\}$ is a basis of U_1 . Likewise there exist vectors $w_1, \dots, w_{n-r}$ such that $S_2 = \{u_1, \dots, u_r, w_1, \dots, w_{n-r}\}$ is a basis of U_2 . Here $m = \dim U_1$ and $n = \dim U_2$.

Now consider the set $\mathcal{S} = \{u_1, \dots, u_r, v_1, \dots, v_{m-r}, w_1, \dots, w_{n-r}\}$ consisting of $r + m - r + n - r = n + m - r$ vectors. If this set is a basis of $U_1 + U_2$, then the claim follows, since then $\dim(U_1 + U_2) = n + m - r = \dim(U_1) + \dim(U_2) - \dim(U_1 \cap U_2)$.

We first show that $\mathcal{S}$ generates $U_1 + U_2$. Let $y \in U_1 + U_2$ so that $y = x_1 + x_2$ for vectors $x_1 \in U_1$ and $x_2 \in U_2$. Since S_1 is a basis of U_1 , we can write x_1 as a linear combination of elements of S_1 . Likewise we can write x_2 as a linear combination of elements of S_2 . It follows that $\mathcal{S}$ generates $U_1 + U_2$.

We need to show that $\mathcal{S}$ is linearly independent. So suppose we have scalars $s_1, \dots, s_r, t_1, \dots, t_{m-r},$ and $r_1, \dots, r_{n-r}$, so that

$$\underbrace{s_1 u_1 + \dots + s_r u_r}_{=u} + \underbrace{t_1 v_1 + \dots + t_{m-r} v_{m-r}}_{=v} + \underbrace{r_1 w_1 + \dots + r_{n-r} w_{n-r}}_{=w} = 0_V.$$

Equivalently, $w = -u - v$ so that $w \in U_1$. Since w is a linear combination of elements of S_2 , we also have $w \in U_2$. Therefore, $w \in U_1 \cap U_2$ and there exist scalars $\hat{s}_1, \dots, \hat{s}_r$ such that

$$w = \underbrace{\hat{s}_1 u_1 + \dots + \hat{s}_r u_r}_{=\hat{u}}$$

This is equivalent to $w - \hat{u} = 0_V$, or written out

$$r_1 w_1 + \dots + r_{n-r} w_{n-r} - \hat{s}_1 u_1 - \dots - \hat{s}_r u_r = 0_V.$$

Since the vectors $\{u_1, \dots, u_r, w_1, \dots, w_{n-r}\}$ are linearly independent, we conclude that $r_1 = \dots = r_{n-r} = \hat{s}_1 = \dots = \hat{s}_r = 0$. It follows that $w = 0_V$ and hence $u + v = 0_V$. Again, since $\{u_1, \dots, u_r, v_1, \dots, v_{m-r}\}$ are linearly independent, we conclude that $s_1 = \dots = s_r = t_1 = \dots = t_{m-r} = 0$ and we are done. $\square$

6.2 Invariants of endomorphisms

Let V be a finite dimensional vector space equipped with an ordered basis $\mathbf{b}$ and $g : V \rightarrow V$ an endomorphism. Recall from [Theorem 3.107](#) that if we consider another ordered basis $\mathbf{b}'$ of V , then

$$\mathbf{M}(g, \mathbf{b}', \mathbf{b}') = \mathbf{C} \mathbf{M}(g, \mathbf{b}, \mathbf{b}) \mathbf{C}^{-1},$$

where we write $\mathbf{C} = \mathbf{C}(\mathbf{b}, \mathbf{b}')$ for the change of basis matrix. This motivates the following definition:

Definition 6.13 (Similar / conjugate matrices) Let $n \in \mathbb{N}$ and $\mathbf{A}, \mathbf{A}' \in M_{n,n}(\mathbb{K})$. The matrices $\mathbf{A}$ and $\mathbf{A}'$ are called *similar* or *conjugate over $\mathbb{K}$* if there exists an invertible matrix $\mathbf{C} \in M_{n,n}(\mathbb{K})$ such that

$$\mathbf{A}' = \mathbf{C} \mathbf{A} \mathbf{C}^{-1}.$$

Similarity of matrices over $\mathbb{K}$ is an *equivalence relation*:

Proposition 6.14 Let $n \in \mathbb{N}$ and $\mathbf{A}, \mathbf{B}, \mathbf{X} \in M_{n,n}(\mathbb{K})$. Then we have

- (i) $\mathbf{A}$ is similar to itself;
- (ii) $\mathbf{A}$ is similar to $\mathbf{B}$ then $\mathbf{B}$ is similar to $\mathbf{A}$;
- (iii) If $\mathbf{A}$ is similar to $\mathbf{B}$ and $\mathbf{B}$ is similar to $\mathbf{X}$, then $\mathbf{A}$ is also similar to $\mathbf{X}$.

Proof (i) We take $\mathbf{C} = \mathbf{1}_n$.

(ii) Suppose $\mathbf{A}$ is similar to $\mathbf{B}$ so that $\mathbf{B} = \mathbf{C} \mathbf{A} \mathbf{C}^{-1}$ for some invertible matrix $\mathbf{C} \in M_{n,n}(\mathbb{K})$. Multiplying with $\mathbf{C}^{-1}$ from the left and $\mathbf{C}$ from the right, we get

$$\mathbf{C}^{-1} \mathbf{B} \mathbf{C} = \mathbf{C}^{-1} \mathbf{C} \mathbf{A} \mathbf{C}^{-1} \mathbf{C} = \mathbf{A},$$

so that the similarity follows for the choice $\hat{\mathbf{C}} = \mathbf{C}^{-1}$.

(iii) We have $\mathbf{B} = \mathbf{C} \mathbf{A} \mathbf{C}^{-1}$ and $\mathbf{X} = \mathbf{D} \mathbf{B} \mathbf{D}^{-1}$ for invertible matrices $\mathbf{C}, \mathbf{D}$. Then we get

$$\mathbf{X} = \mathbf{D} \mathbf{C} \mathbf{A} \mathbf{C}^{-1} \mathbf{D}^{-1},$$

so that the similarity follows for the choice $\hat{\mathbf{C}} = \mathbf{D} \mathbf{C}$. □

Remark 6.15

- Because of (ii) in particular, one can say that two matrices $\mathbf{A}$ and $\mathbf{B}$ are similar without ambiguity.
- [Theorem 3.107](#) shows that $\mathbf{A}$ and $\mathbf{B}$ are similar if and only if there exists an endomorphism g of $\mathbb{K}^n$ such that $\mathbf{A}$ and $\mathbf{B}$ represent g with respect to two ordered bases of $\mathbb{K}^n$.

One might wonder whether there exist functions $f : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ which are *invariant* under conjugation, that is, f satisfies $f(\mathbf{C} \mathbf{A} \mathbf{C}^{-1}) = f(\mathbf{A})$ for all $\mathbf{A} \in M_{n,n}(\mathbb{K})$ and all invertible matrices $\mathbf{C} \in M_{n,n}(\mathbb{K})$. We have already seen an example of such a function, namely the determinant. Indeed using the product rule [Proposition 5.21](#) and [Corollary 5.22](#), we compute

$$\begin{aligned} \det(\mathbf{C} \mathbf{A} \mathbf{C}^{-1}) &= \det(\mathbf{C} \mathbf{A}) \det(\mathbf{C}^{-1}) = \det(\mathbf{C}) \det(\mathbf{A}) \det(\mathbf{C}^{-1}) \\ (6.1) \qquad \qquad &= \det(\mathbf{A}). \end{aligned}$$

Because of this fact, the following definition makes sense:

Definition 6.16 (Determinant of an endomorphism) Let V be a finite dimensional $\mathbb{K}$ -vector space and $g : V \rightarrow V$ an endomorphism. We define

$$\det(g) = \det(\mathbf{M}(g, \mathbf{b}, \mathbf{b}))$$

where $\mathbf{b}$ is any ordered basis of V . By [Theorem 3.107](#) and [\(6.1\)](#), the scalar $\det(g)$ is independent of the chosen ordered basis.

Another example of a scalar that we can associate to an endomorphism is the so-called *trace*. Like for the determinant, we first define the trace for matrices. Luckily, the trace is a lot simpler to define:

Definition 6.17 (Trace of a matrix) Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{K})$. The sum $\sum_{i=1}^n [\mathbf{A}]_{ii}$ of its diagonal entries is called the *trace of $\mathbf{A}$* and denoted by $\text{Tr}(\mathbf{A})$ or $\text{Tr } \mathbf{A}$.

Example 6.18 For all $n \in \mathbb{N}$ we have $\text{Tr}(\mathbf{1}_n) = n$. For

$$\mathbf{A} = \begin{pmatrix} 2 & 1 & 1 \\ 1 & 2 & 1 \\ 1 & 1 & 3 \end{pmatrix}$$

we have $\text{Tr}(\mathbf{A}) = 2 + 2 + 3 = 7$.

The trace of a product of square matrices is independent of the order of multiplication:

Proposition 6.19 Let $n \in \mathbb{N}$ and $\mathbf{A}, \mathbf{B} \in M_{n,n}(\mathbb{K})$. Then we have

$$\text{Tr}(\mathbf{AB}) = \text{Tr}(\mathbf{BA}).$$

Proof Let $\mathbf{A} = (A_{ij})_{1 \leq i,j \leq n}$ and $\mathbf{B} = (B_{ij})_{1 \leq i,j \leq n}$. Then

$$[\mathbf{AB}]_{ij} = \sum_{k=1}^n A_{ik} B_{kj} \quad \text{and} \quad [\mathbf{BA}]_{kj} = \sum_{i=1}^n B_{ki} A_{ij},$$

so that

$$\text{Tr}(\mathbf{AB}) = \sum_{i=1}^n \sum_{k=1}^n A_{ik} B_{ki} = \sum_{k=1}^n \sum_{i=1}^n B_{ki} A_{ik} = \text{Tr}(\mathbf{BA}).$$

□

Using the previous proposition, we obtain

$$(6.2) \quad \text{Tr}(\mathbf{CAC}^{-1}) = \text{Tr}(\mathbf{AC}^{-1}\mathbf{C}) = \text{Tr}(\mathbf{A}).$$

As for the determinant, the following definition thus makes sense:

Definition 6.20 (Trace of an endomorphism) Let V be a finite dimensional $\mathbb{K}$ -vector space and $g : V \rightarrow V$ an endomorphism. We define

$$\text{Tr}(g) = \text{Tr}(\mathbf{M}(g, \mathbf{b}, \mathbf{b}))$$

where $\mathbf{b}$ is any ordered basis of V . By [Theorem 3.107](#) and (6.2), the scalar $\text{Tr}(g)$ is independent of the chosen ordered basis.

The trace and determinant of endomorphisms behave nicely with respect to composition of maps:

Proposition 6.21 *Let V be a finite dimensional $\mathbb{K}$ -vector space. Then, for all endomorphisms $f, g : V \rightarrow V$ we have*

- (i) $\text{Tr}(f \circ g) = \text{Tr}(g \circ f)$;
- (ii) $\det(f \circ g) = \det(f) \det(g)$.

Proof (i) Fix an ordered basis $\mathbf{b}$ of V . Then, using [Corollary 3.101](#) and [Proposition 6.19](#), we obtain

$$\begin{aligned} \text{Tr}(f \circ g) &= \text{Tr}(\mathbf{M}(f \circ g, \mathbf{b}, \mathbf{b})) = \text{Tr}(\mathbf{M}(f, \mathbf{b}, \mathbf{b})\mathbf{M}(g, \mathbf{b}, \mathbf{b})) \\ &= \text{Tr}(\mathbf{M}(g, \mathbf{b}, \mathbf{b})\mathbf{M}(f, \mathbf{b}, \mathbf{b})) = \text{Tr}(\mathbf{M}(g \circ f, \mathbf{b}, \mathbf{b})) = \text{Tr}(g \circ f). \end{aligned}$$

The proof of (ii) is analogous, but we use [Proposition 5.21](#) instead of [Proposition 6.19](#). $\square$

We also have:

Proposition 6.22 *Let V be a finite dimensional $\mathbb{K}$ -vector space and $g : V \rightarrow V$ an endomorphism. Then the following statements are equivalent:*

- (i) g is injective;
- (ii) g is surjective;
- (iii) g is bijective;
- (iv) $\det(g) \neq 0$.

Proof The equivalence of the first three statements follows from [Corollary 3.77](#). We fix an ordered basis $\mathbf{b}$ of V . Suppose g is bijective with inverse $g^{-1} : V \rightarrow V$. Then we have $\det(g \circ g^{-1}) = \det(g) \det(g^{-1}) = \det(\text{Id}_V) = \det(\mathbf{M}(\text{Id}_V, \mathbf{b}, \mathbf{b})) = \det(\mathbf{1}_{\dim V}) = 1$.

It follows that $\det(g) \neq 0$ and moreover that

$$\det(g^{-1}) = \frac{1}{\det g}.$$

Conversely, suppose that $\det g \neq 0$. Then $\det \mathbf{M}(g, \mathbf{b}, \mathbf{b}) \neq 0$ so that $\mathbf{M}(g, \mathbf{b}, \mathbf{b})$ is invertible by [Corollary 5.22](#) and [Proposition 3.102](#) implies that g is bijective. $\square$

Remark 6.23 Notice that [Proposition 6.22](#) is wrong for infinite dimensional vector spaces. Consider $V = \mathbb{K}^\infty$, the $\mathbb{K}$ -vector space of sequences from [Example 3.6](#). The endomorphism $g : V \rightarrow V$ defined by $(x_1, x_2, x_3, \dots) \mapsto (0, x_1, x_2, x_3, \dots)$ is injective but not surjective.

6.3 Eigenvectors and eigenvalues

Mappings g that have the same domain and codomain allow for the notion of a fixed point. Recall that an element x of a set $\mathcal{X}$ is called a *fixed point* of a mapping $g : \mathcal{X} \rightarrow \mathcal{X}$ if $g(x) = x$, that is, x agrees with its image under g . In Linear Algebra, a generalisation of the notion of a fixed point is that of an eigenvector. A vector $v \in V$ is called an *eigenvector* of the linear map $g : V \rightarrow V$ if v is merely scaled when applying g to v , that is, there exists a scalar $\lambda \in \mathbb{K}$ – called *eigenvalue* – such that $g(v) = \lambda v$. Clearly, the zero vector 0_V will satisfy this condition for every choice of scalar λ . For this reason, eigenvectors are usually required to be different from the zero vector. In this terminology, fixed points v of g are simply eigenvectors with eigenvalue 1, since they satisfy $g(v) = v = 1v$.

It is natural to ask whether a linear map $g : V \rightarrow V$ always admits an eigenvector. In the remaining part of this chapter we will answer this question and further develop our theory of linear maps, specifically endomorphisms. We start with some precise definitions.

Definition 6.24 (Eigenvector, eigenspace, eigenvalue — Video) Let $g : V \rightarrow V$ be an endomorphism of a $\mathbb{K}$ -vector space V .

- An *eigenvector* with *eigenvalue* $\lambda \in \mathbb{K}$ is a *non-zero* vector $v \in V$ such that $g(v) = \lambda v$.
- If $\lambda \in \mathbb{K}$ is an eigenvalue of g , the λ -*eigenspace* $\text{Eig}_g(\lambda)$ is the subspace of vectors $v \in V$ satisfying $g(v) = \lambda v$.
- The dimension of $\text{Eig}_g(\lambda)$ is called the *geometric multiplicity* of the eigenvalue λ .
- The set of all eigenvalues of g is called the *spectrum* of g .
- For $\mathbf{A} \in M_{n,n}(\mathbb{K})$ we speak of eigenvalues, eigenvectors, eigenspaces and spectrum to mean those of the endomorphism $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$.

Remark 6.25 By definition, the zero vector 0_V is not an eigenvector, it is however an element of the eigenspace $\text{Eig}_g(\lambda)$ for every eigenvalue λ .

Example 6.26

- The scalar 0 is an eigenvalue of an endomorphism $g : V \rightarrow V$ if and only if the kernel of g is different from $\{0_V\}$. In the case where the kernel of f does not only consist of the zero vector, we have $\text{Ker } g = \text{Eig}_g(0)$ and the geometric multiplicity of 0 is the nullity of g .
- The endomorphism $f_{\mathbf{D}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$ associated to a diagonal matrix with distinct diagonal entries

$$\mathbf{D} = \begin{pmatrix} \lambda_1 & & & \\ & \lambda_2 & & \\ & & \ddots & \\ & & & \lambda_n \end{pmatrix}$$

has spectrum $\{\lambda_1, \dots, \lambda_n\}$ and corresponding eigenspaces $\text{Eig}_{f_{\mathbf{D}}}(\lambda_i) = \text{span}\{\vec{e}_i\}$.

- Consider the $\mathbb{R}$ -vector space $P(\mathbb{R})$ of polynomials and $f = \frac{d}{dx} : P(\mathbb{R}) \rightarrow P(\mathbb{R})$ the derivative by the variable x . The kernel of f consists of the constant polynomials and hence 0 is an eigenvalue for f . For any non-zero scalar λ we

cannot have polynomials p satisfying $\frac{d}{dx}p = \lambda p$, as the left hand of this last expression has a smaller degree than the right hand side.

Previously we defined the trace and determinant for an endomorphism $g : V \rightarrow V$ by observing that the trace and determinant of the matrix representation of g are independent of the chosen basis of V . Similarly, we can consider eigenvalues of g and eigenvalues of the matrix representation of g with respect to some ordered basis of V . Perhaps unsurprisingly, the eigenvalues are the same:

Proposition 6.27 *Let $g : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{K}$ -vector space V . Let $\mathbf{b}$ be an ordered basis of V with corresponding linear coordinate system β . Then $v \in V$ is an eigenvector of g with eigenvalue $\lambda \in \mathbb{K}$ if and only if $\beta(v) \in \mathbb{K}^n$ is an eigenvector with eigenvalue λ of $\mathbf{M}(g, \mathbf{b}, \mathbf{b})$. In particular, conjugate matrices have the same eigenvalues.*

Proof Write $\mathbf{A} = \mathbf{M}(g, \mathbf{b}, \mathbf{b})$. Recall that by an eigenvector of $\mathbf{A} \in M_{n,n}(\mathbb{K})$, we mean an eigenvector of $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$. By Definition 3.92, we have $f_{\mathbf{A}} = \beta \circ g \circ \beta^{-1}$. Suppose $\lambda \in \mathbb{K}$ is an eigenvalue of g so that $g(v) = \lambda v$ for some non-zero vector $v \in V$. Consider the vector $\vec{x} = \beta(v) \in \mathbb{K}^n$ which is non-zero, since $\beta : V \rightarrow \mathbb{K}^n$ is an isomorphism. Then

$$f_{\mathbf{A}}(\vec{x}) = \beta(g(\beta^{-1}(\vec{x}))) = \beta(g(v)) = \beta(\lambda v) = \lambda \beta(v) = \lambda \vec{x},$$

so that $\vec{x}$ is an eigenvector of $f_{\mathbf{A}}$ with eigenvalue λ .

Conversely, if λ is an eigenvalue of $f_{\mathbf{A}}$ with non-zero eigenvector $\vec{x}$, then it follows as above that $v = \beta^{-1}(\vec{x}) \in V$ is an eigenvector of g with eigenvalue λ .

By Remark 6.15, if the matrices $\mathbf{A}, \mathbf{B}$ are similar, then they represent the same endomorphism $g : \mathbb{K}^n \rightarrow \mathbb{K}^n$ and hence have the same eigenvalues. $\square$

The “nicest” endomorphisms are those for which there exists an ordered basis consisting of eigenvectors:

Definition 6.28 (Diagonalisable endomorphism)

- An endomorphism $g : V \rightarrow V$ is called *diagonalisable* if there exists an ordered basis $\mathbf{b}$ of V such that each element of $\mathbf{b}$ is an eigenvector of g .
- For $n \in \mathbb{N}$, a matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ is called diagonalisable over $\mathbb{K}$ if the endomorphism $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$ is diagonalisable.

Example 6.29

- (i) We consider $V = P(\mathbb{R})$ and the endomorphism $g : V \rightarrow V$ which replaces the variable x with $2x$. For instance, we have

$$g(x^2 - 2x + 3) = (2x)^2 - 2(2x) + 3 = 4x^2 - 4x + 3.$$

Then g is diagonalisable. The vector space $P(\mathbb{R})$ has an ordered basis $\mathbf{b} = (1, x, x^2, x^3, \dots)$. Clearly, for all $k \in \mathbb{N} \cup \{0\}$ we have $g(x^k) = 2^k x^k$, so that x^k is an eigenvector of g with eigenvalue 2^k .

(ii) For $\alpha \in (0, \pi)$ consider

$$\mathbf{R}_\alpha = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix}.$$

Recall that the endomorphism $f_{\mathbf{R}_\alpha} : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ rotates vectors counter-clockwise around the origin $0_{\mathbb{R}^2}$ by the angle α . Since $\alpha \in (0, \pi)$, the endomorphism $f_{\mathbf{R}_\alpha}$ has no eigenvectors and hence is not diagonalisable.

Remark 6.30 Applying [Proposition 6.27](#), we conclude that in the case of a finite dimensional $\mathbb{K}$ -vector space V , an endomorphism $g : V \rightarrow V$ is diagonalisable if and only if there exists an ordered basis $\mathbf{b}$ of V such that $\mathbf{M}(g, \mathbf{b}, \mathbf{b})$ is a diagonal matrix. Moreover, $\mathbf{A} \in M_{n,n}(\mathbb{K})$ is diagonalisable if and only if $\mathbf{A}$ is similar over $\mathbb{K}$ to a diagonal matrix.

Recall, if $\mathcal{X}, \mathcal{Y}$ are sets, $f : \mathcal{X} \rightarrow \mathcal{Y}$ a mapping and $\mathcal{Z} \subset \mathcal{X}$ a subset of $\mathcal{X}$, we can consider the *restriction of f to $\mathcal{Z}$* , usually denoted by $f|_{\mathcal{Z}}$, which is the mapping

$$f|_{\mathcal{Z}} : \mathcal{Z} \rightarrow \mathcal{Y}, \quad z \mapsto f(z).$$

So we simply take the same mapping f , but apply it to the elements of the subset only.

Closely related to the notion of an eigenvector is that of a stable subspace. Let $v \in V$ be an eigenvector with eigenvalue λ of the endomorphism $g : V \rightarrow V$. The 1-dimensional subspace $U = \text{span}\{v\}$ is stable under g , that is, $g(U) \subset U$. Indeed, since $g(v) = \lambda v$ and since every vector $u \in U$ can be written as $u = tv$ for some scalar $t \in \mathbb{K}$, we have $g(u) = g(tv) = tg(v) = t\lambda v \in U$. This motivates the following definition:

Definition 6.31 (Stable subspace) A subspace $U \subset V$ is called *stable* or *invariant* under the endomorphism $g : V \rightarrow V$ if $g(U) \subset U$, that is $g(u) \in U$ for all vectors $u \in U$. In this case, the restriction $g|_U$ of g to U is an endomorphism of U .

Remark 6.32 Notice that a finite dimensional subspace $U \subset V$ is stable under g if and only if $g(v_i) \in U$ for $1 \leq i \leq m$, where $\{v_1, \dots, v_m\}$ is a basis of U .

Example 6.33

- (i) Every eigenspace of an endomorphism $g : V \rightarrow V$ is a stable subspace. By definition $g|_{\text{Eig}_g(\lambda)} : \text{Eig}_g(\lambda) \rightarrow \text{Eig}_g(\lambda)$ is multiplication by the scalar $\lambda \in \mathbb{K}$.
(ii) We consider $V = \mathbb{R}^3$ and

$$\mathbf{R}_\alpha = \begin{pmatrix} \cos \alpha & -\sin \alpha & 0 \\ \sin \alpha & \cos \alpha & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

for $\alpha \in (0, \pi)$. The endomorphism $f_{\mathbf{R}_\alpha} : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ is the rotation by the angle $\alpha \in \mathbb{R}$ around the axis spanned by $\vec{e}_3$. Then the plane $U = \{\vec{x} = (x_i)_{1 \leq i \leq 3} \in \mathbb{R}^3 \mid x_3 = 0\}$ is stable under $f = f_{\mathbf{R}_\alpha}$. Here $f|_U : U \rightarrow U$ is the rotation in the plane U around the origin with angle α .

Moreover, the vector $\vec{e}_3$ is an eigenvector with eigenvalue 1 so that

$$\text{Eig}_f(1) = \text{span}\{\vec{e}_3\}.$$

- (iii) We consider again the $\mathbb{R}$ -vector space $P(\mathbb{R})$ of polynomials and $f = \frac{d}{dx} : P(\mathbb{R}) \rightarrow P(\mathbb{R})$ the derivative by the variable x . For $n \in \mathbb{N}$ let U_n denote the subspace of polynomials of degree at most n . Since $U_{n-1} \subset U_n$, the subspace U_n is stable under f .

Stable subspaces correspond to zero blocks in the matrix representation of linear maps. More precisely:

Proposition 6.34 *Let V be a $\mathbb{K}$ -vector space of dimension $n \in \mathbb{N}$ and $g : V \rightarrow V$ an endomorphism. Furthermore, let $U \subset V$ be a subspace of dimension $1 \leq m \leq n$ and $\mathbf{b}$ an ordered basis of U and $\mathbf{c} = (\mathbf{b}, \mathbf{b}')$ an ordered basis of V . Then U is stable under g if and only if the matrix $\mathbf{A} = \mathbf{M}(g, \mathbf{c}, \mathbf{c})$ has the form*

$$\mathbf{A} = \begin{pmatrix} \hat{\mathbf{A}} & * \\ \mathbf{0}_{n-m,m} & * \end{pmatrix}$$

for some matrix $\hat{\mathbf{A}} \in M_{m,m}(\mathbb{K})$. In the case where U is stable under g , we have $\hat{\mathbf{A}} = \mathbf{M}(g|_U, \mathbf{b}, \mathbf{b}) \in M_{m,m}(\mathbb{K})$.

Proof Write $\mathbf{b} = (v_1, \dots, v_m)$ for vectors $v_i \in U$ and $\mathbf{b}' = (w_1, \dots, w_{n-m})$ for vectors $w_i \in V$.

$\Rightarrow$ Since U is stable under g , we have $g(u) \in U$ for all vectors $u \in U$. Since $\mathbf{b}$ is a basis of U , there exist scalars $\hat{A}_{ij} \in \mathbb{K}$ with $1 \leq i, j \leq m$ such that

$$g(v_j) = \sum_{i=1}^m \hat{A}_{ij} v_i$$

for all $1 \leq j \leq m$. By [Proposition 3.93](#), the matrix representation of g with respect to the ordered basis $\mathbf{c} = (\mathbf{b}, \mathbf{b}')$ of V thus takes the form

$$\mathbf{A} = \begin{pmatrix} \hat{\mathbf{A}} & * \\ \mathbf{0}_{n-m,m} & * \end{pmatrix}$$

where we write $\hat{\mathbf{A}} = (\hat{A}_{ij})_{1 \leq i,j \leq m} = \mathbf{M}(g|_U, \mathbf{b}, \mathbf{b})$.

$\Leftarrow$ Suppose

$$\mathbf{A} = \begin{pmatrix} \hat{\mathbf{A}} & * \\ \mathbf{0}_{n-m,m} & * \end{pmatrix} = \mathbf{M}(g, \mathbf{c}, \mathbf{c})$$

is the matrix representation of g with respect to the ordered basis $\mathbf{c}$ of V . Write $\hat{\mathbf{A}} = (\hat{A}_{ij})_{1 \leq i,j \leq m}$. Then, by [Proposition 3.93](#), $g(v_j) = \sum_{i=1}^m \hat{A}_{ij} v_i \in U$ for all $1 \leq j \leq m$, hence U is stable under g , by [Remark 6.32](#). $\square$

From [Proposition 6.34](#) we can conclude:

Remark 6.35 Suppose V is the direct sum of subspaces $U_1, U_2, \dots, U_m$, all of which are stable under the endomorphism $g : V \rightarrow V$. If $\mathbf{b}_i$ is an ordered basis of U_i for $i = 1, \dots, m$. Then the matrix representation of g with respect to the ordered basis

$\mathbf{c} = (\mathbf{b}_1, \dots, \mathbf{b}_m)$ takes the block form

$$\mathbf{A} = \begin{pmatrix} \mathbf{A}_1 & & & \\ & \mathbf{A}_2 & & \\ & & \ddots & \\ & & & \mathbf{A}_m \end{pmatrix}$$

where $\mathbf{A}_i = \mathbf{M}(g|_{U_i}, \mathbf{b}_i, \mathbf{b}_i)$ for $i = 1, \dots, m$.

6.4 The characteristic polynomial

The eigenvalues of an endomorphism are the solutions of a polynomial equation:

Lemma 6.36 *Let V be a finite dimensional $\mathbb{K}$ -vector space and $g : V \rightarrow V$ an endomorphism. Then $\lambda \in \mathbb{K}$ is an eigenvalue of g if and only if*

$$\det(\lambda \text{Id}_V - g) = 0.$$

Moreover if λ is an eigenvalue of g , then $\text{Eig}_g(\lambda) = \text{Ker}(\lambda \text{Id}_V - g)$.

Proof Let $v \in V$. We may write $v = \text{Id}_V(v)$. Hence

$$g(v) = \lambda v \iff 0_V = (\lambda \text{Id}_V - g)(v) \iff v \in \text{Ker}(\lambda \text{Id}_V - g)$$

It follows that $\text{Eig}_g(\lambda) = \text{Ker}(\lambda \text{Id}_V - g)$. Moreover $\lambda \in \mathbb{K}$ is an eigenvalue of g if and only if the kernel of $\lambda \text{Id}_V - g$ is different from $\{0_V\}$ or if and only if $\lambda \text{Id}_V - g$ is not injective. [Proposition 6.22](#) implies that $\lambda \in \mathbb{K}$ is an eigenvalue of g if and only if $\det(\lambda \text{Id}_V - g) = 0$. $\square$

Definition 6.37 (Characteristic polynomial — Video) Let $g : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{K}$ -vector space V . The function

$$\text{char}_g : \mathbb{K} \rightarrow \mathbb{K}, \quad x \mapsto \det(x \text{Id}_V - g)$$

is called the *characteristic polynomial of the endomorphism g* .

In practice, in order to compute the characteristic polynomial of an endomorphism $g : V \rightarrow V$, we choose an ordered basis $\mathbf{b}$ of V and compute the matrix representation $\mathbf{A} = \mathbf{M}(g, \mathbf{b}, \mathbf{b})$ of g with respect to $\mathbf{b}$. We then have

$$\text{char}_g(x) = \det(x \mathbf{1}_n - \mathbf{A}).$$

By the characteristic polynomial of a matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$, we mean the characteristic polynomial of the endomorphism $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$, that is, the function $x \mapsto \det(x \mathbf{1}_n - \mathbf{A})$.

A zero of a polynomial $f : \mathbb{K} \rightarrow \mathbb{K}$ is a scalar $\lambda \in \mathbb{K}$ such that $f(\lambda) = 0$. The *multiplicity of a zero λ* is the largest integer $n \geq 1$ such that there exists a polynomial $\hat{f} : \mathbb{K} \rightarrow \mathbb{K}$ so that $f(x) = (x - \lambda)^n \hat{f}(x)$ for all $x \in \mathbb{K}$. Zeros are also known as *roots*.

Example 6.38 The polynomial $f(x) = x^3 - x^2 - 8x + 12$ can be factorised as $f(x) = (x - 2)^2(x + 3)$ and hence has zero 2 with multiplicity 2 and -3 with multiplicity 1.

Definition 6.39 (Algebraic multiplicity) Let λ be an eigenvalue of the endomorphism $g : V \rightarrow V$. The multiplicity of the zero λ of char_g is called the *algebraic multiplicity* of λ .

Example 6.40

(i) We consider

$$\mathbf{A} = \begin{pmatrix} 1 & 5 \\ 5 & 1 \end{pmatrix}.$$

Then

$$\begin{aligned} \text{char}_{\mathbf{A}}(x) &= \text{char}_{f_{\mathbf{A}}}(x) = \det(x\mathbf{1}_2 - \mathbf{A}) = \det \begin{pmatrix} x-1 & -5 \\ -5 & x-1 \end{pmatrix} \\ &= (x-1)^2 - 25 = x^2 - 2x - 24 = (x+4)(x-6). \end{aligned}$$

Hence we have eigenvalues $\lambda_1 = 6$ and $\lambda_2 = -4$, both with algebraic multiplicity 1. By definition we have

$$\text{Eig}_{\mathbf{A}}(6) = \text{Eig}_{f_{\mathbf{A}}}(6) = \{ \vec{v} \in \mathbb{K}^2 \mid \mathbf{A}\vec{v} = 6\vec{v} \}$$

and we compute that

$$\text{Eig}_{\mathbf{A}}(6) = \text{span} \left\{ \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\}$$

Since $\dim \text{Eig}_{\mathbf{A}}(6) = 1$, the eigenvalue 6 has geometric multiplicity 1. Likewise we compute

$$\text{Eig}_{\mathbf{A}}(-4) = \text{span} \left\{ \begin{pmatrix} -1 \\ 1 \end{pmatrix} \right\}$$

so that the eigenvalue -4 has geometric multiplicity 1 as well. Notice that we have an ordered basis of eigenvectors of $\mathbf{A}$ and hence $\mathbf{A}$ is diagonalisable.

(ii) We consider

$$\mathbf{A} = \begin{pmatrix} 2 & 1 \\ 0 & 2 \end{pmatrix}$$

Then $\text{char}_{\mathbf{A}}(x) = (x-2)^2$ so that we have a single eigenvalue 2 with algebraic multiplicity 2. We compute

$$\text{Eig}_{\mathbf{A}}(2) = \text{span} \left\{ \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right\}$$

so that the eigenvalue 2 has geometric multiplicity 1. Notice that we cannot find an ordered basis consisting of eigenvectors, hence $\mathbf{A}$ is not diagonalisable.

The determinant and trace of an endomorphism do appear as coefficients in its characteristic polynomial:

Lemma 6.41 Let $g : V \rightarrow V$ be an endomorphism of a $\mathbb{K}$ -vector space V of dimension n . Then char_g is a polynomial of degree n and

$$\text{char}_g(x) = x^n - \text{Tr}(g)x^{n-1} + \cdots + (-1)^n \det(g).$$

Proof We fix an ordered basis $\mathbf{b}$ of V . Writing $\mathbf{M}(g, \mathbf{b}, \mathbf{b}) = \mathbf{A} = (A_{ij})_{1 \leq i, j \leq n}$ and using the Leibniz formula (5.8), we have

$$\text{char}_g(x) = \sum_{\sigma \in S_n} \text{sgn}(\sigma) \prod_{i=1}^n B_{i\sigma(i)},$$

where

$$B_{ij} = \begin{cases} x - A_{ii}, & i = j, \\ -A_{ij}, & i \neq j. \end{cases}$$

Therefore, char_g is a finite sum of products containing x at most n times, hence char_g is a polynomial in x of degree at most n . The identity permutation contributes the term $\prod_{i=1}^n B_{ii}$ in the Leibniz formula, hence we obtain

$$\text{char}_g(x) = \prod_{i=1}^n (x - A_{ii}) + \sum_{\sigma \in S_n, \sigma \neq 1} \text{sgn}(\sigma) \prod_{i=1}^n B_{i\sigma(i)}$$

We now use induction to show that

$$\prod_{i=1}^n (x - A_{ii}) = x^n - \text{Tr}(\mathbf{A})x^{n-1} + C_{n-2}x^{n-2} + \cdots + c_1x + c_0$$

for scalars $C_{n-2}, \dots, c_0 \in \mathbb{K}$. For $n = 1$ we obtain $x - A_{11}$, so that the statement is anchored.

Inductive step: Suppose

$$\prod_{i=1}^{n-1} (x - A_{ii}) = x^{n-1} - \left(\sum_{i=1}^{n-1} A_{ii} \right) x^{n-2} + C_{n-2}x^{n-3} + \cdots + c_1x + c_0,$$

for coefficients $C_{n-2}, \dots, c_0$, then

$$\begin{aligned} \prod_{i=1}^n (x - A_{ii}) &= (x - A_{nn}) \left[x^{n-1} - \left(\sum_{i=1}^{n-1} A_{ii} \right) x^{n-2} + C_{n-2}x^{n-3} + \cdots + c_1x + c_0 \right] \\ &= x^n - \left(\sum_{i=1}^n A_{ii} \right) x^{n-1} + \text{lower order terms in } x, \end{aligned}$$

so the induction is complete.

We next argue that $\sum_{\sigma \in S_n, \sigma \neq 1} \text{sgn}(\sigma) \prod_{i=1}^n B_{i\sigma(i)}$ has at most degree $n - 2$. Notice that each factor $B_{i\sigma(i)}$ of $\prod_{i=1}^n B_{i\sigma(i)}$ for which $i \neq \sigma(i)$ does not contain x . So suppose that $\sum_{\sigma \in S_n, \sigma \neq 1} \text{sgn}(\sigma) \prod_{i=1}^n B_{i\sigma(i)}$ has degree bigger or equal than $n - 1$. Then we have $n - 1$ integers i with $1 \leq i \leq n$ such that $i = \sigma(i)$. Let j denote the remaining integer. Since σ is injective, it follows that for any $i \neq j$ we must have $i = \sigma(i) \neq \sigma(j)$. Therefore, $\sigma(j) = j$ and hence $\sigma = 1$, a contradiction.

In summary, we have shown that

$$\text{char}_g(x) = x^n - \text{Tr}(g)x^{n-1} + C_{n-2}x^{n-2} + \cdots + c_1x + c_0$$

for coefficients $C_{n-2}, \dots, c_0 \in \mathbb{K}$. It remains to show that $c_0 = (-1)^n \det(g)$. We have $c_0 = \text{char}_g(0) = \det(-g) = \det(-\mathbf{A})$. Since the determinant is linear in each row of $\mathbf{A}$, this gives $\det(-\mathbf{A}) = (-1)^n \det(\mathbf{A})$, as claimed. $\square$

Remark 6.42 In particular, for $n = 2$ we have $\text{char}_g(x) = x^2 - \text{Tr}(g)x + \det(g)$. Compare with [Example 6.40](#).

6.5 Properties of eigenvalues

We will argue next that an endomorphism $g : V \rightarrow V$ of a finite dimensional $\mathbb{K}$ -vector space V has at most $\dim(V)$ eigenvalues. We first need:

Theorem 6.43 (Little Bézout's theorem) *For a polynomial $f \in P(\mathbb{K})$ of degree $n \geq 1$ and $x_0 \in \mathbb{K}$, there exists a polynomial $g \in P(\mathbb{K})$ of degree $n - 1$ such that for all $x \in \mathbb{K}$ we have $f(x) = f(x_0) + g(x)(x - x_0)$.*

Proof We will give an explicit expression for the polynomial g . If one is not interested in such an expression, a proof using induction can also be given. Write $f(x) = \sum_{k=0}^n a_k x^k$ for coefficients $(a_0, \dots, a_n) \in \mathbb{K}^{n+1}$. For $0 \leq j \leq n - 1$ consider

$$(6.3) \quad b_j = \sum_{k=0}^{n-j-1} a_{k+j+1} x_0^k$$

and the polynomial

$$g(x) = \sum_{j=0}^{n-1} b_j x^j$$

of degree $n - 1$. We have

$$\begin{aligned} g(x)(x - x_0) &= \sum_{j=0}^{n-1} \sum_{k=0}^{n-j-1} (a_{k+j+1} x_0^k x^{j+1}) - \sum_{j=0}^{n-1} \sum_{k=0}^{n-j-1} (a_{k+j+1} x_0^{k+1} x^j) \\ &= \sum_{j=1}^n \sum_{k=0}^{n-j} (a_{k+j} x_0^k x^j) - \sum_{j=0}^{n-1} \sum_{k=1}^{n-j} (a_{k+j} x_0^k x^j) \\ &= a_n x^n + \sum_{j=1}^{n-1} a_j x^j + a_0 - a_0 - \sum_{k=1}^n a_k x_0^k = f(x) - f(x_0). \end{aligned}$$

□

From this we conclude:

Proposition 6.44 *Let $f \in P(\mathbb{K})$ be a polynomial of degree n . Then f has at most n (distinct) zeros or f is the zero polynomial.*

Proof We use induction. The case $n = 0$ is clear, hence the statement is anchored.

Inductive step: Suppose $f \in P(\mathbb{K})$ is a polynomial of degree n with $n + 1$ distinct zeros $\lambda_1, \dots, \lambda_{n+1}$. Since $f(\lambda_{n+1}) = 0$, [Theorem 6.43](#) implies that

$$f(x) = (x - \lambda_{n+1})g(x)$$

for some polynomial g of degree $n - 1$. For $1 \leq i \leq n$, we thus have

$$0 = f(\lambda_i) = (\lambda_i - \lambda_{n+1})g(\lambda_i).$$

Since $\lambda_i \neq \lambda_{n+1}$ it follows that $g(\lambda_i) = 0$. Therefore, g has n distinct zeros and must be the zero polynomial by the induction hypothesis. It follows that f is the zero polynomial as well. □

This gives:

Corollary 6.45 Let $g : V \rightarrow V$ be an endomorphism of a $\mathbb{K}$ -vector space of dimension $n \in \mathbb{N}$. Then g has at most n (distinct) eigenvalues.

Proof By Lemma 6.36 and Lemma 6.41, the eigenvalues of g are the zeros of the characteristic polynomial. The characteristic polynomial of g has degree n . The claim follows by applying Proposition 6.44. $\square$

Proposition 6.46 (Linear independence of eigenvectors) Let V be a finite dimensional $\mathbb{K}$ -vector space and $g : V \rightarrow V$ an endomorphism. Then the eigenspaces $\text{Eig}_g(\lambda)$ of g are in direct sum. In particular, if $v_1, \dots, v_m$ are eigenvectors corresponding to distinct eigenvalues of g , then $\{v_1, \dots, v_m\}$ are linearly independent.

Proof We use induction on the number m of distinct eigenvalues of g . Let $\{\lambda_1, \dots, \lambda_m\}$ be distinct eigenvalues of g . For $m = 1$ the statement is trivially true, so the statement is anchored.

Inductive step: Assume $m - 1$ eigenspaces are in direct sum. We want to show that then m eigenspaces are also in direct sum. Let $v_i, v'_i \in \text{Eig}_g(\lambda_i)$ be eigenvectors such that

$$(6.4) \quad v_1 + v_2 + \dots + v_m = v'_1 + v'_2 + \dots + v'_m.$$

Applying g to this last equation gives

$$(6.5) \quad \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_m v_m = \lambda_1 v'_1 + \lambda_2 v'_2 + \dots + \lambda_m v'_m.$$

Subtracting λ_m times (6.4) from (6.5) gives

$$(\lambda_1 - \lambda_m)v_1 + \dots + (\lambda_{m-1} - \lambda_m)v_{m-1} = (\lambda_1 - \lambda_m)v'_1 + \dots + (\lambda_{m-1} - \lambda_m)v'_{m-1}.$$

Since $m - 1$ eigenspaces are in direct sum, this implies that $(\lambda_i - \lambda_m)v_i = (\lambda_i - \lambda_m)v'_i$ for $1 \leq i \leq m - 1$. Since the eigenvalues are distinct, we have $\lambda_i - \lambda_m \neq 0$ for all $1 \leq i \leq m - 1$ and hence $v_i = v'_i$ for all $1 \leq i \leq m - 1$. Now (6.5) implies that $v_m = v'_m$ as well and the inductive step is complete.

Since the eigenspaces are in direct sum, the linear independence of eigenvectors with respect to distinct eigenvalues follows from Remark 6.7. $\square$

In the case where all the eigenvalues are distinct, we conclude that g is diagonalisable.

Proposition 6.47 Let $g : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{K}$ -vector space V . Suppose the characteristic polynomial of g has $\dim(V)$ distinct zeros (that is, the algebraic multiplicity of each eigenvalue is 1), then g is diagonalisable.

Proof Let $n = \dim(V)$. Let $\lambda_1, \dots, \lambda_n$ denote the distinct eigenvalues of g . Let $0_V \neq v_i \in \text{Eig}_g(\lambda_i)$ for $i = 1, \dots, n$. Then, by Proposition 6.46, the eigenvectors are linearly independent, it follows that $(v_1, \dots, v_n)$ is an ordered basis of V consisting of eigenvectors, hence g is diagonalisable. $\square$

Remark 6.48 Proposition 6.47 gives a sufficient condition for an endomorphism $g : V \rightarrow V$ to be diagonalisable, it is however not necessary. The identity endomorphism is diagonalisable, but its spectrum consists of the single eigenvalue 1 with algebraic multiplicity $\dim(V)$.

Every polynomial in $P(\mathbb{C})$ of degree at least 1 has at least one zero. This fact is known as the *fundamental theorem of algebra*. The name is well-established, but quite misleading, as there is no purely algebraic proof. You will encounter a proof of this statement in the module M07. As a consequence we obtain the following important existence theorem:

Theorem 6.49 (Existence of eigenvalues) *Let $g : V \rightarrow V$ be an endomorphism of a complex vector space V of dimension $n \geq 1$. Then g admits at least one eigenvalue. Moreover, the sum of the algebraic multiplicities of the eigenvalues of g is equal to n . In particular, if $\mathbf{A} \in M_{n,n}(\mathbb{C})$ is a matrix, then there is at least one eigenvalue of $\mathbf{A}$.*

Proof By Lemma 6.36 and Lemma 6.41, the eigenvalues of g are the zeros of the characteristic polynomial and this is an element of $P(\mathbb{C})$. The first statement thus follows by applying the fundamental theorem of algebra to the characteristic polynomial of g .

Applying Theorem 6.43 and the fundamental theorem of algebra repeatedly, we find $k \in \mathbb{N}$ and multiplicities $m_1, \dots, m_k \in \mathbb{N}$ such that

$$\text{char}_g(x) = (x - \lambda_1)^{m_1} (x - \lambda_2)^{m_2} \cdots (x - \lambda_k)^{m_k}$$

where $\lambda_1, \dots, \lambda_k$ are zeros of char_g . Since char_g has degree n , it follows that $\sum_{i=1}^k m_i = n$. □

Example 6.50

- Recall that the *discriminant* of a quadratic polynomial $x \mapsto ax^2 + bx + c \in P(\mathbb{K})$ is $b^2 - 4ac$, provided $a \neq 0$. If $\mathbb{K} = \mathbb{C}$ and $b^2 - 4ac$ is non-zero, then the polynomial $ax^2 + bx + c$ has two distinct zeros. The characteristic polynomial of a 2-by-2 matrix $\mathbf{A}$ satisfies $\text{char}_{\mathbf{A}}(x) = x^2 - \text{Tr}(\mathbf{A})x + \det(\mathbf{A})$. Therefore, if $\mathbf{A}$ has complex entries and satisfies $(\text{Tr } \mathbf{A})^2 - 4 \det \mathbf{A} \neq 0$, then it is diagonalisable. If $\mathbf{A}$ has real entries and satisfies $(\text{Tr } \mathbf{A})^2 - 4 \det \mathbf{A} \geq 0$, then it has a least one eigenvalue. If $(\text{Tr } \mathbf{A})^2 - 4 \det \mathbf{A} > 0$ then it is diagonalisable.
- Recall that, by Proposition 5.24, an upper triangular matrix $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n}$ satisfies $\det \mathbf{A} = \prod_{i=1}^n A_{ii}$. It follows that

$$\text{char}_{\mathbf{A}}(x) = \prod_{i=1}^n (x - A_{ii}) = (x - A_{11})(x - A_{22}) \cdots (x - A_{nn}).$$

Consequently, an upper triangular matrix has spectrum $\{A_{11}, A_{22}, \dots, A_{nn}\}$ and is diagonalisable if all its diagonal entries are distinct. Notice that by Example 6.40 (ii) not every upper triangular matrix is diagonalisable.

Example 6.51 (Fibonacci sequences) We revisit the Fibonacci sequences, now equipped with the theory of endomorphisms. A Fibonacci sequence is a sequence $\xi : \mathbb{N} \cup \{0\} \rightarrow \mathbb{K}$ satisfying the recursive relation $\xi_{n+2} = \xi_n + \xi_{n+1}$. Consider the matrix

$$\mathbf{A} = \begin{pmatrix} \xi_0 & \xi_1 \\ \xi_1 & \xi_2 \end{pmatrix}.$$

Then, using induction, we can show that

$$\mathbf{A}^n = \begin{pmatrix} \xi_{n-1} & \xi_n \\ \xi_n & \xi_{n+1} \end{pmatrix}$$

for all $n \in \mathbb{N}$. We would like to compute $\mathbf{A}^n$ for the initial conditions $\xi_0 = 0$ and $\xi_1 = 1$. Suppose we can find an invertible matrix $\mathbf{C}$ so that $\mathbf{A} = \mathbf{C} \mathbf{D} \mathbf{C}^{-1}$ for some

diagonal matrix $\mathbf{D}$. Then

$$\mathbf{A}^n = \mathbf{C}\mathbf{D}\mathbf{C}^{-1}\mathbf{C}\mathbf{D}\mathbf{C}^{-1}\dots\mathbf{C}\mathbf{D}\mathbf{C}^{-1} = \mathbf{C}\mathbf{D}^n\mathbf{C}^{-1}$$

and we can easily compute $\mathbf{A}^n$, as the n -th power of a diagonal matrix $\mathbf{D}$ is the diagonal matrix whose diagonal entries are given by the n -th powers of diagonal entries of $\mathbf{D}$. We thus want to diagonalise the matrix

$$\mathbf{A} = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix}.$$

We obtain $\text{char}_{\mathbf{A}}(x) = x^2 - x - 1$ and hence eigenvalues $\lambda_1 = (1 + \sqrt{5})/2$ and $\lambda_2 = (1 - \sqrt{5})/2$. From this we compute

$$\text{Eig}_{\mathbf{A}}(\lambda_1) = \text{span} \left\{ \begin{pmatrix} 1 \\ \lambda_1 \end{pmatrix} \right\} \quad \text{and} \quad \text{Eig}_{\mathbf{A}}(\lambda_2) = \text{span} \left\{ \begin{pmatrix} 1 \\ \lambda_2 \end{pmatrix} \right\}$$

Let $\mathbf{e} = (\vec{e}_1, \vec{e}_2)$ denote the standard basis of $\mathbb{R}^2$ and consider the ordered basis

$$\mathbf{b} = \left(\begin{pmatrix} 1 \\ \lambda_1 \end{pmatrix}, \begin{pmatrix} 1 \\ \lambda_2 \end{pmatrix} \right)$$

of eigenvectors of $f_{\mathbf{A}}$. We have

$$\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} = \mathbf{D}$$

and the change of base matrix is

$$\mathbf{C} = \mathbf{C}(\mathbf{b}, \mathbf{e}) = \begin{pmatrix} 1 & 1 \\ \lambda_1 & \lambda_2 \end{pmatrix}$$

and

$$\mathbf{C}^{-1} = \mathbf{C}(\mathbf{e}, \mathbf{b}) = \frac{1}{\lambda_2 - \lambda_1} \begin{pmatrix} \lambda_2 & -1 \\ -\lambda_1 & 1 \end{pmatrix}.$$

Therefore $\mathbf{A} = \mathbf{C}\mathbf{D}\mathbf{C}^{-1}$ and hence $\mathbf{A}^n = \mathbf{C}\mathbf{D}^n\mathbf{C}^{-1}$ so that

$$\mathbf{A}^n = \frac{1}{\lambda_2 - \lambda_1} \begin{pmatrix} 1 & 1 \\ \lambda_1 & \lambda_2 \end{pmatrix} \begin{pmatrix} \lambda_1^n & 0 \\ 0 & \lambda_2^n \end{pmatrix} \begin{pmatrix} \lambda_2 & -1 \\ -\lambda_1 & 1 \end{pmatrix} = \begin{pmatrix} \xi_{n-1} & \xi_n \\ \xi_n & \xi_{n+1} \end{pmatrix}.$$

This yields the formula

$$\xi_n = \frac{\lambda_1^n - \lambda_2^n}{\lambda_1 - \lambda_2}.$$

Proposition 6.52 Let $g : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{K}$ -vector space V of dimension $n \geq 1$.

- (i) Let λ be an eigenvalue of g . Then its algebraic multiplicity is at least as big as its geometric multiplicity.
- (ii) If $\mathbb{K} = \mathbb{C}$, then g is diagonalisable if and only if for all eigenvalues of g , the algebraic and geometric multiplicity are the same.

Proof (i) Let $\dim \text{Eig}_g(\lambda) = m$ and $\mathbf{b}$ be an ordered basis of $\text{Eig}_g(\lambda)$. Furthermore, let $\mathbf{b}'$ be an ordered tuple of vectors such that $\mathbf{c} = (\mathbf{b}, \mathbf{b}')$ is an ordered basis of V . The eigenspace $\text{Eig}_g(\lambda)$ is stable under g and

$$\mathbf{M}(g|_{\text{Eig}_g(\lambda)}, \mathbf{b}, \mathbf{b}) = \lambda \mathbf{1}_m.$$

By Proposition 6.34, the matrix representation of g with respect to the basis $\mathbf{c}$ takes the form

$$\mathbf{M}(g, \mathbf{c}, \mathbf{c}) = \begin{pmatrix} \lambda \mathbf{1}_m & * \\ \mathbf{0}_{m-n, m} & \mathbf{B} \end{pmatrix}$$

for some matrix $\mathbf{B} \in M_{n-m, n-m}(\mathbb{K})$. We thus obtain

$$\text{char}_g(x) = \det \begin{pmatrix} (x - \lambda)\mathbf{1}_m & * \\ \mathbf{0}_{m-n, m} & x\mathbf{1}_{n-m} - \mathbf{B} \end{pmatrix}$$

Applying the Laplace expansion (5.5) with respect to the first column, we have

$$\text{char}_g(x) = (x - \lambda) \det \begin{pmatrix} (x - \lambda)\mathbf{1}_{m-1} & * \\ \mathbf{0}_{m-n, m-1} & x\mathbf{1}_{n-m} - \mathbf{B} \end{pmatrix}$$

Applying the Laplace expansion again with respect to the first column, m -times in total, we get

$$\text{char}_g(x) = (x - \lambda)^m \det(x\mathbf{1}_{n-m} - \mathbf{B}) = (x - \lambda)^m \text{char}_{\mathbf{B}}(x).$$

The algebraic multiplicity of λ is thus at least m .

(ii) Suppose $\mathbb{K} = \mathbb{C}$ and that $g : V \rightarrow V$ is diagonalisable. Hence we have an ordered basis $(v_1, \dots, v_n)$ of V consisting of eigenvectors of g . Therefore,

$$\text{char}_g(x) = \prod_{i=1}^n (x - \lambda_i)$$

where λ_i is the eigenvalue of the eigenvector v_i , $1 \leq i \leq n$. For any eigenvalue λ_j , its algebraic multiplicity is the number of indices i with $\lambda_i = \lambda_j$. For each such index i , the eigenvector v_i satisfies $g(v_i) = \lambda_i v_i = \lambda_j v_i$ and hence is an element of the eigenspace $\text{Eig}_g(\lambda_j)$. The geometric multiplicity of each eigenvalue is thus at least as big as the algebraic multiplicity, but by the previous statement, the latter cannot be bigger than the former, hence they are equal.

Conversely, suppose that for all eigenvalues of g , the algebraic and geometric multiplicity are the same. Since $\mathbb{K} = \mathbb{C}$, by Theorem 6.49, the sum of the algebraic multiplicities is n . The sum of the geometric multiplicities is by assumption also n . Since, by Proposition 6.46, the eigenspaces with respect to different eigenvalues are in direct sum, we obtain a basis of V consisting of eigenvectors of g . $\square$

6.6 Special endomorphisms

6.6.1 Involutions

A mapping $\iota : \mathcal{X} \rightarrow \mathcal{X}$ from a set $\mathcal{X}$ into itself is called an *involution*, if $\iota \circ \iota = \text{Id}_{\mathcal{X}}$. In the case where $\mathcal{X}$ is a vector space and ι is linear, then ι is called a *linear involution*.

Example 6.53 (Involutions)

- (i) Let V be a $\mathbb{K}$ -vector space. Then the identity mapping $\text{Id}_V : V \rightarrow V$ is a linear involution.
- (ii) For all $n \in \mathbb{N}$, the transpose $M_{n,n}(\mathbb{K}) \rightarrow M_{n,n}(\mathbb{K})$ is a linear involution.
- (iii) For $n \in \mathbb{N}$, let $\mathcal{X}$ denote the set of invertible $n \times n$ matrices. Then the matrix inverse $^{-1} : \mathcal{X} \rightarrow \mathcal{X}$ is an involution. Notice that $\mathcal{X}$ is not a vector space.
- (iv) For any $\mathbb{K}$ -vector space V , the mapping $\iota : V \rightarrow V, v \mapsto -v$ is a linear involution. Considering $F(I, \mathbb{K})$, the $\mathbb{K}$ -vector space of functions on the interval $I \subset \mathbb{R}$, we obtain a linear involution of $F(V, \mathbb{K})$ by sending a function f to $f \circ \iota$.
- (v) If $\mathbf{A} \in M_{n,n}(\mathbb{K})$ satisfies $\mathbf{A}^2 = \mathbf{1}_n$, then $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$ is a linear involution.

The spectrum of an involution is a subset of $\{-1, 1\}$.

Proposition 6.54 *Let V be a $\mathbb{K}$ -vector space and $\iota : V \rightarrow V$ a linear involution. Then the spectrum of ι is contained in $\{-1, 1\}$. Moreover $V = \text{Eig}_\iota(1) \oplus \text{Eig}_\iota(-1)$ and ι is diagonalisable.*

Proof Suppose $\lambda \in \mathbb{K}$ is an eigenvalue of ι so that $\iota(v) = \lambda v$ for some non-zero vector $v \in V$. Then $\iota(\iota(v)) = v = \lambda \iota(v) = \lambda^2 v$. Hence $(1 - \lambda^2)v = 0_V$ and since v is non-zero, we conclude that $\lambda = \pm 1$. By Proposition 6.46, the eigenspaces $\text{Eig}_\iota(1)$ and $\text{Eig}_\iota(-1)$ are in direct sum.

For $v \in V$ we write

$$v = \underbrace{\frac{1}{2}(v + \iota(v))}_{\in \text{Eig}_\iota(1)} + \underbrace{\frac{1}{2}(v - \iota(v))}_{\in \text{Eig}_\iota(-1)}$$

hence $V = \text{Eig}_\iota(1) \oplus \text{Eig}_\iota(-1)$. Take an ordered basis $\mathbf{b}_+$ of $\text{Eig}_\iota(1)$ and an ordered basis $\mathbf{b}_-$ of $\text{Eig}_\iota(-1)$. Then $(\mathbf{b}_+, \mathbf{b}_-)$ is an ordered basis of V consisting of eigenvectors of ι . $\square$

6.6.2 Projections

A linear mapping $\Pi : V \rightarrow V$ satisfying $\Pi \circ \Pi = \Pi$ is called a *projection*.

Example 6.55 Consider $V = \mathbb{R}^3$ and

$$\mathbf{A} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Clearly, $\mathbf{A}^2 = \mathbf{A}$ and $f_{\mathbf{A}} : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ projects a vector $\vec{x} = (x_i)_{1 \leq i \leq 3}$ onto the plane $\{\vec{x} \in \mathbb{R}^3 \mid x_3 = 0\}$.

In a sense there is only one type of projection. Recall from the exercises that for a projection $\Pi : V \rightarrow V$, we have $V = \text{Ker } \Pi \oplus \text{Im } \Pi$. Given two subspaces U_1, U_2 of V such that $V = U_1 \oplus U_2$, there is a projection $\Pi : V \rightarrow V$ whose kernel is U_1 and whose image is U_2 . Indeed, every vector $v \in V$ can be written as $v = u_1 + u_2$ for unique vectors $u_i \in U_i$ for $i = 1, 2$. Hence we obtain a projection by defining $\Pi(v) = u_2$ for all $v \in V$.

Denote by $\mathcal{X}$ the set of projections from V to V and by $\mathcal{Y}$ the set of pairs (U_1, U_2) of subspaces of V that are in direct sum and satisfy $V = U_1 \oplus U_2$. Then we obtain a mapping $\Lambda : \mathcal{X} \rightarrow \mathcal{Y}$ defined by $f \mapsto (\text{Ker } f, \text{Im } f)$.

Similar to Proposition 6.54, we obtain:

Proposition 6.56 *Let V be a $\mathbb{K}$ -vector space and $\Pi : V \rightarrow V$ a projection. Then the spectrum of Π is contained in $\{0, 1\}$. Moreover $V = \text{Eig}_\Pi(0) \oplus \text{Eig}_\Pi(1)$, Π is diagonalisable and $\text{Im } \Pi = \text{Eig}_\Pi(1)$.*

Proof Let $v \in V$ be an eigenvector of the projection Π with eigenvalue λ . Hence we obtain $\Pi(\Pi(v)) = \lambda^2 v = \Pi(v) = \lambda v$, equivalently, $\lambda(\lambda - 1)v = 0_V$. Since v is non zero, it follows that $\lambda = 0$ or $\lambda = 1$. Since Π is a projection, we have $V = \text{Ker } \Pi \oplus \text{Im } \Pi$. Since $\text{Ker } \Pi = \text{Eig}_\Pi(0)$, we thus only need to show that $\text{Im } \Pi = \text{Eig}_\Pi(1)$. Let $v \in \text{Im } \Pi$ so that $v = \Pi(\hat{v})$ for some vector $\hat{v} \in V$. Hence $\Pi(v) = \Pi(\Pi(\hat{v})) = \Pi(\hat{v}) = v$ and v is an

eigenvector with eigenvalue 1. Conversely, suppose $v \in V$ is an eigenvector of Π with eigenvalue 1. Then $\Pi(v) = v = \Pi(\Pi(v))$ and hence $v \in \text{Im } \Pi$. We thus conclude that $\text{Im } \Pi = \text{Eig}_\Pi(1)$. Choosing an ordered basis of $\text{Ker } \Pi$ and an ordered basis of $\text{Im } \Pi$ gives a basis of V consisting of eigenvectors, hence Π is diagonalisable. $\square$

Exercises

Exercise 6.57 Derive the formula (6.3) for the coefficients b_j .

Exercise 6.58 Show that Λ is a bijection.

Exercise 6.59 Show that if $\Pi : V \rightarrow V$ is a projection then $\text{Id}_V - \Pi : V \rightarrow V$ is a projection with kernel equal to the image of Π and image equal to the kernel of Π .

Quotient vector spaces

7.1 Affine mappings and affine spaces

WEEK 13

Previously we saw that we can take the sum of subspaces of a vector space. In this final chapter of the Linear Algebra I module we introduce the concept of a quotient of a vector space by a subspace.

Translations are among the simplest non-linear mappings.

Definition 7.1 (Translation) Let V be a $\mathbb{K}$ -vector space and $v_0 \in V$. The mapping

$$T_{v_0} : V \rightarrow V, \quad v \mapsto v + v_0$$

is called the *translation* by the vector v_0 .

Remark 7.2 Notice that for $v_0 \neq 0_V$, a translation is not linear, since $T_{v_0}(0_V) = 0_V + v_0 = v_0 \neq 0_V$.

Taking $s_1 = 1$ and $s_2 = -1$ in (3.6), we see that a linear map $f : V \rightarrow W$ between $\mathbb{K}$ -vector spaces V, W satisfies $f(v_1 - v_2) = f(v_1) - f(v_2)$ for all $v_1, v_2 \in V$. In particular, linear maps are affine maps in the following sense:

Definition 7.3 (Affine mapping) A mapping $f : V \rightarrow W$ is called *affine* if there exists a linear map $g : V \rightarrow W$ so that $f(v_1) - f(v_2) = g(v_1 - v_2)$ for all $v_1, v_2 \in V$. We call g the *linear map associated to f* .

Affine mappings are compositions of linear mappings and translations:

Proposition 7.4 A mapping $f : V \rightarrow W$ is affine if and only if there exists a linear map $g : V \rightarrow W$ and a translation $T_{w_0} : W \rightarrow W$ so that $f = T_{w_0} \circ g$.

Proof $\Leftarrow$ Let $g : V \rightarrow W$ be linear and $T_{w_0} : W \rightarrow W$ be a translation for some vector $w_0 \in W$ so that $T_{w_0}(w) = w + w_0$ for all $w \in W$. Let $f = T_{w_0} \circ g$ so that $f(v) = g(v) + w_0$ for all $v \in V$. Then

$$f(v_1) - f(v_2) = g(v_1) + w_0 - g(v_2) - w_0 = g(v_1) - g(v_2) = g(v_1 - v_2),$$

hence f is affine.

$\Rightarrow$ Let $f : V \rightarrow W$ be affine and $g : V \rightarrow W$ its associated linear map. Since f is affine we have for all $v \in V$

$$f(v) - f(0_V) = g(v - 0_V) = g(v) - g(0_V) = g(v)$$

where we use the linearity of g and [Lemma 3.15](#). Writing $w_0 = f(0_V)$ we thus have

$$f(v) = g(v) + w_0$$

so that f is the composition of the linear map g and the translation $T_{w_0} : W \rightarrow W$, $w \mapsto w + w_0$. $\square$

Example 7.5 Let $\mathbf{A} \in M_{m,n}(\mathbb{K})$, $\vec{b} \in \mathbb{K}^m$ and

$$f_{\mathbf{A},\vec{b}} : \mathbb{K}^n \rightarrow \mathbb{K}^m, \quad \vec{x} \mapsto \mathbf{A}\vec{x} + \vec{b}.$$

Then $f_{\mathbf{A},\vec{b}}$ is an affine map whose associated linear map is $f_{\mathbf{A}}$. Conversely, combining [Lemma 3.18](#) and [Proposition 7.4](#), we see that every affine map $\mathbb{K}^n \rightarrow \mathbb{K}^m$ is of the form $f_{\mathbf{A},\vec{b}}$ for some matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and vector $\vec{b} \in \mathbb{K}^m$.

An affine subspace of a $\mathbb{K}$ -vector space V is a translation of a subspace by some fixed vector v_0 .

Definition 7.6 (Affine subspace) Let V be a $\mathbb{K}$ -vector space. An affine subspace of V is a subset of the form

$$U + v_0 = \{u + v_0 \mid u \in U\},$$

where $U \subset V$ is a subspace and $v_0 \in V$. We call U the associated vector space to the affine subspace $U + v_0$ and we say that $U + v_0$ is parallel to U .

Example 7.7 Let $V = \mathbb{R}^2$ and $U = \text{span}\{\vec{e}_1 + \vec{e}_2\} = \{s(\vec{e}_1 + \vec{e}_2) \mid s \in \mathbb{R}\}$ where here, as usual, $\{\vec{e}_1, \vec{e}_2\}$ denotes the standard basis of $\mathbb{R}^2$. So U is the line through the origin $0_{\mathbb{R}^2}$ defined by the equation $y = x$. By definition, for all $\vec{v} \in \mathbb{R}^2$ we have

$$U + \vec{v} = \{\vec{v} + s\vec{w} \mid s \in \mathbb{R}\},$$

where we write $\vec{w} = \vec{e}_1 + \vec{e}_2$. So for each $\vec{v} \in \mathbb{R}^2$, the affine subspace $U + \vec{v}$ is a line in $\mathbb{R}^2$, the translation by the vector $\vec{v}$ of the line defined by $y = x$.

7.2 Quotient vector spaces

Let U be a subspace of a $\mathbb{K}$ -vector space V . We want to make sense of the notion of dividing V by U . It turns out that there is a natural way to do this and moreover, the quotient V/U again carries the structure of a $\mathbb{K}$ -vector space. The idea is to define V/U to be the set of all translations of the subspace U , that is, we consider the set of subsets

$$V/U = \{U + v \mid v \in V\}.$$

We have to define what it means to add affine subspaces $U + v_1$ and $U + v_2$ and what it means to scale $U + v$ by a scalar $s \in \mathbb{K}$. Formally, it is tempting to define $0_{V/U} = U + 0_V$ and

$$(7.1) \quad (U + v_1) +_{V/U} (U + v_2) = U + (v_1 + v_2)$$

for all $v_1, v_2 \in V$ as well as

$$(7.2) \quad s \cdot_{V/U} (U + v) = U + (sv)$$

for all $v \in V$ and $s \in \mathbb{K}$. However, we have to make sure that these operations are well defined. We do this with the help of the following lemma.

Lemma 7.8 *Let $U \subset V$ be a subspace. Then any vector $v \in V$ belongs to a unique affine subspace parallel to U , namely $U + v$. In particular, two affine subspaces $U + v_1$ and $U + v_2$ are either equal or have empty intersection.*

Proof Since $0_V \in U$, we have $v \in (U + v)$, hence we only need to show that if $v \in (U + \hat{v})$ for some vector $\hat{v}$, then $U + v = U + \hat{v}$. Assume $v \in (U + \hat{v})$ so that $v = u + \hat{v}$ for some vector $u \in U$. Suppose $w \in (U + \hat{v})$. We need to show that then also $w \in (U + v)$. Since $w \in (U + \hat{v})$ we have $w = \hat{u} + \hat{v}$ for some vector $\hat{u} \in U$. Using that $\hat{v} = v - u$, we obtain

$$w = \hat{u} + v - u = \hat{u} - u + v$$

Since U is a subspace we have $\hat{u} - u \in U$ and hence $w \in (U + v)$.

Conversely, suppose $w \in (U + v)$, it follows exactly as before that then $w \in (U + \hat{v})$ as well. $\square$

We are now going to show that (7.1) and (7.2) are well defined. We start with (7.1). Let $v_1, v_2 \in V$ and $w_1, w_2 \in V$ such that

$$U + v_1 = U + w_1 \quad \text{and} \quad U + v_2 = U + w_2.$$

We need to show that $U + (v_1 + v_2) = U + (w_1 + w_2)$. By Lemma 7.8 it suffices to show that $w_1 + w_2$ is an element of $U + (v_1 + v_2)$. Since $U + w_1 = U + v_1$ it follows that $w_1 \in (U + v_1)$ so that $w_1 = u_1 + v_1$ for some element $u_1 \in U$. Likewise it follows that $w_2 = u_2 + v_2$ for some element $u_2 \in U$. Hence

$$w_1 + w_2 = u_1 + u_2 + v_1 + v_2.$$

Since U is a subspace, we have $u_1 + u_2 \in U$ and thus it follows that $w_1 + w_2$ is an element of $U + (v_1 + v_2)$.

For (7.2) we need to show that if $v \in V$ and $w \in V$ are such that $U + v = U + w$, then $U + (sv) = U + (sw)$ for all $s \in \mathbb{K}$. Again, applying Lemma 7.8 we only need to show that $sw \in U + (sv)$. Since $U + v = U + w$ it follows that there exists $u \in U$ with $w = u + v$. Hence $sw = su + sv$ and U being a subspace, we have $su \in U$ and thus sw lies in $U + (sv)$, as claimed.

Having equipped V/U with addition $+_{V/U}$ defined by (7.1) and scalar multiplication $\cdot_{V/U}$ defined by (7.2), we need to show that V/U with zero vector $U + 0_V$ is indeed a $\mathbb{K}$ -vector space. All the properties of Definition 3.1 for V/U are however simply a consequence of the corresponding property for V . For instance commutativity of vector addition in V/U follows from the commutativity of vector in addition in V , that is, for all $v_1, v_2 \in V$ we have

$$(U + v_1) +_{V/U} (U + v_2) = U + (v_1 + v_2) = U + (v_2 + v_1) = (U + v_2) +_{V/U} (U + v_1).$$

The remaining properties follow similarly.

Notice that we have a surjective mapping

$$p: V \rightarrow V/U, \quad v \mapsto U + v.$$

which satisfies

$$p(v_1 + v_2) = U + (v_1 + v_2) = (U + v_1) +_{V/U} (U + v_2) = p(v_1) +_{V/U} p(v_2)$$

for all $v_1, v_2 \in V$ and

$$p(sv) = U + (sv) = s \cdot_{V/U} (U + v) = s \cdot_{V/U} p(v).$$

for all $v \in V$ and $s \in \mathbb{K}$. Therefore, the mapping p is linear.

Definition 7.9 (Quotient vector space) The vector space V/U is called the *quotient (vector) space of V by U* . The linear map $p : V \rightarrow V/U$ is called the *canonical surjection* from V to V/U .

The mapping $p : V \rightarrow V/U$ satisfies

$$p(v) = 0_{V/U} = U + 0_V \iff v \in U$$

and hence $\text{Ker}(p) = U$. This gives:

Proposition 7.10 Suppose the $\mathbb{K}$ -vector space V is finite dimensional. Then V/U is finite dimensional as well and

$$\dim(V/U) = \dim(V) - \dim(U).$$

Proof Since p is surjective it follows that V/U is finite dimensional as well. Hence we can apply [Theorem 3.76](#) and obtain

$$\dim V = \dim \text{Ker}(p) + \dim \text{Im}(p) = \dim U + \dim(V/U),$$

where we use that $\text{Im}(p) = V/U$ and $\text{Ker}(p) = U$. □

Example 7.11 (Special cases)

- (i) In the case where $U = V$ we obtain $V/U = \{0_{V/U}\}$.
- (ii) In the case where $U = \{0_V\}$ we obtain that V/U is isomorphic to V .

Exercises

Exercise 7.12 Show that the image of an affine subspace under an affine map is again an affine subspace and that the preimage of an affine subspace under an affine map is again an affine subspace or empty (cf. [Proposition 3.26](#)).

Part 2

Linear Algebra II

Symmetry and groups

WEEK 1

8.1 Symmetry

The notion of a group arose by trying to formalise the concept of *symmetry*. Roughly speaking, given a non-empty set $\mathcal{X}$ with some extra structure, a *symmetry or symmetry transformation of $\mathcal{X}$* is a bijective transformation $\sigma : \mathcal{X} \rightarrow \mathcal{X}$ that respects the extra structure. For simplicity, we ignore any extra structure, so for us a symmetry of a set $\mathcal{X}$ is simply a bijective mapping from $\mathcal{X}$ to itself.

Example 8.1

- (i) Let $n \in \mathbb{N}$. A permutation is a symmetry of the set $\mathcal{X} = \{1, 2, \dots, n\}$.
- (ii) Let V be a $\mathbb{K}$ -vector space and $v_0 \in V$. The translation $T_{v_0} : V \rightarrow V, v \mapsto v + v_0$ by the vector v_0 is a symmetry of V .
- (iii) Let $\mathcal{X}$ be any non-empty set. The identity transformation $\text{Id}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X}$ defined by $\text{Id}_{\mathcal{X}}(x) = x$ for all $x \in \mathcal{X}$ is a symmetry of $\mathcal{X}$.

Often the set $\mathcal{X}$ is a subset of some larger set $\mathcal{Z}$ and the symmetries of $\mathcal{X}$ arise as bijective mappings $\sigma : \mathcal{Z} \rightarrow \mathcal{Z}$ that leave $\mathcal{X}$ *invariant*, that is, $\sigma(x) \in \mathcal{X}$ for all $x \in \mathcal{X}$. We illustrate this with two examples:

Example 8.2

- (i) Consider $\mathcal{Z} = \mathbb{R}^2$ and $\mathcal{X}$ to be the circle of radius $r > 0$ centred at the origin $0_{\mathbb{R}^2}$, that is, $\mathcal{X} = \{\vec{x} = (x_i)_{1 \leq i \leq 2} \in \mathbb{R}^2 \mid (x_1)^2 + (x_2)^2 = r^2\}$. Let $\theta \in \mathbb{R}$ and

$$\mathbf{R}_{\theta} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$$

so that $f_{\mathbf{R}_{\theta}} : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ is the counter-clockwise rotation around the origin $0_{\mathbb{R}^2}$ with angle θ . A rotation does not change the length of a vector and hence $f_{\mathbf{R}_{\theta}}(\vec{x}) \in \mathcal{X}$ for each element $\vec{x} \in \mathcal{X}$. The restriction $\sigma = f_{\mathbf{R}_{\theta}}|_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X}$ of the rotation $f_{\mathbf{R}_{\theta}}$ to the circle $\mathcal{X}$ is thus a symmetry of the circle. Notice that not all symmetries of the circle are restrictions of rotations. The linear mapping

$$f : \mathbb{R}^2 \rightarrow \mathbb{R}^2, \quad \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \mapsto \begin{pmatrix} x_1 \\ -x_2 \end{pmatrix}$$

is the reflection along the x_1 -axis and hence restricts to be a bijective mapping from the circle $\mathcal{X}$ onto itself. It is thus also a symmetry of the circle.

- (ii) Let $n \in \mathbb{N}$ with $n \geq 3$. We consider a regular polygon $\mathcal{X}$ with n sides centred at the origin in $\mathcal{Z} = \mathbb{R}^2$ an so that $(1, 0) \in \mathcal{X}$. Clearly, not every rotation of $\mathbb{R}^2$ restricts to be a symmetry of $\mathcal{X}$, but only rotations with angle $2\pi k/n$ where $k \in \{0, 1, 2, \dots, n-1\}$. We thus have n rotation symmetries arising from the

matrices

$$\begin{pmatrix} \cos \frac{2\pi k}{n} & -\sin \frac{2\pi k}{n} \\ \sin \frac{2\pi k}{n} & \cos \frac{2\pi k}{n} \end{pmatrix}.$$

In addition, the reflection along the x_1 -axis is a symmetry of $\mathcal{X}$.

Remark 8.3 The composition of two symmetries of a set $\mathcal{X}$ is again a symmetry of $\mathcal{X}$ and composing symmetries satisfies the following *fundamental* properties:

- If $\sigma, \pi, \tau : \mathcal{X} \rightarrow \mathcal{X}$ are symmetries, then

$$(\sigma \circ \pi) \circ \tau = \sigma \circ (\pi \circ \tau)$$

- The identity transformation $\text{Id}_{\mathcal{X}}$ is a symmetry of $\mathcal{X}$ and for all symmetries $\sigma : \mathcal{X} \rightarrow \mathcal{X}$, we have

$$\sigma \circ \text{Id}_{\mathcal{X}} = \sigma = \text{Id}_{\mathcal{X}} \circ \sigma$$

- For each symmetry $\sigma : \mathcal{X} \rightarrow \mathcal{X}$ there exists an inverse symmetry $\sigma^{-1} : \mathcal{X} \rightarrow \mathcal{X}$ so that

$$\sigma \circ \sigma^{-1} = \text{Id}_{\mathcal{X}} = \sigma^{-1} \circ \sigma.$$

8.2 Groups

We have defined the permutations S_n to be the bijective mappings of the set $\mathcal{X}_n = \{1, 2, \dots, n\}$, hence by definition, they are symmetries of $\mathcal{X}_n$. Recall that in addition, every permutation $\sigma \in S_n$ also gives rise to a bijective (linear) mapping from $\mathbb{K}^n \rightarrow \mathbb{K}^n$ defined by $\vec{e}_i \mapsto \vec{e}_{\sigma(i)}$, where $\{\vec{e}_1, \dots, \vec{e}_n\}$ denotes the standard basis of $\mathbb{K}^n$. Hence, every permutation also gives a symmetry of $\mathbb{K}^n$. The permutations thus make an appearance as symmetries of two different sets, $\mathcal{X}_n$ and $\mathbb{K}^n$. This suggests that a more detailed picture of a symmetry is needed. It turns out that a symmetry is the interplay of two mathematical notions, the notion of a *group* and the *action of a group on a set* $\mathcal{X}$. We start with the definition of a group, c.f. [Remark 8.3](#):

Definition 8.4 (Group) A *group* is a pair $(G, *_G)$ consisting of a set G together with a binary operation $*_G : G \times G \rightarrow G$, called *group operation*, so that the following properties hold:

- (i) The group operation $*_G$ is associative, that is,

$$(a *_G b) *_G c = a *_G (b *_G c) \quad \text{for all } a, b, c \in G.$$

- (ii) There exists an element $e_G \in G$ such that

$$e_G *_G a = a = a *_G e_G \quad \text{for all } a \in G.$$

The element e_G is unique (see below) and is called the *identity element* of G .

- (iii) For each $a \in G$ there exists an element $b \in G$ such that

$$a *_G b = e_G = b *_G a.$$

The element b is unique (see below) and called the *inverse* of a and is commonly denoted by a^{-1} .

Example 8.5 (Examples of groups)

- (i) The symmetries of a set $\mathcal{X}$ form a group G , often denoted by $\text{Sym}(\mathcal{X})$, where $*_G = \circ$ is the composition of mappings. The identity element is the identity mapping $e_G = \text{Id}_{\mathcal{X}}$ and the inverse of each symmetry σ is the mapping inverse σ^{-1} . In particular, for $n \in \mathbb{N}$, the permutations of $\mathcal{X}_n = \{1, 2, \dots, n\}$ form a group $G = S_n$ with $*_G = \circ$ and $e_G = 1$, the identity permutation.
- (ii) A field $\mathbb{K}$ gives rise to two groups. The *additive group of the field* where $G = \mathbb{K}$ and $*_G = +_{\mathbb{K}}$ and the *multiplicative group of the field* where $G = \mathbb{K}^*$ and $*_G = \cdot_{\mathbb{K}}$. For the additive group we have $e_G = 0_{\mathbb{K}}$ and the inverse of $x \in \mathbb{K}$ is $-x$. For the multiplicative group we have $e_G = 1_{\mathbb{K}}$ and the inverse of $x \in \mathbb{K}^*$ is $\frac{1}{x}$.
- (iii) A $\mathbb{K}$ -vector space V gives rise to a group where $G = V$ and $*_G = +_V$. Here the identity element is the zero vector $e_G = 0_V$ and the inverse of $v \in V$ is $-v$.
- (iv) Let $n \in \mathbb{N}$. The invertible $n \times n$ matrices with entries in $\mathbb{K}$ form a group G commonly denoted by $\text{GL}_n(\mathbb{K})$ or $\text{GL}(n, \mathbb{K})$. Here $*_G$ is matrix multiplication, $e_G = \mathbf{1}_n$, the identity matrix of size n and the inverse of a group element is the matrix inverse. GL is an abbreviation of *general linear*.

Remark 8.6

- A group with finitely many elements is called *finite*. The group of permutations S_n is an example of a finite group. A finite field gives rise to two finite groups.
- Notice that we do not require the group operation $*_G : G \times G \rightarrow G$ to be commutative, so in general $a *_G b \neq b *_G a$. As an example consider $G = \text{GL}(n, \mathbb{K})$ where $*_G$ is matrix multiplication. If the group operation $*_G$ is commutative, then the group is called *Abelian or commutative*. The examples (ii) and (iii) above are examples of Abelian groups. The permutation group S_n is Abelian only for $n = 1, 2$.
- Often we write $+_G$ instead of $*_G$ and 0_G instead of e_G when the group is Abelian.
- Some authors write 1_G instead of e_G and/or $\cdot_G$ instead of $*_G$.
- As always, the subscript G is often omitted so that we write $*$ instead of $*_G$ and e or 1 instead of e_G . Like for fields, $*$ or $*_G$ is often omitted entirely so that we write ab instead of $a *_G b$.

Similar to fields, the definition of a group implies some basic properties:

Proposition 8.7 Let $(G, *_G)$ be a group. Then

- (i) the identity element e_G is unique;
- (ii) for all $a \in G$, the inverse a^{-1} is unique.

Proof

- (i) Suppose e_G and $\hat{e}_G$ are identity elements for G . Then

$$e_G = e_G *_G \hat{e}_G = \hat{e}_G.$$

- (ii) Suppose $a \in G$ and both b and c are inverse elements for a . Then

$$b = b *_G e_G = b *_G (a *_G c) = (b *_G a) *_G c = e_G *_G c = c.$$

□

Similar to vector spaces and fields, groups allow for the notion of a subgroup.

Definition 8.8 (Subgroup) A non-empty subset H of a group G is called a *subgroup* if for all $a, b \in H$, we have $a *_G b \in H$ and for all $a \in H$, we have $a^{-1} \in H$.

Notice that if $H \subset G$ is a subgroup, the non-emptiness condition implies that there exists $a \in H$. Therefore, $a^{-1} \in H$ and hence $a *_G a^{-1} = e_G \in H$. We can thus equip H with the structure of a group as well by defining $e_H = e_G$ and $a *_H b = a *_G b$ for all $a, b \in H$.

Example 8.9 The set of integers $\mathbb{Z}$ is a subgroup of the Abelian group $(\mathbb{Q}, +)$, where $+$ denotes usual addition of rational numbers. Indeed $0 \in \mathbb{Z}$ and the sum of two integers is again an integer. Recall that for $m \in \mathbb{Z}$, the notation m^{-1} refers to the inverse element of m with respect to the group operation. So here m^{-1} is the additive inverse of $m \in \mathbb{Z}$, that is $-m$. Since $-m \in \mathbb{Z}$ for all $m \in \mathbb{Z}$, we conclude that $\mathbb{Z}$ is an (Abelian) subgroup of $(\mathbb{Q}, +)$.

Example 8.10 A subspace $U \subset V$ of a $\mathbb{K}$ -vector space V is a subgroup of the Abelian group $(V, +_V)$.

Example 8.11 Let $\text{SL}(n, \mathbb{K})$ denote the subset of $\text{GL}(n, \mathbb{K})$ consisting of matrices of determinant 1. The set $\text{SL}(n, \mathbb{K})$ is non-empty since it contains $\mathbf{1}_n$. Furthermore, the product rule for the determinant [Proposition 5.21](#) implies that if $\mathbf{A}, \mathbf{B} \in \text{SL}(n, \mathbb{K})$, then so is the matrix product $\mathbf{AB}$. [Corollary 5.22](#) furthermore implies that if $\mathbf{A} \in \text{SL}(n, \mathbb{K})$, then so is $\mathbf{A}^{-1}$. It follows that $\text{SL}(n, \mathbb{K})$ – commonly also denoted by $\text{SL}_n(\mathbb{K})$ – is a subgroup of $\text{GL}(n, \mathbb{K})$ called the *special linear group*.

Example 8.12 The trigonometric identities for $\sin$ and $\cos$ imply that $\mathbf{R}_\theta \mathbf{R}_\vartheta = \mathbf{R}_{\theta+\vartheta}$, where $\theta, \vartheta \in \mathbb{R}$. Since $\mathbf{R}_0 = \mathbf{1}_2 \in \text{SL}(2, \mathbb{R})$ and $\det \mathbf{R}_\theta = 1$ for all $\theta \in \mathbb{R}$, we conclude that the rotations $\{\mathbf{R}_\theta | \theta \in \mathbb{R}\}$ around the origin $0_{\mathbb{R}^2}$ form a subgroup of $\text{SL}(2, \mathbb{R})$. The group of rotations in $\mathbb{R}^2$ is denoted by $\text{SO}(2)$. Later on we will encounter the *orthogonal group* $\text{O}(n)$ and the *special orthogonal group* $\text{SO}(n)$, the latter of which generalises $\text{SO}(2)$ to higher dimensions.

8.3 Group actions

In order to tie the notion of a group more closely to the notion of a symmetry, we need the concept of a group G acting on a set $\mathcal{X}$. This section – which we include for the interested reader – goes beyond the usual material in a Linear Algebra course and is *not examinable*.

Definition 8.13 (Group action) Let G be a group and $\mathcal{X}$ a non-empty set. A (*left*) *group action* of G on $\mathcal{X}$ is a mapping $\phi : G \times \mathcal{X} \rightarrow \mathcal{X}$ such that for all $x \in \mathcal{X}$

$$\phi(e_G, x) = x$$

and

$$\phi(a *_G b, x) = \phi(a, \phi(b, x))$$

for all $a, b \in G$ and $x \in \mathcal{X}$.

Remark 8.14

- The first condition simply requests that the identity element e_G of G acts *trivially*, that is, nothing happens to the elements of $\mathcal{X}$ when acting with e_G .
- The second condition requests that acting with $a *_G b$ corresponds to first acting with b and then acting with a .
- Notice that for each fixed $a \in G$ we obtain a mapping $\phi_a : \mathcal{X} \rightarrow \mathcal{X}$ defined by $\phi_a(x) = \phi(a, x)$. The above properties imply that for all $a \in G$ we have $\phi_a \circ \phi_{a^{-1}} = \phi_{a^{-1}} \circ \phi_a = \text{Id}_{\mathcal{X}}$, hence $\phi_a : \mathcal{X} \rightarrow \mathcal{X}$ is bijective and hence a symmetry of $\mathcal{X}$.

Example 8.15

(i) Every group G acts on itself. We take $\mathcal{X} = G$ and define

$$\phi : G \times G \rightarrow G, \quad (a, b) \mapsto \phi(a, b) = a *_G b.$$

Then for all $a \in G$ we have $\phi(e_G, a) = e_G *_G a = a$. Furthermore, for all $a, b, c \in G$ we have

$$\phi(a *_G b, c) = (a *_G b) *_G c = a *_G (b *_G c) = a *_G \phi(b, c) = \phi(a, \phi(b, c))$$

so that ϕ does indeed define an action of G on itself.

(ii) Consider $\mathcal{X} = \mathbb{R}^2$ and $G = \text{SO}(2)$. We define an action

$$\phi : G \times \mathcal{X} \rightarrow \mathcal{X}, \quad (\mathbf{R}_\theta, \vec{x}) \mapsto \phi(\mathbf{R}_\theta, \vec{x}) = \mathbf{R}_\theta \vec{x}$$

which rotates a vector $\vec{x} \in \mathbb{R}^2$ counter-clockwise around the origin $0_{\mathbb{R}^2}$ by the angle θ . Here $*_G$ is just matrix multiplication, so we have for all $\vec{x} \in \mathbb{R}^2$ and $\mathbf{R}_\theta, \mathbf{R}_\vartheta \in \text{SO}(2)$

$$\phi(\mathbf{R}_\theta \mathbf{R}_\vartheta, \vec{x}) = \mathbf{R}_\theta \mathbf{R}_\vartheta \vec{x} = \mathbf{R}_\theta \phi(\mathbf{R}_\vartheta, \vec{x}) = \phi(\mathbf{R}_\theta, \phi(\mathbf{R}_\vartheta, \vec{x})).$$

Furthermore, since $e_{\text{SO}(2)} = \mathbf{R}_0 = \mathbf{1}_2$, we have for all $\vec{x} \in \mathbb{R}^2$

$$\phi(e_{\text{SO}(2)}, \vec{x}) = \mathbf{1}_2 \vec{x} = \vec{x}.$$

It follows that ϕ does indeed define an action of $\text{SO}(2)$ on $\mathbb{R}^2$.

(iii) Let $n \in \mathbb{N}$ and $\mathcal{X} = M_{n,n}(\mathbb{K})$. The general linear group $\text{GL}(n, \mathbb{K})$ acts on $\mathcal{X}$ by conjugation. We define

$$\phi : \text{GL}(n, \mathbb{K}) \times M_{n,n}(\mathbb{K}) \rightarrow M_{n,n}(\mathbb{K}), \quad (\mathbf{C}, \mathbf{A}) \mapsto \phi(\mathbf{C}, \mathbf{A}) = \mathbf{C} \mathbf{A} \mathbf{C}^{-1}.$$

Then for all $\mathbf{A} \in M_{n,n}(\mathbb{K})$ we have

$$\phi(e_G, \mathbf{A}) = \phi(\mathbf{1}_n, \mathbf{A}) = \mathbf{1}_n \mathbf{A} (\mathbf{1}_n)^{-1} = \mathbf{A}$$

where we use that $e_{\text{GL}(n, \mathbb{K})} = \mathbf{1}_n$. Moreover, for $\mathbf{C}, \mathbf{C}' \in \text{GL}(n, \mathbb{K})$ and $\mathbf{A} \in M_{n,n}(\mathbb{K})$, we have

$$\begin{aligned} \phi(\mathbf{C} \mathbf{C}', \mathbf{A}) &= \mathbf{C} \mathbf{C}' \mathbf{A} (\mathbf{C} \mathbf{C}')^{-1} = \mathbf{C} \mathbf{C}' \mathbf{A} (\mathbf{C}')^{-1} \mathbf{C}^{-1} \\ &= \mathbf{C} \phi(\mathbf{C}', \mathbf{A}) \mathbf{C}^{-1} = \phi(\mathbf{C}, \phi(\mathbf{C}', \mathbf{A})), \end{aligned}$$

where we use that for all $\mathbf{C}, \mathbf{C}' \in \text{GL}(n, \mathbb{K})$, we have $(\mathbf{C} \mathbf{C}')^{-1} = (\mathbf{C}')^{-1} \mathbf{C}^{-1}$. It follows that ϕ does indeed define an action of $\text{GL}(n, \mathbb{K})$ on $M_{n,n}(\mathbb{K})$.

- (iv) Let V be a $\mathbb{K}$ -vector space and $U \subset V$ a subspace. Taking $G = U$ with $*_G = +_U$ and $\mathcal{X} = V$, the group G acts by translation. We define

$$\phi : U \times V \rightarrow V, \quad (u, v) \mapsto \phi(u, v) = u +_V v.$$

Since $e_G = 0_U = 0_V$, we have for all $v \in V$

$$\phi(e_G, v) = 0_V +_V v = v.$$

Moreover, for all $u_1, u_2 \in U$ and $v \in V$ we have

$$\phi(u_1 +_U u_2, v) = (u_1 +_U u_2) +_V v = u_1 +_V \phi(u_2, v) = \phi(u_1, \phi(u_2, v)),$$

where we use that $+_U : U \times U \rightarrow U$ is the restriction of $+_V : V \times V \rightarrow V$ to $U \times U \subset V \times V$. We conclude that ϕ defines an action of the subspace U on V .

- (v) Let $n \in \mathbb{N}$. A permutation $\sigma \in S_n$ acts on $\mathcal{X}_n = \{1, 2, \dots, n\}$ by

$$\phi : S_n \times \mathcal{X}_n \rightarrow \mathcal{X}_n \quad (\sigma, m) \mapsto \phi(\sigma, m) = \sigma(m).$$

We leave it to the reader to check that this is indeed an action. In addition, a permutation $\sigma \in S_n$ does also act on $\mathbb{K}^n$ by the rule

$$\phi(\sigma, \vec{x}) = \mathbf{P}_\sigma \vec{x},$$

where $\vec{x} \in \mathbb{K}^n$ and $\mathbf{P}_\sigma$ is the permutation matrix associated to $\sigma \in S_n$, c.f. [Definition 5.28](#).

A particularly important class of group actions arises when $(G, *_G)$ is the Abelian group $(\mathbb{R}, +)$ or its subgroup $(\mathbb{Z}, +)$. This case arises for instance when the set $\mathcal{X}$ is a *phase space* (roughly speaking, the set of different physical states) of a physical system and the action describes the evolution of the system under the progression of time.

Definition 8.16 (Dynamical system) Let $\mathcal{X}$ be a non-empty set. A *time-discrete dynamical system* is an action of $(\mathbb{Z}, +)$ on $\mathcal{X}$. A *time-continuous dynamical system* is an action of $(\mathbb{R}, +)$ on $\mathcal{X}$.

Often the term dynamical system is also used when the action is only defined for all non-negative times $\mathbb{R}_0^+ = \{t \in \mathbb{R} | t \geq 0\}$ or $\mathbb{N}_0 = \{t \in \mathbb{Z} | t \geq 0\}$.

Example 8.17

- (i) Let $\mathcal{X} \subset \mathbb{R}^3$ denote the set of all points in our solar system. An asteroid initially at rest at the position $x_0 \in \mathcal{X}$ will move under the influence of gravity. Let x_t denote the position of the asteroid after time $t \in \mathbb{R}$ has passed. The mapping

$$\phi : \mathbb{R}_0^+ \times \mathcal{X} \rightarrow \mathcal{X}, \quad (t, x_0) \mapsto \phi(t, x_0) = x_t$$

describing the movement of the asteroid is then a time-continuous dynamical system.

- (ii) Let $\mathcal{X} = \{0, 1\}^N$ denote the carrier status of a contagious disease of each individual of a population of size $N \in \mathbb{N}$. So $x \in \mathcal{X}$ is a list of length N containing 0s and 1s, where the k -th entry reflects the carrier status of the k -th member of the population, 0 for non-carriers and 1 for carriers. Let $x_0 \in \mathcal{X}$ denote the carrier status at some initial time $t = 0$ and for $m \in \mathbb{N}_0$ let x_m denote the carrier status after m days have passed. The mapping

$$\phi : \mathbb{N}_0 \times \mathcal{X} \rightarrow \mathcal{X}, \quad (m, x_0) \mapsto \phi(m, x_0) = x_m$$

describing the progression of the disease in the population is then a time-discrete dynamical system.

Given a group action on some set $\mathcal{X}$ and some point $x \in \mathcal{X}$, we consider the subset of elements of $\mathcal{X}$ that can be reached by acting with all the groups elements of G . This subset is known as the orbit of x . More precisely:

Definition 8.18 (Orbit) Let $\mathcal{X}$ be a non-empty set, $\phi : G \times \mathcal{X} \rightarrow \mathcal{X}$ an action of the group $(G, *_G)$ on $\mathcal{X}$ and $x \in \mathcal{X}$. The *orbit of $x \in \mathcal{X}$ under G* (or sometimes *G -orbit of x*) is the subset

$$G *_G x = \{\phi(a, x) \in \mathcal{X} | a \in G\}.$$

The set of all G -orbits in $\mathcal{X}$ is denoted by $\mathcal{X}/G$.

In the time-continuous dynamical system above, the orbit of $x_0 \in \mathcal{X}$ consists of the points x_t where $t \in \mathbb{R}_0^+$ and x_t is the time t position of the asteroid with initial position x_0 . The orbit is thus the trajectory of the asteroid as time progresses. Therefore, the mathematical concept of orbit is a generalisation of the standard use of the term orbit.

Example 8.19

- (i) Consider the action of $SO(2)$ on $\mathbb{R}^2$ from above. The orbit of $\vec{x} \neq 0_{\mathbb{R}^2}$ consists of all points in $\mathbb{R}^2$ obtained by rotating $\vec{x}$ counter-clockwise around the origin. Since the rotation angle can be chosen arbitrarily, the orbit of $\vec{x}$ is the circle of all points of $\mathbb{R}^2$ that have the same length as $\vec{x}$. On the other hand, the orbit of $0_{\mathbb{R}^2}$ only consists of $0_{\mathbb{R}^2}$, that is, we have

$$SO(2) *_G 0_{\mathbb{R}^2} = \{0_{\mathbb{R}^2}\}.$$

In this particular case we have a complete picture of all possible orbits, an orbit is either the zero vector or else a circle centred at the origin, hence

$$\mathcal{X}/G = \mathbb{R}^2/SO(2) = \{0_{\mathbb{R}^2}\} \cup \{\text{circle of radius } r \text{ centred at } 0_{\mathbb{R}^2} | r > 0\}.$$

- (ii) Consider the action of $GL(n, \mathbb{K})$ on $M_{n,n}(\mathbb{K})$ from above. Let $\mathbf{D} = \text{diag}(\lambda_1, \dots, \lambda_n)$ be a diagonal matrix with entries $\lambda_1, \dots, \lambda_n \in \mathbb{K}$. The $GL(n, \mathbb{K})$ -orbit of $\mathbf{D}$ then consists of all $n \times n$ -matrices with entries in $\mathbb{K}$ that are diagonalisable with eigenvalues $\lambda_1, \dots, \lambda_n$. A complete description of the set of orbits $M_{n,n}(\mathbb{K})/GL(n, \mathbb{K})$ is out of reach for us at this point, we will however have more to say about this in the Linear Algebra II module.
- (iii) We consider the action of S_2 on $\mathbb{R}^2$ as defined above. The orbit of a vector $\vec{v} = \begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$ with $x \neq y$ is the subset

$$S_2 *_G \vec{v} = \left\{ \begin{pmatrix} x \\ y \end{pmatrix}, \begin{pmatrix} y \\ x \end{pmatrix} \right\}.$$

On the other hand, the orbit of a vector $\vec{v} = \begin{pmatrix} x \\ x \end{pmatrix} \in \mathbb{R}^2$ is just $\{\vec{v}\}$. For a vector of the first type, either $x > y$ or $x < y$. The orbit of each such vector can thus be represented uniquely by a vector $\vec{v} = \begin{pmatrix} x \\ y \end{pmatrix}$ with $y > x$. The vectors of the second type lie on the axis defined by the equation $y = x$. We thus have a bijective mapping from $\mathbb{R}^2/S_2$ to the *half plane* $\mathbb{H} = \{(x, y) \in \mathbb{R}^2 | y \geq x\}$.

Remark 8.20 (Quotient vector space) Given a subspace $U \subset V$, we have seen that the Abelian group $(G, *_G) = (U, +_U)$ acts on $\mathcal{X} = V$ by translation. In this sense $U + v$ is simply the orbit of v under this action and V/U is the set of orbits $\mathcal{X}/G$. Furthermore, [Lemma 7.8](#) is a special case of a more general statement about orbits: If a group $(G, *_G)$ acts on a non-empty set $\mathcal{X}$, then every element $x \in \mathcal{X}$ belongs to a unique G -orbit, namely $G *_G x$. In particular two orbits $G *_G x_1$ and $G *_G x_2$ are either equal or have empty intersection.

Exercises

Exercise 8.21 Show that mapping $S_n \times \mathbb{K}^n \rightarrow \mathbb{K}^n$ given in [Example 8.15](#) does indeed define an action.

Exercise 8.22 Prove the statement about orbits from [Remark 8.20](#).

Bilinear forms

WEEK 2

9.1 Definitions and basic properties

So far in Linear Algebra we have dealt with vector spaces without thinking much about geometric aspects. For example, for an abstract vector space we cannot say what the angle between two vectors is. Likewise, we are not able to talk about the distance between elements of a vector space. To make sense of these notions, the vector space needs further structure given by an *inner product*.

An inner product is a special case of a bilinear form. The prototypical example of a bilinear form is the *standard scalar product* on $\mathbb{R}^n$ that you might already know. Recall that for $\vec{x} = (x_i)_{1 \leq i \leq n}$ and $\vec{y} = (y_i)_{1 \leq i \leq n} \in \mathbb{R}^n$, we define

$$(9.1) \quad \vec{x} \cdot \vec{y} = \sum_{i=1}^n x_i y_i = x_1 y_1 + \cdots + x_n y_n.$$

It is also common to write $\langle \vec{x}, \vec{y} \rangle$ instead of $\vec{x} \cdot \vec{y}$. As we have already seen in [Example 5.3](#), the standard scalar product is an example of a 2-multilinear map.

Definition 9.1 (Bilinear form) Let V be a $\mathbb{K}$ -vector space. A *bilinear form* on V is a 2-multilinear map with values in $\mathbb{K}$

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{K}, \quad (v_1, v_2) \mapsto \langle v_1, v_2 \rangle.$$

That is, for all $s_1, s_2 \in \mathbb{K}$ and all $v_1, v_2, v_3 \in V$ we have

$$\langle s_1 v_1 + s_2 v_2, v_3 \rangle = s_1 \langle v_1, v_3 \rangle + s_2 \langle v_2, v_3 \rangle$$

as well as

$$\langle v_3, s_1 v_1 + s_2 v_2 \rangle = s_1 \langle v_3, v_1 \rangle + s_2 \langle v_3, v_2 \rangle.$$

We say that $\langle \cdot, \cdot \rangle$ is *symmetric* if $\langle v_1, v_2 \rangle = \langle v_2, v_1 \rangle$ for all $v_1, v_2 \in V$ and *alternating* if $\langle v, v \rangle = 0$ for all $v \in V$.

Example 9.2 (Bilinear forms)

- (i) The standard scalar product defined by the rule (9.1) is a bilinear form on $\mathbb{R}^n$.
- (ii) Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{K})$ be a matrix. Using matrix multiplication, we define a mapping

$$(9.2) \quad \langle \cdot, \cdot \rangle_{\mathbf{A}} : \mathbb{K}^n \times \mathbb{K}^n \rightarrow \mathbb{K}, \quad (\vec{x}_1, \vec{x}_2) \mapsto \langle \vec{x}_1, \vec{x}_2 \rangle_{\mathbf{A}} = \vec{x}_1^T \mathbf{A} \vec{x}_2.$$

Notice that $\mathbf{A} \vec{x}_2 \in M_{n,1}(\mathbb{K})$ and $\vec{x}_1^T \in M_{1,n}(\mathbb{K})$ so that $\vec{x}_1^T \mathbf{A} \vec{x}_2 \in M_{1,1}(\mathbb{K}) = \mathbb{K}$. The properties of the transpose and matrix multiplication imply that $\langle \cdot, \cdot \rangle_{\mathbf{A}}$ is indeed a bilinear form on $\mathbb{K}^n$. Also, observe that the standard scalar product on $\mathbb{R}^n$ arises by taking $\mathbf{A}$ to be the identity matrix $\mathbf{1}_n \in M_{n,n}(\mathbb{R})$. That is, for all $\vec{x}, \vec{y} \in \mathbb{R}^n$, we have $\vec{x} \cdot \vec{y} = \langle \vec{x}, \vec{y} \rangle_{\mathbf{1}_n}$.

(iii) Using the determinant of a 2×2 -matrix, we obtain a map

$$\langle \cdot, \cdot \rangle : \mathbb{K}_2 \times \mathbb{K}_2 \rightarrow \mathbb{K} \quad (\vec{\xi}_1, \vec{\xi}_2) \mapsto \langle \vec{\xi}_1, \vec{\xi}_2 \rangle = \det \begin{pmatrix} \vec{\xi}_1 \\ \vec{\xi}_2 \end{pmatrix}.$$

The properties of the determinant then imply that $\langle \cdot, \cdot \rangle$ is an alternating bilinear form on the $\mathbb{K}$ -vector space $\mathbb{K}_2$.

(iv) For $n \in \mathbb{N}$ we consider $V = M_{n,n}(\mathbb{K})$, the $\mathbb{K}$ -vector space of $n \times n$ -matrices with entries in $\mathbb{K}$. We define $\langle \cdot, \cdot \rangle : M_{n,n}(\mathbb{K}) \times M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ by the rule

$$(9.3) \quad (\mathbf{A}, \mathbf{B}) \mapsto \langle \mathbf{A}, \mathbf{B} \rangle = \text{Tr}(\mathbf{AB}).$$

Definition 6.17 implies that

$$\text{Tr}(s_1 \mathbf{A}_1 + s_2 \mathbf{A}_2) = s_1 \text{Tr}(\mathbf{A}_1) + s_2 \text{Tr}(\mathbf{A}_2)$$

for all $s_1, s_2 \in \mathbb{K}$ and all $\mathbf{A}_1, \mathbf{A}_2 \in M_{n,n}(\mathbb{K})$, that is, the trace is a linear map from $M_{n,n}(\mathbb{K})$ into $\mathbb{K}$. Hence we obtain for all $s_1, s_2 \in \mathbb{K}$ and all $\mathbf{A}_1, \mathbf{A}_2, \mathbf{B} \in M_{n,n}(\mathbb{K})$

$$\begin{aligned} \langle s_1 \mathbf{A}_1 + s_2 \mathbf{A}_2, \mathbf{B} \rangle &= \text{Tr}((s_1 \mathbf{A}_1 + s_2 \mathbf{A}_2) \mathbf{B}) = s_1 \text{Tr}(\mathbf{A}_1 \mathbf{B}) + s_2 \text{Tr}(\mathbf{A}_2 \mathbf{B}) \\ &= s_1 \langle \mathbf{A}_1, \mathbf{B} \rangle + s_2 \langle \mathbf{A}_2, \mathbf{B} \rangle. \end{aligned}$$

showing that $\langle \cdot, \cdot \rangle$ is linear in the first argument. **Proposition 6.19** implies that $\langle \mathbf{A}, \mathbf{B} \rangle = \langle \mathbf{B}, \mathbf{A} \rangle$ for all $\mathbf{A}, \mathbf{B} \in M_{n,n}(\mathbb{K})$, hence $\langle \cdot, \cdot \rangle$ is symmetric and therefore also linear in the second variable. We conclude that (9.3) defines a symmetric bilinear form on the vector space $M_{n,n}(\mathbb{K})$.

(v) We consider $V = P(\mathbb{K})$, the $\mathbb{K}$ -vector space of polynomials. For some fixed scalar $x_0 \in \mathbb{K}$ we may define

$$\langle \cdot, \cdot \rangle : P(\mathbb{K}) \times P(\mathbb{K}) \rightarrow \mathbb{K}, \quad (p, q) \mapsto \langle p, q \rangle = p(x_0)q(x_0).$$

Then we have for all $s_1, s_2 \in \mathbb{K}$ and polynomials $p_1, p_2, q \in P(\mathbb{K})$

$$\begin{aligned} \langle s_1 \cdot_{P(\mathbb{K})} p_1 +_{P(\mathbb{K})} s_2 \cdot_{P(\mathbb{K})} p_2, q \rangle &= (s_1 \cdot_{P(\mathbb{K})} p_1 +_{P(\mathbb{K})} s_2 \cdot_{P(\mathbb{K})} p_2)(x_0)q(x_0) \\ &= (s_1 p_1(x_0) + s_2 p_2(x_0))q(x_0) \\ &= s_1 p_1(x_0)q(x_0) + s_2 p_2(x_0)q(x_0) \\ &= s_1 \langle p_1, q \rangle + s_2 \langle p_2, q \rangle. \end{aligned}$$

Hence $\langle \cdot, \cdot \rangle$ is linear in the first variable. Clearly $\langle \cdot, \cdot \rangle$ is also symmetric and therefore defines a symmetric bilinear form on $V = P(\mathbb{K})$.

(vi) We consider $V = C([-1, 1], \mathbb{R})$, the $\mathbb{R}$ -vector space of continuous real-valued functions defined on the interval $[-1, 1]$. Recall from M03 Analysis I that continuous functions are integrable, hence we can define

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}, \quad (f, g) \mapsto \langle f, g \rangle = \int_{-1}^1 f(x)g(x)dx.$$

The properties of integration imply that this defines a symmetric bilinear form on $C([-1, 1], \mathbb{R})$.

Recall that the choice of an ordered basis of a finite dimensional $\mathbb{K}$ -vector space V allowed to associate a matrix to every endomorphism $f : V \rightarrow V$. Similarly, an ordered basis also allows to associate a matrix to a bilinear form $\langle \cdot, \cdot \rangle$ on V .

Definition 9.3 (Matrix representation of a bilinear form) Let V be a finite dimensional $\mathbb{K}$ -vector space, $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V and $\langle \cdot, \cdot \rangle$ a bilinear form on V . The *matrix representation of $\langle \cdot, \cdot \rangle$ with respect to $\mathbf{b}$* is the matrix $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$

satisfying

$$\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) = (\langle v_i, v_j \rangle)_{1 \leq i, j \leq n}.$$

Remark 9.4 ($\heartsuit$ – not examinable) Let $\text{Bil}(V)$ denote the set of bilinear forms on some $\mathbb{K}$ -vector space V . By definition, $\text{Bil}(V)$ is a subset of the vector space of functions from $V \times V$ into $\mathbb{K}$. By [Definition 3.21](#), it follows that $\text{Bil}(V)$ is itself a $\mathbb{K}$ -vector space. Moreover, if $\dim V = n \in \mathbb{N}$ and V is equipped with an ordered basis $\mathbf{b}$, the mapping from $\text{Bil}(V)$ into $M_{n,n}(\mathbb{K})$ which sends a bilinear form to its matrix representation with respect to $\mathbf{b}$

$$\langle \cdot, \cdot \rangle \mapsto \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$$

is an isomorphism. In particular, $\dim \text{Bil}(V) = n^2$. The proof is left to the interested reader.

Example 9.5

- (i) Let $\langle \cdot, \cdot \rangle$ denote the standard scalar product on $\mathbb{R}^n$ and $\mathbf{e} = (\vec{e}_1, \dots, \vec{e}_n)$ the standard basis of $\mathbb{K}^n$. Then, one easily computes that

$$\langle \vec{e}_i, \vec{e}_j \rangle = \vec{e}_i^T \vec{e}_j = \delta_{ij}$$

and hence $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{e}) = (\delta_{ij})_{1 \leq i, j \leq n} = \mathbf{1}_n$.

- (ii) Likewise, if $\mathbf{A} \in M_{n,n}(\mathbb{K})$, then $\mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{e}) = \mathbf{A}$. Indeed, writing $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n}$, we have

$$\mathbf{A} \vec{e}_j = \sum_{k=1}^n A_{kj} \vec{e}_k$$

and thus

$$\langle \vec{e}_i, \vec{e}_j \rangle_{\mathbf{A}} = \vec{e}_i^T \mathbf{A} \vec{e}_j = \vec{e}_i^T \sum_{k=1}^n A_{kj} \vec{e}_k = \sum_{k=1}^n A_{kj} \vec{e}_i^T \vec{e}_k = \sum_{k=1}^n A_{kj} \delta_{ik} = A_{ij}.$$

- (iii) Let $\langle \cdot, \cdot \rangle$ denote the alternating bilinear form on $\mathbb{K}_2$ from [Example 9.2](#) above and

$$\mathbf{b} = ((1 \ 0), (0 \ 1)).$$

The alternating property of $\langle \cdot, \cdot \rangle$ implies that the diagonal entries of $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$ vanish. Hence we obtain

$$\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) = \begin{pmatrix} 0 & \det \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \\ \det \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} & 0 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}.$$

Proposition 9.6 Let V be a finite dimensional $\mathbb{K}$ -vector space, $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V with associated linear coordinate system $\beta : V \rightarrow \mathbb{K}^n$ and $\langle \cdot, \cdot \rangle$ a bilinear form on V . Then

- (i) for all $w_1, w_2 \in V$ we have

$$\langle w_1, w_2 \rangle = (\beta(w_1))^T \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) \beta(w_2).$$

- (ii) $\langle \cdot, \cdot \rangle$ is symmetric if and only if $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$ is symmetric;

(iii) if $\mathbf{b}'$ is another ordered basis of V , then

$$\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}') = \mathbf{C}^T \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) \mathbf{C},$$

where $\mathbf{C} = \mathbf{C}(\mathbf{b}', \mathbf{b})$ denotes the change of basis matrix, see [Definition 3.104](#).

Proof (i) Since $\mathbf{b}$ is a basis of V , it follows that for all $w_1, w_2 \in V$ there exist unique scalars $s_1, \dots, s_n$ and $t_1, \dots, t_n$ so that

$$w_1 = \sum_{i=1}^n s_i v_i \quad \text{and} \quad w_2 = \sum_{i=1}^n t_i v_i.$$

Recall that this means that

$$\beta(w_1) = \begin{pmatrix} s_1 \\ \vdots \\ s_n \end{pmatrix} \quad \text{and} \quad \beta(w_2) = \begin{pmatrix} t_1 \\ \vdots \\ t_n \end{pmatrix}.$$

Using the bilinearity of $\langle \cdot, \cdot \rangle$, this gives

$$\begin{aligned} \langle w_1, w_2 \rangle &= \left\langle \sum_{i=1}^n s_i v_i, \sum_{j=1}^n t_j v_j \right\rangle = \sum_{i=1}^n s_i \sum_{j=1}^n t_j \langle v_i, v_j \rangle \\ &= \sum_{i=1}^n s_i \sum_{j=1}^n [\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})]_{ij} t_j = (\beta(w_1))^T \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) \beta(w_2). \end{aligned}$$

(ii) Suppose $\langle \cdot, \cdot \rangle$ is symmetric. Then for all $1 \leq i, j \leq n$, we have

$$[\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})]_{ij} = \langle v_i, v_j \rangle = \langle v_j, v_i \rangle = [\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})]_{ji}$$

so that $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$ is symmetric. Conversely, suppose $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$ is symmetric. Using notation as in (i), we obtain for all $w_1, w_2 \in V$

$$\begin{aligned} \langle w_1, w_2 \rangle &= \sum_{i=1}^n \sum_{j=1}^n s_i [\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})]_{ij} t_j = \sum_{j=1}^n \sum_{i=1}^n t_j [\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})]_{ji} s_i \\ &= \langle w_2, w_1 \rangle \end{aligned}$$

so that $\langle \cdot, \cdot \rangle$ is symmetric as well.

(iii) Let $\mathbf{b}' = (v'_1, \dots, v'_n)$ be another ordered basis of V . Since $\mathbf{b}$ is a basis of V there exist unique scalars C_{ij} , $1 \leq i, j \leq n$ such that

$$v'_j = \sum_{i=1}^n C_{ij} v_i$$

and, by [Definition 3.104](#), we have $\mathbf{C}(\mathbf{b}', \mathbf{b}) = (C_{ij})_{1 \leq i, j \leq n}$. Writing $\mathbf{C} = (C_{ij})_{1 \leq i, j \leq n}$ and using the bilinearity of $\langle \cdot, \cdot \rangle$, we calculate

$$\begin{aligned} [\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}')]_{ij} &= \langle v'_i, v'_j \rangle = \left\langle \sum_{k=1}^n C_{ki} v_k, \sum_{l=1}^n C_{lj} v_l \right\rangle = \sum_{k=1}^n \sum_{l=1}^n C_{ki} C_{lj} \langle v_k, v_l \rangle \\ &= \sum_{k=1}^n C_{ki} \sum_{l=1}^n \langle v_k, v_l \rangle C_{lj} = \sum_{k=1}^n C_{ki} [\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) C]_{kj} \\ &= \sum_{k=1}^n [C^T]_{ik} [\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) C]_{kj} = [C^T \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) C]_{ij} \end{aligned}$$

as claimed. □

Example 9.7 We consider the symmetric bilinear form $\langle \cdot, \cdot \rangle_{\mathbf{A}}$ on $\mathbb{R}^2$ arising from the matrix

$$\mathbf{A} = \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix}.$$

via the rule (9.2). Let $\mathbf{e} = (\vec{e}_1, \vec{e}_2)$ denote the ordered standard basis of $\mathbb{R}^2$ and $\mathbf{b} = (\vec{e}_1 + \vec{e}_2, \vec{e}_2 - \vec{e}_1)$. In Example 9.5 we have seen that $\mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{e}) = \mathbf{A}$. In Example 3.106 we computed that

$$\mathbf{C}(\mathbf{b}, \mathbf{e}) = \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix}.$$

By definition, we have

$$[\mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{b})]_{11} = (1 \ 1) \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} \begin{pmatrix} 1 \\ 1 \end{pmatrix} = 12,$$

$$[\mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{b})]_{12} = (1 \ 1) \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} \begin{pmatrix} -1 \\ 1 \end{pmatrix} = 0,$$

$$[\mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{b})]_{22} = (-1 \ 1) \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} \begin{pmatrix} -1 \\ 1 \end{pmatrix} = 8,$$

so that

$$\mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{b}) = \begin{pmatrix} 12 & 0 \\ 0 & 8 \end{pmatrix}.$$

Indeed, writing $\mathbf{C} = \mathbf{C}(\mathbf{b}, \mathbf{e})$, we have

$$\mathbf{C}^T \mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{e}) \mathbf{C} = \begin{pmatrix} 1 & 1 \\ -1 & 1 \end{pmatrix} \begin{pmatrix} 5 & 1 \\ 1 & 5 \end{pmatrix} \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix} = \begin{pmatrix} 12 & 0 \\ 0 & 8 \end{pmatrix} = \mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{b}),$$

in agreement with Proposition 9.6.

Remark 9.8 You may remember from school that two non-zero vectors $\vec{x}_1, \vec{x}_2 \in \mathbb{R}^n$ are *perpendicular* if and only if $\vec{x}_1 \cdot \vec{x}_2 = 0$. In particular, no non-zero vector in $\mathbb{R}^n$ is perpendicular to all vectors, or phrased differently, if $\vec{x} \cdot \vec{x}_0 = 0$ for all vectors $\vec{x}$, then $\vec{x}_0 = 0_{\mathbb{R}^n}$.

This condition also makes sense for a bilinear form:

Definition 9.9 (Non-degenerate bilinear form) Let $\langle \cdot, \cdot \rangle$ be a bilinear form on a finite dimensional $\mathbb{K}$ -vector space V . Then $\langle \cdot, \cdot \rangle$ is called *non-degenerate*, if whenever a vector $v_0 \in V$ satisfies $\langle v, v_0 \rangle = 0$ for all vectors $v \in V$, then we must have $v_0 = 0_V$.

Non-degeneracy of a bilinear form $\langle \cdot, \cdot \rangle$ can be characterized in terms of its matrix representation, more precisely:

Proposition 9.10 Let $\langle \cdot, \cdot \rangle$ be a bilinear form on a finite dimensional $\mathbb{K}$ -vector space V and $\mathbf{b}$ an ordered basis of V . Then $\langle \cdot, \cdot \rangle$ is non-degenerate if and only if $\det \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) \neq 0$.

Proof Let $n = \dim V$. First observe that a vector $\vec{y} \in \mathbb{K}^n$ satisfies $\vec{x}^T \vec{y} = 0$ for all $\vec{x} \in \mathbb{K}^n$ if and only if $\vec{y} = 0_{\mathbb{K}^n}$.

The statement of the proposition is equivalent to the statement that $\det \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) = 0$ if and only there exists a non-zero vector $v_0 \in V$ so that $\langle v, v_0 \rangle = 0$ for all $v \in V$. We write $\mathbf{A} = \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$. By [Proposition 6.22](#), $\det \mathbf{A} = 0$ is equivalent to the mapping $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$ not being injective and hence by [Lemma 3.31](#) equivalent to the existence of a non-zero vector $\vec{x}_0 \in \mathbb{K}^n$ with $\mathbf{A}\vec{x}_0 = 0_{\mathbb{K}^n}$. Let $v_0 \in V$ be the non-zero vector whose coordinate representation is $\vec{x}_0$, that is, $\beta(v_0) = \vec{x}_0$, where $\beta : V \rightarrow \mathbb{K}^n$ denotes the linear coordinate system associated to $\mathbf{b}$. By [Proposition 9.6](#) we have for all $v \in V$

$$(9.4) \quad \langle v, v_0 \rangle = (\beta(v))^T \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) \beta(v_0) = (\beta(v))^T \mathbf{A} \vec{x}_0.$$

Writing $\vec{y} = \mathbf{A} \vec{x}_0$, the observation at the beginning of the proof shows that (9.4) is 0 for all $v \in V$ if and only if $\mathbf{A} \vec{x}_0 = 0_{\mathbb{K}^n}$. $\square$

Exercises

Exercise 9.11 We consider $V = M_{2,2}(\mathbb{R})$ and define

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}, \quad (\mathbf{A}, \mathbf{B}) \mapsto \langle \mathbf{A}, \mathbf{B} \rangle = \frac{1}{4} (\det(\mathbf{A} + \mathbf{B}) - \det(\mathbf{A} - \mathbf{B})).$$

Show that $\langle \cdot, \cdot \rangle$ defines a symmetric bilinear form on $V = M_{2,2}(\mathbb{R})$.

9.2 Symmetric bilinear forms

We now restrict to the case $\mathbb{K} = \mathbb{R}$. Perpendicular vectors are orthogonal in the following sense:

Definition 9.12 (Orthogonal vectors) Let V be an $\mathbb{R}$ -vector space equipped with a symmetric bilinear form $\langle \cdot, \cdot \rangle$. Two vectors $v_1, v_2 \in V$ are called *orthogonal with respect to $\langle \cdot, \cdot \rangle$* if $\langle v_1, v_2 \rangle = 0$. We write $v_1 \perp v_2$ if the vectors $v_1, v_2 \in V$ are orthogonal. A subset $S \subset V$ is called *orthogonal with respect to $\langle \cdot, \cdot \rangle$* if all pairs of distinct vectors of S are orthogonal with respect to $\langle \cdot, \cdot \rangle$. A basis of V which is also an orthogonal subset is called an *orthogonal basis*.

Example 9.13

- (i) Perpendicular vectors in $\mathbb{R}^n$ are orthogonal with respect to the standard scalar product defined by the rule (9.1).
- (ii) Example 9.7 continued: As we computed above, the vectors $\vec{v}_1 = \vec{e}_1 + \vec{e}_2$ and $\vec{v}_2 = \vec{e}_2 - \vec{e}_1$ satisfy $\langle \vec{v}_1, \vec{v}_2 \rangle_A = 0$ and hence are orthogonal with respect to $\langle \cdot, \cdot \rangle_A$.
- (iii) Example 9.2 (vi) continued: Let $f_1 \in V$ be the function $x \mapsto x$ and $f_3 \in V$ be the function $x \mapsto \frac{1}{2}(5x^3 - 3x)$. Then

$$\langle f_1, f_3 \rangle = \int_{-1}^1 x \frac{1}{2}(5x^3 - 3x) dx = \frac{1}{2} (x^5 - x^3) \Big|_{-1}^1 = 0,$$

so that f_1 and f_3 are orthogonal with respect to $\langle \cdot, \cdot \rangle$.

Definition 9.14 (Orthonormal vectors) Let V be an $\mathbb{R}$ -vector space equipped with a symmetric bilinear form $\langle \cdot, \cdot \rangle$. A subset $S \subset V$ is called *orthonormal with respect to $\langle \cdot, \cdot \rangle$* if S is orthogonal with respect to $\langle \cdot, \cdot \rangle$ and if for all vectors $v \in S$ we have $\langle v, v \rangle = 1$. A basis of V which is also a orthonormal subset is called an *orthonormal basis*.

Remark 9.15

- Often when $\langle \cdot, \cdot \rangle$ is clear from the context we will simply speak of orthogonal or orthonormal vectors without explicitly mentioning $\langle \cdot, \cdot \rangle$.
- Notice that an ordered basis $\mathbf{b}$ of V is orthonormal with respect to $\langle \cdot, \cdot \rangle$ if and only if

$$\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) = \mathbf{1}_n,$$

where $n = \dim V$.

Example 9.16

- (i) The standard basis $\{\vec{e}_1, \dots, \vec{e}_n\}$ of $\mathbb{R}^n$ satisfies

$$\vec{e}_i \cdot \vec{e}_j = \delta_{ij}$$

and hence is a orthonormal basis with respect to the standard scalar product on $\mathbb{R}^n$.

(ii) **Example 9.2** (vi) continued: Let $S = \{f_1, f_2, f_3\} \subset C([-1, 1], \mathbb{R})$ be the subset defined by the functions

$$f_1 : x \mapsto \sqrt{\frac{3}{2}}x, \quad f_2 : x \mapsto \frac{1}{2}\sqrt{\frac{5}{2}}(3x^2 - 1), \quad f_3 : x \mapsto \frac{1}{2}\sqrt{\frac{7}{2}}(5x^3 - 3x).$$

Then S is orthonormal with respect to $\langle \cdot, \cdot \rangle$ as can be verified by direct computation.

Given a subspace $U \subset V$, its orthogonal subspace consists of all vectors in V that are orthogonal to all vectors of U .

Definition 9.17 (Orthogonal subspace) Let V be an $\mathbb{R}$ -vector space equipped with a symmetric bilinear form $\langle \cdot, \cdot \rangle$ and $U \subset V$ a subspace. The set

$$U^\perp = \{v \in V \mid \langle v, u \rangle = 0 \ \forall u \in U\}$$

is called the *orthogonal subspace to U* .

Remark 9.18

- It is common to write $\langle v, U \rangle = 0$ instead of $\langle v, u \rangle = 0 \ \forall u \in U$.
- Notice that the orthogonal subspace is indeed a subspace. The bilinearity of $\langle \cdot, \cdot \rangle$ implies that $\langle 0_V, u \rangle = 0$ for all $u \in U$, hence $0_V \in U^\perp$ and $U^\perp$ is non-empty. Moreover, if $v_1, v_2 \in U^\perp$, then we have for all $u \in U$ and all $s_1, s_2 \in \mathbb{R}$

$$\langle s_1 v_1 + s_2 v_2, u \rangle = s_1 \langle v_1, u \rangle + s_2 \langle v_2, u \rangle = 0$$

where we use the bilinearity of $\langle \cdot, \cdot \rangle$ and that $v_1, v_2 \in U^\perp$. By [Definition 3.21](#) it follows that $U^\perp$ is indeed a subspace.

- Notice also that a symmetric bilinear form $\langle \cdot, \cdot \rangle$ on V is non-degenerate if and only if $V^\perp = \{0_V\}$.

Example 9.19

(i) Let $\mathbb{R}^3$ be equipped with the standard scalar product. If U is a line through the origin in $\mathbb{R}^3$, then $U^\perp$ consists of the plane through the origin that is perpendicular to U , see [Figure 9.1](#).

(ii) **Example 9.2** (iv) continued. Let $U = \{s\mathbf{1}_n \mid s \in \mathbb{R}\}$ then

$$U^\perp = \{\mathbf{A} \in M_{n,n}(\mathbb{R}) \mid \text{Tr}(\mathbf{A}s\mathbf{1}_n) = 0 \ \forall s \in \mathbb{R}\}.$$

Since $\text{Tr}(\mathbf{A}s\mathbf{1}_n) = s \text{Tr}(\mathbf{A}\mathbf{1}_n) = s \text{Tr}(\mathbf{A})$, we conclude that the orthogonal subspace to U consists of the matrices whose trace is zero

$$U^\perp = \{\mathbf{A} \in M_{n,n}(\mathbb{R}) \mid \text{Tr}(\mathbf{A}) = 0\}.$$

Previously in [Corollary 3.65](#) we saw that every finite dimensional vector space V admits a basis. We can now upgrade this fact in the case where V is equipped with a symmetric bilinear form:

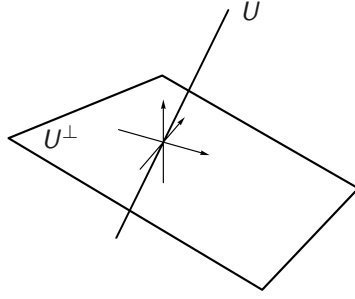


FIGURE 9.1. The orthogonal complement of a line through the origin.

Theorem 9.20 (Existence of an orthogonal basis) *Let V be a finite dimensional $\mathbb{R}$ -vector space equipped with a symmetric bilinear form $\langle \cdot, \cdot \rangle$. Then V admits an orthogonal basis with respect to $\langle \cdot, \cdot \rangle$.*

For the proof of [Theorem 9.20](#) we need two lemmas.

Lemma 9.21 *Let V be an $\mathbb{R}$ -vector space and $\langle \cdot, \cdot \rangle$ a symmetric bilinear form on V . Suppose there exist vectors $v_1, v_2 \in V$ such that $\langle v_1, v_2 \rangle \neq 0$. Then there exists a vector $v \in V$ with $\langle v, v \rangle \neq 0$.*

Proof If $\langle v_1, v_1 \rangle \neq 0$ or $\langle v_2, v_2 \rangle \neq 0$ we are done, hence assume $\langle v_1, v_1 \rangle = \langle v_2, v_2 \rangle = 0$. Let $v = v_1 + v_2$, then we obtain

$$\langle v, v \rangle = \langle v_1 + v_2, v_1 + v_2 \rangle = \langle v_1, v_1 \rangle + 2\langle v_1, v_2 \rangle + \langle v_2, v_2 \rangle = 2\langle v_1, v_2 \rangle.$$

By assumption we have $\langle v_1, v_2 \rangle \neq 0$ and hence also $\langle v, v \rangle \neq 0$. □

Lemma 9.22 *Let V be a finite dimensional $\mathbb{R}$ -vector space equipped with a symmetric bilinear form $\langle \cdot, \cdot \rangle$. Suppose $v \in V$ satisfies $\langle v, v \rangle \neq 0$, then $V = U \oplus U^\perp$ where $U = \{sv \mid s \in \mathbb{R}\}$.*

Proof Applying [Remark 6.7](#), we need to show that $U \cap U^\perp = \{0_V\}$ and that $U + U^\perp = V$.

We first show that $U \cap U^\perp = \{0_V\}$. Suppose $u \in U$ and $u \in U^\perp$. Since $u \in U$ we have $u = sv$ for some scalar s . Since $u \in U^\perp$ we must also have $0 = \langle u, v \rangle = s\langle v, v \rangle$. Since $\langle v, v \rangle \neq 0$, this implies $s = 0$ and hence $u = 0_V$.

We next show that $U + U^\perp = V$. Let $w \in V$. We want to write $w = sv + \hat{v}$ for some $\hat{v}$ satisfying $\langle \hat{v}, v \rangle = 0$. Since $\hat{v} = w - sv$, this condition becomes

$$0 = \langle v, w - sv \rangle = \langle v, w \rangle - s\langle v, v \rangle$$

and since $\langle v, v \rangle \neq 0$, this gives $s = \frac{\langle v, w \rangle}{\langle v, v \rangle}$. Taking

$$\hat{v} = w - \frac{\langle v, w \rangle}{\langle v, v \rangle} v$$

thus gives $w = sv + \hat{v}$. □

Proof of Theorem 9.20 Let $n = \dim V$. Suppose $\langle \cdot, \cdot \rangle$ is degenerate and consider $V^\perp$. By [Corollary 6.11](#) there exists a subspace $V' \subset V$ such that $V = V^\perp \oplus V'$. By construction,

the restriction of $\langle \cdot, \cdot \rangle$ onto V' is non-degenerate. If $v_1, \dots, v_m$ is an orthogonal basis of V' and $v_{m+1}, \dots, v_n$ a basis of $V^\perp$, then $\{v_1, \dots, v_n\}$ is an orthogonal basis of V , since the vectors $v_{m+1}, \dots, v_n$ are orthogonal to all vectors of V .

It is thus sufficient to prove the existence of an orthogonal basis for the case when $\langle \cdot, \cdot \rangle$ is non-degenerate.

Let us therefore assume that $\langle \cdot, \cdot \rangle$ is non-degenerate on V . We are going to prove the statement by using induction on the dimension of the vector space. If $\dim V = 0$ there is nothing to show, hence the statement is anchored. We will argue next that if every $(n-1)$ -dimensional $\mathbb{R}$ -vector space equipped with a non-degenerate symmetric bilinear form admits an orthogonal basis, then so does every n -dimensional $\mathbb{R}$ -vector space equipped with a non-degenerate symmetric bilinear form.

Let $v_1 \in V$ be any non-zero vector. Since $\langle \cdot, \cdot \rangle$ is non-degenerate v_1 cannot be orthogonal to all vectors of V and hence there exists a vector $v_2 \in V$ such that $\langle v_1, v_2 \rangle \neq 0$. Therefore, by [Lemma 9.21](#) there exists a non-zero vector $v \in V$ with $\langle v, v \rangle \neq 0$. Writing $U = \{sv \mid s \in \mathbb{R}\}$, we have that $V = U \oplus U^\perp$ by [Lemma 9.22](#). Since $\dim U = 1$, we must have $\dim U^\perp = n-1$ by [Proposition 6.12](#). The restriction of $\langle \cdot, \cdot \rangle$ onto $U^\perp$ is non-degenerate. Indeed, if there were a vector in $U^\perp$ which is orthogonal to all vectors in $U^\perp$, then – since it lies in $U^\perp$ – it is also orthogonal to all vectors of U and hence to all vectors of V . This contradicts the assumption that $\langle \cdot, \cdot \rangle$ is non-degenerate on V . Since the restriction of $\langle \cdot, \cdot \rangle$ on $U^\perp$ is non-degenerate and $\dim U^\perp = n-1$, the induction hypothesis implies that there exists a basis $\{w_2, \dots, w_n\}$ of $U^\perp$ which is orthogonal with respect to $\langle \cdot, \cdot \rangle$. Setting $w_1 = v$ gives an orthogonal basis $\{w_1, w_2, \dots, w_n\}$ of V . $\square$

We also have:

Lemma 9.23 *Let V be a finite dimensional $\mathbb{R}$ -vector space equipped with a symmetric bilinear form $\langle \cdot, \cdot \rangle$. Furthermore, let $U \subset V$ be a subspace and $\{u_1, \dots, u_k\}$ be a basis of U . Then the following two statements are equivalent*

- (i) *a vector $v \in V$ is an element of $U^\perp$;*
- (ii) *for $1 \leq i \leq k$ we have $\langle v, u_i \rangle = 0$.*

Proof Exercise. $\square$

As a corollary to [Theorem 9.20](#) we obtain a generalisation of [Lemma 9.22](#).

Corollary 9.24 *Let V be a finite dimensional $\mathbb{R}$ -vector space and $\langle \cdot, \cdot \rangle$ a symmetric bilinear form on V . Suppose $U \subset V$ is a subspace such that the restriction of $\langle \cdot, \cdot \rangle$ to U is non-degenerate. Then U and $U^\perp$ are in direct sum and we have*

$$V = U \oplus U^\perp.$$

Proof The proof is similar to [Lemma 9.22](#). We first show that $U \cap U^\perp = \{0_V\}$. Suppose $u_0 \in U \cap U^\perp$. Recall that

$$U^\perp = \{v \in V \mid \langle v, u \rangle = 0 \ \forall u \in U\}$$

Since $u_0 \in U^\perp$ we thus have for all $u \in U$

$$\langle u_0, u \rangle = 0.$$

Since the restriction of $\langle \cdot, \cdot \rangle$ to U is non-degenerate, this implies that $u_0 = 0_V$, hence $U \cap U^\perp = \{0_V\}$.

We next show that $U + U^\perp = V$. By [Theorem 9.20](#), the subspace U admits an ordered basis $\mathbf{b} = (v_1, \dots, v_k)$ that is orthogonal with respect to $\langle \cdot, \cdot \rangle$, that is, $\langle v_i, v_j \rangle = 0$ for $i \neq j$. In particular, the matrix representation of $\langle \cdot, \cdot \rangle$ with respect to $\mathbf{b}$ is diagonal and the diagonal entries are given by $\langle v_i, v_i \rangle$ for $1 \leq i \leq k$. By [Proposition 5.24](#) we have

$$\det \mathbf{M}(\langle \cdot, \cdot \rangle|_U, \mathbf{b}) = \prod_{i=1}^k \langle v_i, v_i \rangle$$

where $\langle \cdot, \cdot \rangle|_U$ denotes the restriction of $\langle \cdot, \cdot \rangle$ onto $U \times U$. Since $\langle \cdot, \cdot \rangle|_U$ is non-degenerate, we have $\det \mathbf{M}(\langle \cdot, \cdot \rangle|_U, \mathbf{b}) \neq 0$ by [Proposition 9.10](#), hence $\langle v_i, v_i \rangle \neq 0$ for $1 \leq i \leq k$.

Finally, we argue that any vector $w \in V$ can be written as $w = \hat{v} + \sum_{i=1}^k s_i v_i$ for a suitable vector $\hat{v} \in U^\perp$ and scalars s_i . As in the proof of [Lemma 9.22](#), we define

$$s_i = \frac{\langle v_i, w \rangle}{\langle v_i, v_i \rangle}$$

and $\hat{v} = w - \sum_{i=1}^k s_i v_i$. Then $w = \hat{v} + \sum_{i=1}^k s_i v_i$ and moreover $\langle \hat{v}, v_i \rangle = 0$ for $1 \leq i \leq k$, since $\langle v_i, v_j \rangle = 0$ for $i \neq j$. Since $\mathbf{b}$ is a basis of U [Lemma 9.22](#) implies that $\hat{v}$ is an element of $U^\perp$. $\square$

Remark 9.25 In the case where the restriction of a symmetric bilinear form to a subspace U is non-degenerate, we have seen that $U^\perp$ is a complement to U . The subspace $U^\perp$ is called the *orthogonal complement* of U .

The process of scaling a vector v so that $\langle v, v \rangle$ equals some specific value – typically 1 – is known as *normalising the vector*.

Remark 9.26 (Normalisation) By definition, the matrix representation of a symmetric bilinear form $\langle \cdot, \cdot \rangle$ with respect to an ordered orthogonal basis $\mathbf{b} = (v_1, \dots, v_n)$ of V is diagonal. Notice that if we define

$$v'_i = \begin{cases} v_i, & \langle v_i, v_i \rangle = 0 \\ \frac{v_i}{\sqrt{|\langle v_i, v_i \rangle|}}, & \langle v_i, v_i \rangle \neq 0 \end{cases}$$

for $1 \leq i \leq n$, then $\mathbf{b}' = (v'_1, \dots, v'_n)$ is also an ordered basis of V and either $\langle v'_i, v'_i \rangle = 0$ or

$$\langle v'_i, v'_i \rangle = \left\langle \frac{v_i}{\sqrt{|\langle v_i, v_i \rangle|}}, \frac{v_i}{\sqrt{|\langle v_i, v_i \rangle|}} \right\rangle = \frac{\langle v_i, v_i \rangle}{|\langle v_i, v_i \rangle|} = \pm 1.$$

Therefore, the matrix representation of $\langle \cdot, \cdot \rangle$ with respect to $\mathbf{b}'$ is diagonal as well and the diagonal entries are elements of the set $\{-1, 0, 1\}$.

This observation allows to reformulate [Theorem 9.20](#):

Theorem 9.27 (Matrix version of [Theorem 9.20](#)) Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{R})$ be a symmetric $n \times n$ -matrix. Then there exists an invertible $n \times n$ -matrix $\mathbf{C} \in \text{GL}(n, \mathbb{R})$

and integers p, q, s such that

$$(9.5) \quad \mathbf{C}^T \mathbf{A} \mathbf{C} = \begin{pmatrix} \mathbf{1}_p & & \\ & -\mathbf{1}_q & \\ & & \mathbf{0}_s \end{pmatrix}.$$

Proof Let $\mathbf{A} \in M_{n,n}(\mathbb{R})$ be a symmetric $n \times n$ -matrix and let $\langle \cdot, \cdot \rangle$ denote the symmetric bilinear form on $\mathbb{R}^n$ defined by the rule $\langle \vec{x}_1, \vec{x}_2 \rangle = \vec{x}_1^T \mathbf{A} \vec{x}_2$ for all $\vec{x}_1, \vec{x}_2 \in \mathbb{R}^n$. By [Example 9.5](#) (ii), we have that $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{e}) = \mathbf{A}$, where $\mathbf{e}$ denotes the standard ordered basis of $\mathbb{R}^n$. [Theorem 9.20](#) implies that $\mathbb{R}^n$ admits an orthogonal basis with respect to $\langle \cdot, \cdot \rangle$. After carrying out the normalisation procedure described in [Remark 9.26](#) and possibly renumbering the basis vectors, we thus obtain an ordered basis $\mathbf{b}$ of $\mathbb{R}^n$ such that

$$\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) = \begin{pmatrix} \mathbf{1}_p & & \\ & -\mathbf{1}_q & \\ & & \mathbf{0}_s \end{pmatrix}.$$

Defining $\mathbf{C} = \mathbf{C}(\mathbf{b}, \mathbf{e})$, [Proposition 9.6](#) thus implies that $\mathbf{C}^T \mathbf{A} \mathbf{C} = \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$ as claimed. Finally, the matrix $\mathbf{C}$ is invertible by [Remark 3.105](#). $\square$

Remark 9.28 (Sylvester’s law of inertia)

- Sylvester’s law of inertia states that the numbers p and q in (9.5) (and hence also s) are uniquely determined by the bilinear form $\langle \cdot, \cdot \rangle$. That is, they do not depend on the choice of matrix $\mathbf{C} \in \text{GL}(n, \mathbb{R})$ such that $\mathbf{C}^T \mathbf{A} \mathbf{C}$ is diagonal with diagonal entries from the set $\{-1, 0, 1\}$. We will not prove this fact, but a proof can be found in most textbooks about Linear Algebra.
- The pair (p, q) is known as the *signature of the bilinear form* $\langle \cdot, \cdot \rangle$.

Euclidean spaces

WEEK 4

10.1 Inner products

A symmetric bilinear form $\langle \cdot, \cdot \rangle$ on an $\mathbb{R}$ -vector space V allows to talk about vectors being orthogonal, but so far we have not defined the *length* of a vector or the *distance* between two vectors. In $\mathbb{R}^n$ equipped with the standard scalar product $\langle \cdot, \cdot \rangle$, the length of a vector $\vec{x} = (x_i)_{1 \leq i \leq n} \in \mathbb{R}^n$ is denoted by $\|\vec{x}\|$ and defined as

$$\|\vec{x}\| = \sqrt{\langle \vec{x}, \vec{x} \rangle} = \sqrt{\sum_{i=1}^n (x_i)^2}.$$

Over the real numbers we can only take square roots of non-negative numbers. Hence if we want to analogously define the length of vectors in an abstract vectors space V that is equipped with a bilinear form $\langle \cdot, \cdot \rangle$, then we need that $\langle v, v \rangle \geq 0$ for all vectors $v \in V$. This is known as positivity. Clearly, having a positive symmetric bilinear form on an $\mathbb{R}$ -vector space, we can define the length of vectors as in the case of $\mathbb{R}^n$ equipped with the standard scalar product. It might however still happen that there are vectors $v \in V$ different from the zero vector 0_V that satisfy $\langle v, v \rangle = 0$. Naturally, one might ask that the zero vector is the only vector with length zero. This leads to the notion of definiteness.

Definition 10.1 (Properties of bilinear forms) A bilinear form $\langle \cdot, \cdot \rangle$ on an $\mathbb{R}$ -vector space V is called

- *positive* if $\langle v, v \rangle \geq 0$ for all vectors $v \in V$;
- *definite* if $\langle v, v \rangle = 0$ if and only if $v = 0_V$.

Combining positivity, definiteness and symmetry, we arrive at the notion of an inner product:

Definition 10.2 (Inner product) Let V be an $\mathbb{R}$ -vector space. A bilinear form $\langle \cdot, \cdot \rangle$ on V that is positive definite and symmetric is called an *inner product*.

Remark 10.3 Notice that an inner product $\langle \cdot, \cdot \rangle$ on an $\mathbb{R}$ -vector space V is always a non-degenerate bilinear form. Indeed, if $v_0 \in V$ satisfies $\langle v, v_0 \rangle = 0$ for all vectors $v \in V$, then we also have $\langle v_0, v_0 \rangle = 0$ and hence $v_0 = 0_V$, since $\langle \cdot, \cdot \rangle$ is positive definite.

Example 10.4 (Inner products)

- (i) The standard scalar product on $\mathbb{R}^n$ defined by the rule (9.1) is indeed an inner product. Clearly, $\langle \cdot, \cdot \rangle$ is symmetric and from the Analysis I module we know that $y^2 \geq 0$ for all real numbers y and that $y^2 = 0$ if and only if $y = 0$. Since

$$\langle \vec{x}, \vec{x} \rangle = \sum_{i=1}^n (x_i)^2,$$

we conclude that $\langle \cdot, \cdot \rangle$ is positive definite and hence an inner product. The vector space $\mathbb{R}^n$ equipped with the standard scalar product is sometimes denoted by $\mathbb{E}^n$ (the letter E is to remind of the Greek Geometer Euclid).

- (ii) We consider $V = M_{3,3}(\mathbb{R})$ and let $U \subset V$ be the subspace consisting of anti-symmetric matrices. On U we define a symmetric bilinear form (notice the minus sign)

$$\langle \cdot, \cdot \rangle : U \times U \rightarrow \mathbb{R}, \quad (\mathbf{A}, \mathbf{B}) \mapsto \langle \mathbf{A}, \mathbf{B} \rangle = -\text{Tr}(\mathbf{AB}).$$

An element $\mathbf{A}$ of U satisfies $\mathbf{A}^T = -\mathbf{A}$ and hence can be written as

$$\mathbf{A} = \begin{pmatrix} 0 & x & y \\ -x & 0 & z \\ -y & -z & 0 \end{pmatrix}$$

for real numbers x, y, z . We obtain

$$\langle \mathbf{A}, \mathbf{A} \rangle = -\text{Tr} \begin{pmatrix} -x^2 - y^2 & -yz & xz \\ -yz & -x^2 - z^2 & -xy \\ xz & -xy & -y^2 - z^2 \end{pmatrix} = 2x^2 + 2y^2 + 2z^2.$$

We conclude that $\langle \mathbf{A}, \mathbf{A} \rangle \geq 0$ and $\langle \mathbf{A}, \mathbf{A} \rangle = 0$ if and only if $\mathbf{A} = \mathbf{0}_3$. Therefore, $\langle \cdot, \cdot \rangle$ is an inner product on U .

- (iii) Let $a < b$ be real numbers and consider $V = C([a, b], \mathbb{R})$, the $\mathbb{R}$ -vector space of continuous real-valued functions on the interval $[a, b]$. As in Example 9.2, (vi) we obtain a symmetric bilinear form on V via the definition

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}, \quad (f, g) \mapsto \langle f, g \rangle = \int_a^b f(x)g(x)dx.$$

The properties of integration from the Analysis module imply that $\langle \cdot, \cdot \rangle$ is also positive definite and hence an inner product.

Remark 10.5 (Naming convention) As we have seen, the standard scalar product on $\mathbb{R}^n$ is an example of an inner product. It is common to refer to inner products as scalar products as well. In these notes we will reserve the term scalar product for the standard scalar product on $\mathbb{R}^n$ and use inner product for a general positive definite symmetric bilinear form.

Notice that a symmetric bilinear form can be positive, but not positive definite:

Example 10.6 For any $x_0 \in \mathbb{R}$, the symmetric bilinear form on $V = P(\mathbb{R})$ defined by

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}, \quad (p, q) \mapsto p(x_0)q(x_0)$$

satisfies $\langle p, p \rangle = p(x_0)^2 \geq 0$ and hence is positive. It is however not an inner product. The polynomial f defined by the rule $x \mapsto f(x) = (x - x_0)$ for all $x \in \mathbb{R}$ is different from the zero polynomial $o : x \mapsto 0 \forall x \in \mathbb{R}$, but also satisfies $\langle f, f \rangle = 0$.

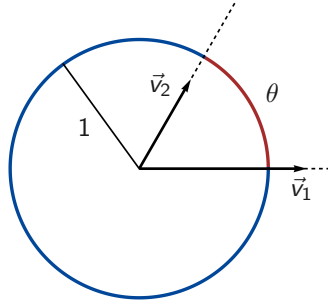


FIGURE 10.1. Angle between two vectors

An inner product $\langle \cdot, \cdot \rangle$ on an abstract $\mathbb{R}$ -vector space V allows to define geometric notions like length and distance on V .

Definition 10.7 (Euclidean space)

- (i) A pair $(V, \langle \cdot, \cdot \rangle)$ consisting of an $\mathbb{R}$ -vector space V and an inner product $\langle \cdot, \cdot \rangle$ on V is called a *Euclidean space*.
- (ii) The mapping $\| \cdot \| : V \rightarrow \mathbb{R}$ defined by the rule

$$v \mapsto \|v\| = \sqrt{\langle v, v \rangle}$$

for all $v \in V$ is called the *norm induced by $\langle \cdot, \cdot \rangle$* . Moreover, for any vector $v \in V$, the real number $\|v\|$ is called the *length of the vector v* .

- (iii) The mapping $d : V \times V \rightarrow \mathbb{R}$ defined by the rule

$$(v_1, v_2) \mapsto d(v_1, v_2) = \|v_1 - v_2\|$$

for all $v_1, v_2 \in V$ is called the *metric induced by $\langle \cdot, \cdot \rangle$* (or also metric induced by the norm $\| \cdot \|$). Furthermore, for any vectors $v_1, v_2 \in V$, the real number $d(v_1, v_2)$ is called the *distance from the vector v_1 to the vector v_2* .

Recall that the angle between two non-zero vectors $\vec{v}_1, \vec{v}_2 \in \mathbb{R}^2$ is defined as follows. The half lines spanned by $\vec{v}_1$ and $\vec{v}_2$ will each intersect the circle of radius 1 centred at the origin in exactly one point. Consequently, the circle of radius 1 is divided into two segments, depicted in red and blue in Figure 10.1. The minimum of the lengths of the two circle segments is the angle θ between $\vec{v}_1$ and $\vec{v}_2$. It is tempting to use (10.3) as a definition of the angle between two vectors v_1, v_2 in an abstract Euclidean space $(V, \langle \cdot, \cdot \rangle)$. That is, for any non-zero vectors $v_1, v_2 \in V$ define the angle between v_1, v_2 to be the unique real-number $\theta \in [0, \pi]$ such that

$$\cos \theta = \frac{\langle v_1, v_2 \rangle}{\|v_1\| \|v_2\|}.$$

Since the cosine is a bijective mapping from $[0, \pi]$ into $[-1, 1]$, this definition only makes sense if the quotient $\langle v_1, v_2 \rangle / (\|v_1\| \|v_2\|)$ lies in the interval $[-1, 1]$ for all pairs $v_1, v_2 \in V$ of non-zero vectors. That this is indeed the case follows from one of the most important inequalities in mathematics (recall that for $x \in \mathbb{R}$ we write $|x|$ for the absolute value of x):

Proposition 10.8 (Cauchy–Schwarz inequality) *Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space. Then, for any two vectors $v_1, v_2 \in V$, we have*

$$(10.1) \quad |\langle v_1, v_2 \rangle| \leq \|v_1\| \|v_2\|.$$

Furthermore, $|\langle v_1, v_2 \rangle| = \|v_1\| \|v_2\|$ if and only if $\{v_1, v_2\}$ are linearly dependent.

By the Cauchy–Schwarz inequality we thus have for all non-zero vectors $v_1, v_2 \in V$

$$0 \leq \frac{|\langle v_1, v_2 \rangle|}{\|v_1\| \|v_2\|} \leq 1,$$

so that $\langle v_1, v_2 \rangle / (\|v_1\| \|v_2\|) \in [-1, 1]$. This allows to define:

Definition 10.9 (Angle between two vectors) Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space and $v_1, v_2 \in V$ two non-zero vectors. The *angle between the vectors v_1 and v_2* is the unique real number $\theta \in [0, \pi]$ such that

$$\cos \theta = \frac{\langle v_1, v_2 \rangle}{\|v_1\| \|v_2\|}.$$

Remark 10.10 Notice that two non-zero vectors v_1, v_2 in $\mathbb{R}^2$ are orthogonal in the sense that $\langle v_1, v_2 \rangle = 0$ if and only if they are perpendicular, that is, the angle between v_1 and v_2 is $\pi/2$.

Proof of Proposition 10.8 First consider the case where $v_2 = 0_V$. Then both sides of (10.1) are 0, hence the inequality holds and, moreover, v_1 and v_2 are linearly dependent.

Let therefore be $v_1, v_2 \in V$ with $v_2 \neq 0_V$ and consider the function $p : \mathbb{R} \rightarrow \mathbb{R}$ defined by the rule

$$p(t) = \langle v_1 + tv_2, v_1 + tv_2 \rangle$$

for all $t \in \mathbb{R}$. Using the bilinearity and the symmetry of $\langle \cdot, \cdot \rangle$, we expand

$$p(t) = \langle v_1, v_1 \rangle + 2t\langle v_1, v_2 \rangle + t^2\langle v_2, v_2 \rangle = \|v_1\|^2 + 2t\langle v_1, v_2 \rangle + t^2\|v_2\|^2.$$

Since $v_2 \neq 0_V$, the function p is a polynomial of degree 2 in the variable t . If the discriminant of p is positive, then p has two distinct zeros and attains both positive and negative values. The bilinear form $\langle \cdot, \cdot \rangle$ is positive definite, hence we have $p(t) \geq 0$ for all $t \in \mathbb{R}$ and the discriminant Δ of p must be non-positive

$$\Delta = 4(\langle v_1, v_2 \rangle^2 - \|v_1\|^2 \|v_2\|^2) \leq 0.$$

Taking the square root implies (10.1).

If v_1, v_2 are linearly dependent and since $v_2 \neq 0_V$, there exists a scalar s such that $v_1 = sv_2$. Hence $\|v_2\| \|v_1\| = |s| \|v_2\|^2 = |\langle sv_2, v_2 \rangle|$ and equality holds in (10.1).

Conversely, suppose that $\langle v_1, v_2 \rangle = \pm \|v_1\| \|v_2\|$. Then we obtain

$$p(t) = \|v_1\|^2 \pm 2t\|v_1\| \|v_2\| + t^2\|v_2\|^2 = (\|v_1\| \pm t\|v_2\|)^2.$$

Taking $t_0 = \mp \|v_1\| / \|v_2\|$ gives $p(t_0) = \langle v_1 + t_0 v_2, v_1 + t_0 v_2 \rangle = 0$. Since $\langle \cdot, \cdot \rangle$ is positive definite, this implies that $v_1 + t_0 v_2 = 0_V$ and hence v_1, v_2 are linearly dependent. $\square$

Example 10.11 (Cauchy–Schwarz inequality)

- (i) Consider $V = \mathbb{R}^n$ equipped with the standard scalar product $\langle \cdot, \cdot \rangle$. The Cauchy–Schwarz inequality translates to the statement that for all $\vec{x} = (x_i)_{1 \leq i \leq n}$ and $\vec{y} = (y_i)_{1 \leq i \leq n} \in \mathbb{R}^n$, we have

$$\left| \sum_{i=1}^n x_i y_i \right| \leq \sqrt{\sum_{i=1}^n x_i^2} \sqrt{\sum_{i=1}^n y_i^2}.$$

- (ii) For $V = C([a, b], \mathbb{R})$ and inner product defined as in [Example 10.4](#) (iii) above, taking the square of the Cauchy–Schwarz inequality, we obtain that for all $f, g \in V$

$$\left| \int_a^b f(x)g(x)dx \right|^2 \leq \int_a^b f(x)^2 dx \int_a^b g(x)^2 dx.$$

The norm induced by an inner product satisfies a few elementary properties:

Proposition 10.12 (Properties of the norm) *Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space with induced norm $\|\cdot\| : V \rightarrow \mathbb{R}$. Then*

- (i) *for all $v \in V$ we have $\|v\| \geq 0$ and $\|v\| = 0$ if and only if $v = 0_V$;*
- (ii) *for all $s \in \mathbb{R}$ and all $v \in V$ we have $\|sv\| = |s|\|v\|$;*
- (iii) *for all vectors $v_1, v_2 \in V$, we have the so-called triangle inequality*

$$\|v_1 + v_2\| \leq \|v_1\| + \|v_2\|.$$

Proof The first two properties follow immediately from the definition of $\|\cdot\|$ and the positive definiteness of $\langle \cdot, \cdot \rangle$. Using the Cauchy–Schwarz inequality ([10.1](#)), we obtain for all $v_1, v_2 \in V$

$$\begin{aligned} \|v_1 + v_2\|^2 &= \langle v_1 + v_2, v_1 + v_2 \rangle = \langle v_1, v_1 \rangle + 2\langle v_1, v_2 \rangle + \langle v_2, v_2 \rangle \\ &\leq \|v_1\|^2 + 2|\langle v_1, v_2 \rangle| + \|v_2\|^2 \leq \|v_1\|^2 + 2\|v_1\|\|v_2\| + \|v_2\|^2 \\ &= (\|v_1\| + \|v_2\|)^2, \end{aligned}$$

and where we also used that $\langle v_1, v_2 \rangle \leq |\langle v_1, v_2 \rangle|$. Since both $\|v_1 + v_2\| \geq 0$ and $\|v_1\| + \|v_2\| \geq 0$, taking the square root implies

$$\|v_1 + v_2\| \leq \|v_1\| + \|v_2\|,$$

as claimed. □

Likewise, we obtain:

Proposition 10.13 (Properties of the metric) *Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space with induced metric $d : V \times V \rightarrow \mathbb{R}$. Then for all $v_1, v_2, v_3 \in V$ we have*

- (i) *$d(v_1, v_2) = 0$ if and only if $v_1 = v_2$;*
- (ii) *$d(v_1, v_2) = d(v_2, v_1)$ (symmetry);*
- (iii) *$d(v_1, v_3) \leq d(v_1, v_2) + d(v_2, v_3)$ (triangle inequality).*

Proof Exercise. □

10.2 The orthogonal projection

In the Euclidean setting, the restriction of an inner product to a subspace is again an inner product:

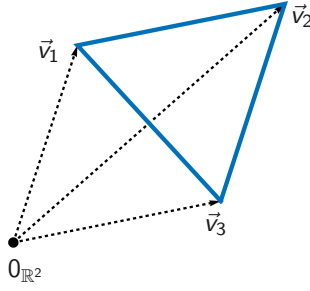


FIGURE 10.2. In $\mathbb{R}^2$ the triangle inequality states that for any triangle, the sum of the lengths of any two sides must be greater than or equal to the length of the remaining side.

Lemma 10.14 *Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space and $U \subset V$ a subspace. Then the restriction $\langle \cdot, \cdot \rangle|_U$ of $\langle \cdot, \cdot \rangle$ to U is an inner product and hence $(U, \langle \cdot, \cdot \rangle|_U)$ is a Euclidean space as well.*

Proof Symmetry and positive definiteness holds for all vectors or pairs of vectors in V , hence also for all vectors or pairs of vectors in $U \subset V$. $\square$

Remark 10.15 Since an inner product is a map $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}$, it would be more precise to write $\langle \cdot, \cdot \rangle|_{U \times U}$ and speak of the restriction of $\langle \cdot, \cdot \rangle$ to $U \times U$. For simplicity, we will use the terminology of [Lemma 10.14](#).

Recall that a projection is an endomorphism $\Pi : V \rightarrow V$ which satisfies $\Pi \circ \Pi = \Pi$ and that for a projection $\Pi : V \rightarrow V$ we have $V = \text{Ker } \Pi \oplus \text{Im } \Pi$. Given two subspaces U_1 and U_2 of V such that $V = U_1 \oplus U_2$, we can write every vector $v \in V$ uniquely as a sum $v = u_1 + u_2$ where $u_i \in U_i$ for $i = 1, 2$. The mapping $\Pi : V \rightarrow V$ defined by the rule $\Pi(v) = u_1$ for all $v \in V$ thus is a projection with $\text{Im } \Pi = U_1$ and $\text{Ker } \Pi = U_2$. Notice that Π is the unique projection whose image is U_1 and whose kernel is U_2 . If $\hat{\Pi} : V \rightarrow V$ is another projection with this property, then we have for all $v \in V$

$$\hat{\Pi}(v) = \hat{\Pi}(u_1 + u_2) = \hat{\Pi}(u_1)$$

Since $u_1 \in U_1 = \text{Im } \hat{\Pi}$, we can write $u_1 = \hat{\Pi}(w)$ for some vector $w \in V$, hence $\hat{\Pi}(u_1) = \hat{\Pi}(\hat{\Pi}(w)) = \hat{\Pi}(w) = u_1$. We thus have

$$\hat{\Pi}(v) = \hat{\Pi}(u_1) = u_1 = \Pi(v)$$

so that $\hat{\Pi} = \Pi$. This shows that there is precisely one projection with $\text{Im } \Pi = U_1$ and $\text{Ker } \Pi = U_2$.

Remark 10.16 By [Lemma 10.14](#) and [Remark 10.3](#), the restriction of an inner product $\langle \cdot, \cdot \rangle$ on a finite dimensional vector space V to a subspace $U \subset V$ is always non-degenerate. Therefore, by [Corollary 9.24](#), the orthogonal subspace $U^\perp$ is always a complement to U , so that $V = U \oplus U^\perp$ and

$$\dim U^\perp = \dim V - \dim U$$

by [Remark 6.7](#) and [Proposition 6.12](#).

This allows to define:

Definition 10.17 (Orthogonal projection) Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space and $U \subset V$ a subspace. The projection whose image is U and whose kernel is $U^\perp$ is called the *orthogonal projection onto the subspace U* and will be denoted by $\Pi_U^\perp$.

While the existence of the orthogonal projection onto a subspace U of a Euclidean space $(V, \langle \cdot, \cdot \rangle)$ follows abstractly from the above considerations, it is illustrative to give an explicit geometric construction. We first consider the case where U is spanned by a non-zero vector $u \in V$. We define a linear map $\Pi_U^\perp : V \rightarrow V$ by the rule

$$\Pi_U^\perp(v) = \frac{\langle v, u \rangle}{\langle u, u \rangle} u$$

for all $v \in V$. Then $\Pi_U^\perp(u) = u$ and $\text{Ker } \Pi_U^\perp = \{v \in V \mid \langle v, u \rangle = 0\} = U^\perp$. Since $\Pi_U^\perp(v) = su$ for some scalar $s \in \mathbb{K}$, we conclude that $\Pi_U^\perp \circ \Pi_U^\perp = \Pi_U^\perp$, hence $\Pi_U^\perp$ is the orthogonal projection onto U .

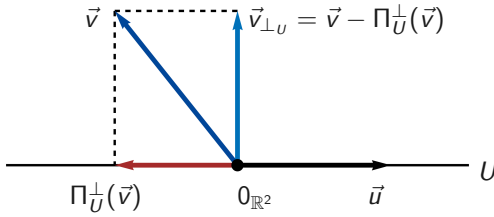


FIGURE 10.3. Orthogonal projection of the vector $\vec{v} \in \mathbb{E}^2$ onto the subspace U spanned by $\vec{u}$. Notice that the vector $\vec{v}_{\perp U} = \vec{v} - \Pi_U^\perp(\vec{v})$ is orthogonal to the vector $\vec{u}$.

In general, we have:

Proposition 10.18 Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space and $U \subset V$ a subspace of dimension $k \in \mathbb{N}$. Let $\{u_1, \dots, u_k\}$ be an orthogonal basis of U , then the map $\Pi_U^\perp : V \rightarrow V$ defined by the rule

$$(10.2) \quad \Pi_U^\perp(v) = \sum_{i=1}^k \frac{\langle v, u_i \rangle}{\langle u_i, u_i \rangle} u_i$$

for all $v \in V$ is the orthogonal projection onto U .

Proof Let $n = \dim V$. Notice that since $\langle \cdot, \cdot \rangle$ is positive definite, we must have $\langle u_i, u_i \rangle > 0$ for $1 \leq i \leq k$, hence the map $\Pi_U^\perp$ is well defined. For $1 \leq j \leq k$ we obtain

$$\Pi_U^\perp(u_j) = \sum_{i=1}^k \frac{\langle u_j, u_i \rangle}{\langle u_i, u_i \rangle} u_i = \frac{\langle u_j, u_j \rangle}{\langle u_j, u_j \rangle} u_j = u_j,$$

where we use the orthogonality of the basis $\{u_1, \dots, u_k\}$. By definition, for all $v \in V$ we have $\Pi_U^\perp(v) = \sum_{i=1}^k s_i u_i$ for scalars $s_i = \frac{\langle v, u_i \rangle}{\langle u_i, u_i \rangle}$. Since $\Pi_U^\perp(u_i) = u_i$, we obtain

$$\Pi_U^\perp(\Pi_U^\perp(v)) = \Pi_U^\perp\left(\sum_{i=1}^k s_i u_i\right) = \sum_{i=1}^k s_i \Pi_U^\perp(u_i) = \sum_{i=1}^k s_i u_i = \Pi_U^\perp(v).$$

Hence we have $\Pi_U^\perp \circ \Pi_U^\perp = \Pi_U^\perp$ and $\Pi_U^\perp$ is a projection.

By [Remark 10.16](#) we can write $V = U \oplus U^\perp$ and by [Theorem 3.64](#) we can find a basis $\{u_{k+1}, \dots, u_n\}$ of $U^\perp$ so that $\{u_1, \dots, u_k, u_{k+1}, \dots, u_n\}$ is a basis of V . Let $v \in V$. We write

$v = \sum_{j=1}^n t_j u_j$ for scalars t_j , $1 \leq j \leq n$. Then v lies in the kernel of $\Pi_U^\perp$ if and only if we have

$$0_V = \Pi_U^\perp(v) = \sum_{i=1}^k \frac{\left\langle \sum_{j=1}^n t_j u_j, u_i \right\rangle}{\langle u_i, u_i \rangle} u_i = \sum_{i=1}^k \sum_{j=1}^n t_j \frac{\langle u_j, u_i \rangle}{\langle u_i, u_i \rangle} u_i = \sum_{i=1}^k t_i u_i,$$

where we use that the vectors $\{u_1, \dots, u_k\}$ are orthogonal and that $\{u_{k+1}, \dots, u_n\} \in U^\perp$. The vector v thus lies in the kernel of $\Pi_U^\perp$ if and only if $v = \sum_{i=k+1}^n t_i u_i$, that is, if and only if $v \in U^\perp$. The map $\Pi_U^\perp$ thus is the orthogonal projection on U . $\square$

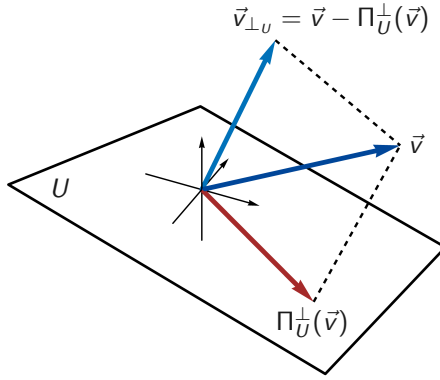


FIGURE 10.4. Orthogonal projection onto the plane U in $\mathbb{E}^3$.

Remark 10.19 Let $(V, \langle \cdot, \cdot \rangle)$ be a finite-dimensional Euclidean space and $U \subset V$ a subspace. Then for all $v \in V$ we can write $v = v - \Pi_U^\perp(v) + \Pi_U^\perp(v)$. Since $\Pi_U^\perp(v) \in U$ and $V = U \oplus U^\perp$, it follows that the vector

$$v_{\perp U} = v - \Pi_U^\perp(v) \in U^\perp$$

and moreover, $v_{\perp U} = 0_V$ if and only if $v \in U$.

Exercises

Exercise 10.20 Let $\vec{v}_1, \vec{v}_2 \in \mathbb{E}^2$ be two non-zero vectors. Show that the angle θ between $\vec{v}_1$ and $\vec{v}_2$ satisfies

$$(10.3) \quad \langle \vec{v}_1, \vec{v}_2 \rangle = \|\vec{v}_1\| \|\vec{v}_2\| \cos \theta.$$

10.3 Gram–Schmidt orthonormalisation

Using the orthogonal projection onto a subspace, we can now describe an explicit computational algorithm which constructs an orthonormal basis from a given ordered basis $\mathbf{b} = (v_1, \dots, v_n)$ of a Euclidean space $(V, \langle \cdot, \cdot \rangle)$. This algorithm is known as *Gram–Schmidt orthonormalisation*.

We first consider the case of a 3-dimensional Euclidean space $(V, \langle \cdot, \cdot \rangle)$ equipped with an ordered basis $\mathbf{b} = (v_1, v_2, v_3)$. We take

$$u_1 = \frac{v_1}{\|v_1\|}$$

as the first vector of our new orthonormal basis. We then construct a vector from v_2 that is orthogonal to the subspace $U_1 = \text{span}\{u_1\}$

$$w_2 = v_2 - \Pi_{U_1}^\perp(v_2) = v_2 - \frac{\langle v_2, u_1 \rangle}{\langle u_1, u_1 \rangle} u_1 = v_2 - \langle v_2, u_1 \rangle u_1,$$

where we use that $\langle u_1, u_1 \rangle = 1$. As our second basis vector we can thus take

$$u_2 = \frac{w_2}{\|w_2\|}.$$

We then define $U_2 = \text{span}\{u_1, u_2\}$ and set

$$w_3 = v_3 - \Pi_{U_2}^\perp(v_3) = v_3 - \frac{\langle v_3, u_1 \rangle}{\langle u_1, u_1 \rangle} u_1 - \frac{\langle v_3, u_2 \rangle}{\langle u_2, u_2 \rangle} u_2 = v_3 - \langle v_3, u_1 \rangle u_1 - \langle v_3, u_2 \rangle u_2$$

As our third basis vector we can thus take

$$u_3 = \frac{w_3}{\|w_3\|}.$$

Setting $\mathbf{b}' = (u_1, u_2, u_3)$, we have obtained an orthonormal basis $\mathbf{b}'$ of $(V, \langle \cdot, \cdot \rangle)$.

Example 10.21 We consider $V = \mathbb{R}^3$ with the standard scalar product $\langle \cdot, \cdot \rangle$ and the ordered basis $\mathbf{b} = (\vec{v}_1, \vec{v}_2, \vec{v}_3)$, where

$$\vec{v}_1 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \quad \vec{v}_2 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}, \quad \vec{v}_3 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}.$$

We apply Gram–Schmidt orthonormalisation to $\mathbf{b}$. We obtain

$$\vec{u}_1 = \frac{\vec{v}_1}{\|\vec{v}_1\|} = \frac{1}{\sqrt{3}} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$$

and

$$\vec{w}_2 = \vec{v}_2 - \langle \vec{v}_2, \vec{u}_1 \rangle \vec{u}_1 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} - \frac{2}{\sqrt{3}} \frac{1}{\sqrt{3}} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} = \frac{1}{3} \begin{pmatrix} 1 \\ -2 \\ 1 \end{pmatrix}$$

so that

$$\vec{u}_2 = \frac{\vec{w}_2}{\|\vec{w}_2\|} = \begin{pmatrix} \frac{1}{\sqrt{6}} \\ -\frac{\sqrt{2}}{\sqrt{3}} \\ \frac{1}{\sqrt{6}} \end{pmatrix}.$$

Likewise,

$$\begin{aligned} \vec{w}_3 &= \vec{v}_3 - \langle \vec{v}_3, \vec{u}_1 \rangle \vec{u}_1 - \langle \vec{v}_3, \vec{u}_2 \rangle \vec{u}_2 \\ &= \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} - \frac{2}{\sqrt{3}} \frac{1}{\sqrt{3}} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} - \left(-\frac{1}{\sqrt{6}}\right) \begin{pmatrix} \frac{1}{\sqrt{6}} \\ -\frac{\sqrt{2}}{\sqrt{3}} \\ \frac{1}{\sqrt{6}} \end{pmatrix} = \frac{1}{2} \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix} \end{aligned}$$

so that

$$\vec{u}_3 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}$$

and we have indeed $\langle \vec{u}_i, \vec{u}_j \rangle = \delta_{ij}$ for $1 \leq i, j \leq 3$. Hence the ordered basis $\mathbf{b}' = (\vec{u}_1, \vec{u}_2, \vec{u}_3)$ is orthonormal with respect to $\langle \cdot, \cdot \rangle$.

The careful reader might object that we have not argued above that $\mathbf{b}'$ is indeed well defined and an ordered basis. This is however the case:

Theorem 10.22 (Gram–Schmidt orthonormalisation) *Let $(V, \langle \cdot, \cdot \rangle)$ be an n -dimensional Euclidean space and $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V . For $2 \leq i \leq n$ we define recursively*

$$w_i = v_i - \Pi_{U_{i-1}}^\perp(v_i) \quad \text{and} \quad u_i = \frac{w_i}{\|w_i\|},$$

where $U_{i-1} = \text{span}\{u_1, \dots, u_{i-1}\}$ and $u_1 = v_1/\|v_1\|$. Then $\mathbf{b}' = (u_1, \dots, u_n)$ is well defined and an orthonormal ordered basis of V . Moreover, $\mathbf{b}'$ is the unique orthonormal ordered basis of V so that the change of basis matrix $\mathbf{C}(\mathbf{b}', \mathbf{b})$ is an upper triangular matrix with positive diagonal entries.

Proof We will use induction on the dimension n of the Euclidean space $(V, \langle \cdot, \cdot \rangle)$. In the case where $\dim V = 1$ we have a single basis vector $v_1 \neq 0_V$. We set $u_1 = v_1/\|v_1\|$. Then $\mathbf{b}' = (u_1)$ is an ordered basis of V which is orthonormal. The change of basis matrix is $\mathbf{C}(\mathbf{b}', \mathbf{b}) = (1/\|v_1\|)$ and hence is an upper triangular matrix with positive diagonal entries. The only other ordered basis of V which is orthonormal is $(-u_1)$, but the change of basis matrix for this basis has a negative diagonal entry. Therefore, the statement is anchored.

Inductive step: Suppose $n \geq 2$ and that the statement is true for an $(n-1)$ -dimensional Euclidean space. Let $(V, \langle \cdot, \cdot \rangle)$ be an n -dimensional Euclidean space and $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V . Consider the subspace $U_{n-1} = \text{span}\{v_1, \dots, v_{n-1}\}$ of dimension $n-1$ for which $\mathbf{c} = (v_1, \dots, v_{n-1})$ is an ordered basis. By the induction hypothesis, there exists a unique ordered basis $\mathbf{c}' = (u_1, \dots, u_{n-1})$ of U_{n-1} which is orthonormal and such that the change of basis matrix $\mathbf{C}(\mathbf{c}', \mathbf{c})$ is an upper triangular matrix with positive diagonal entries. Set $w_n = v_n - \Pi_{U_{n-1}}^\perp(v_n)$ so that $w_n \in U_{n-1}^\perp$. Since $\mathbf{b}$ is a basis it follows that $v_n \notin U_{n-1}$, therefore Remark 10.19 implies that $w_n \neq 0_V$ and we conclude that $\{u_1, \dots, u_{n-1}, w_n\}$ is orthogonal as well as linearly independent. Let $u_n = w_n/\|w_n\|$, then $\mathbf{b}' = (u_1, \dots, u_n)$ is an ordered basis of V which is orthonormal. By definition, we have

$$u_n = \frac{v_n - \Pi_{U_{n-1}}^\perp(v_n)}{\|v_n - \Pi_{U_{n-1}}^\perp(v_n)\|} = \frac{v_n}{\|v_n - \Pi_{U_{n-1}}^\perp(v_n)\|} + \sum_{i=1}^{n-1} s_i v_i$$

for suitable scalars $s_1, \dots, s_{n-1}$. Writing $\vec{s} = (s_i)_{1 \leq i \leq n-1}$, the change of basis matrix thus takes the form

$$\mathbf{C}(\mathbf{b}', \mathbf{b}) = \begin{pmatrix} \mathbf{C}(\mathbf{c}', \mathbf{c}) & \vec{s} \\ 0_{\mathbb{R}_{n-1}} & \frac{1}{\|v_n - \Pi_{U_{n-1}}^\perp(v_n)\|} \end{pmatrix}.$$

Since $\mathbf{C}(\mathbf{c}', \mathbf{c})$ is an upper triangular matrix with positive entries, it follows that $\mathbf{C}(\mathbf{b}', \mathbf{b})$ is an upper triangular matrix with positive entries as well.

Finally, we argue that $\mathbf{b}'$ is the unique ordered basis of V satisfying the conditions of the theorem. Notice that u_n must be an element of $U_{n-1}^\perp$. Now $\dim V = n$ and $\dim U_{n-1} = n-1$ and since $V = U_{n-1} \oplus U_{n-1}^\perp$ by Corollary 9.24, we must have $\dim U_{n-1}^\perp = 1$ by

Remark 10.16. This implies that u_n is uniquely determined up to multiplication by ± 1 , but the above choice is the only one resulting in a change of basis matrix which has positive diagonal entries. Since by the induction hypothesis the basis $\mathbf{c}'$ is unique, it follows that $\mathbf{b}'$ is the unique ordered basis of V satisfying the conditions of the theorem. $\square$

It follows from [Theorem 3.64](#) that an ordered basis of a subspace U of a finite dimensional vector space V can always be extended to a basis of V . A corresponding statement is also true for orthonormal bases:

Corollary 10.23 *Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space and $U \subset V$ a subspace. Suppose $\mathbf{b}$ is an ordered orthonormal basis of U , then there exists an ordered orthonormal basis of V which contains $\mathbf{b}$.*

Proof Let $k = \dim U$, $n = \dim V$ and $\mathbf{b} = (v_1, \dots, v_k)$. Choose any ordered basis $\mathbf{c}$ of $U^\perp$ and apply Gram–Schmidt orthonormalisation to $\mathbf{c}$ to obtain an orthonormal basis $\mathbf{b}' = (v_{k+1}, \dots, v_n)$ of $U^\perp$. Since all vectors of U are orthogonal to all vectors of $U^\perp$, the ordered basis $(v_1, \dots, v_n)$ is an orthonormal ordered basis for $(V, \langle \cdot, \cdot \rangle)$. $\square$

Notice that if we carry out the Gram–Schmidt procedure without normalising the vectors w_i at each step – sometimes referred to as *Gram–Schmidt orthogonalisation* – then we still obtain an ordered orthogonal basis $(w_1, \dots, w_n)$.

Example 10.24 (Legendre polynomials) We consider again the vector space $V = C([-1, 1], \mathbb{R})$ of continuous real-valued functions defined on the interval $[-1, 1]$, equipped with the bilinear form defined by the rule

$$\langle f, g \rangle = \int_{-1}^1 f(x)g(x)dx$$

for all $f, g \in V$. For $n \in \mathbb{N} \cup \{0\}$ let U_n denote the subspace of V consisting of polynomials of degree n . An ordered basis of U_n is given by the polynomials $\mathbf{b} = (1, x, x^2, x^3, \dots, x^n)$. Applying Gram–Schmidt orthogonalisation we obtain an ordered orthogonal basis $(p_0, p_1, \dots, p_n)$ of U_n . That is, for $i \neq j$, the polynomials satisfy

$$\langle p_i, p_j \rangle = \int_{-1}^1 p_i(x)p_j(x)dx = 0.$$

The polynomials p_i are known as the *Legendre polynomials*. There are different ways to normalise the Legendre polynomials. Besides the standard normalisation which makes the polynomials orthonormal, that is, $\langle p_i, p_i \rangle = 1$, it is also common to request that $\langle p_i, p_i \rangle = 2/(2i + 1)$. The reason for this normalisation is that it allows to give a neat formula for p_i known as *Rodrigues' formula* (which we will not prove)

$$p_i(x) = \frac{1}{2^i i!} \frac{d^i}{dx^i} (x^2 - 1)^i,$$

where $\frac{d^i}{dx^i}$ stands for the i -th derivative with respect to the variable x . Using this formula we obtain for the first four Legendre polynomials

$$p_0(x) = \frac{1}{2^0 0!} \frac{d^0}{dx^0} (x^2 - 1)^0 = (x^2 - 1)^0 = 1,$$

$$p_1(x) = \frac{1}{2^1 1!} \frac{d}{dx} (x^2 - 1)^1 = \frac{1}{2} (2x) = x,$$

$$p_2(x) = \frac{1}{2^2 2!} \frac{d^2}{dx^2} (x^2 - 1)^2 = \frac{1}{8} \frac{d^2}{dx^2} (x^4 - 2x^2 + 1) = \frac{1}{2} (3x^2 - 1),$$

$$p_3(x) = \frac{1}{2^3 3!} \frac{d^3}{dx^3} (x^2 - 1)^3 = \frac{1}{48} \frac{d^3}{dx^3} (x^6 - 3x^4 + 3x^2 - 1) = \frac{1}{2} (5x^3 - 3x).$$

The Gram–Schmidt orthonormalisation [Theorem 10.22](#) has a matrix version known as the Cholesky decomposition. In order to phrase it, we make the following definition.

Definition 10.25 (Positive definite matrix) Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{R})$. The matrix $\mathbf{A}$ is called *positive definite* if the bilinear form $\langle \cdot, \cdot \rangle_{\mathbf{A}}$ on $\mathbb{R}^n$ is positive definite.

Theorem 10.26 (Cholesky decomposition) Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{R})$ be a symmetric positive definite matrix. Then there exists a unique upper triangular matrix $\mathbf{C} \in M_{n,n}(\mathbb{R})$ with positive diagonal entries such that $\mathbf{A} = \mathbf{C}^T \mathbf{C}$.

Proof Since $\mathbf{A}$ is positive definite and symmetric, $\langle \cdot, \cdot \rangle_{\mathbf{A}}$ is an inner product on $\mathbb{R}^n$. Let $\mathbf{e} = (\vec{e}_1, \dots, \vec{e}_n)$ denote the standard ordered basis of $\mathbb{R}^n$. Recall that we have $\mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{e}) = \mathbf{A}$. [Theorem 10.22](#) implies the existence of a unique ordered basis $\mathbf{b}'$ of $\mathbb{R}^n$ which is orthonormal with respect to $\langle \cdot, \cdot \rangle_{\mathbf{A}}$. Therefore, using [Proposition 9.6](#), we obtain

$$\mathbf{A} = \mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{e}) = \mathbf{C}^T \mathbf{M}(\langle \cdot, \cdot \rangle_{\mathbf{A}}, \mathbf{b}') \mathbf{C} = \mathbf{C}^T \mathbf{C},$$

where $\mathbf{C} = \mathbf{C}(\mathbf{e}, \mathbf{b}')$ and where we use that the matrix representation of an inner product with respect to an orthonormal basis is the identity matrix. [Theorem 10.22](#) implies that $\mathbf{C}(\mathbf{b}', \mathbf{e})$ is an upper triangular matrix with positive diagonal entries. By [Remark 3.105](#) we have $\mathbf{C}(\mathbf{e}, \mathbf{b}') = \mathbf{C}(\mathbf{b}', \mathbf{e})^{-1}$ and hence [Corollary 5.46](#) implies that $\mathbf{C}$ is an upper triangular matrix as well. Now for $2 \leq i \leq n-1$, we have (the cases $i = 1$ and $i = n$ are similar)

$$\begin{aligned} 1 &= [\mathbf{C}\mathbf{C}^{-1}]_{ii} = \sum_{k=1}^n [\mathbf{C}]_{ik} [\mathbf{C}^{-1}]_{ki} = [\mathbf{C}]_{ii} [\mathbf{C}^{-1}]_{ii} + \sum_{k=1}^{i-1} [\mathbf{C}]_{ik} [\mathbf{C}^{-1}]_{ki} + \sum_{k=i+1}^n [\mathbf{C}]_{ik} [\mathbf{C}^{-1}]_{ki} \\ &= [\mathbf{C}]_{ii} [\mathbf{C}^{-1}]_{ii}, \end{aligned}$$

where we use that $\mathbf{C}$ and $\mathbf{C}^{-1}$ are upper triangular matrices. It follows that $[\mathbf{C}]_{ii}$ has the same sign as $[\mathbf{C}^{-1}]_{ii}$ for $1 \leq i \leq n$. Therefore we conclude that $\mathbf{C}$ has positive diagonal entries as well.

Suppose that $\hat{\mathbf{C}} \in M_{n,n}(\mathbb{R})$ is another upper triangular matrix with positive diagonal entries so that $\mathbf{A} = \hat{\mathbf{C}}^T \hat{\mathbf{C}}$. Using [Lemma 3.109](#) we conclude that there exists an ordered basis $\mathbf{c}'$ of $\mathbb{R}^n$ such that $\hat{\mathbf{C}} = \mathbf{C}(\mathbf{e}, \mathbf{c}')$. Since $\mathbf{A} = \hat{\mathbf{C}}^T \hat{\mathbf{C}}$ the basis $\mathbf{c}'$ is orthonormal with respect to $\langle \cdot, \cdot \rangle_{\mathbf{A}}$. The uniqueness statement of [Theorem 10.22](#) implies that $\mathbf{c}' = \mathbf{b}'$ and hence $\mathbf{C} = \hat{\mathbf{C}}$. $\square$

Remark 10.27 Notice that every invertible matrix $\mathbf{C} \in M_{n,n}(\mathbb{R})$ gives rise to a symmetric positive definite matrix $\mathbf{A} = \mathbf{C}^T \mathbf{C}$. Indeed, by [Remark 2.18](#) we have $\mathbf{A}^T = (\mathbf{C}^T \mathbf{C})^T = \mathbf{C}^T (\mathbf{C}^T)^T = \mathbf{C}^T \mathbf{C} = \mathbf{A}$ so that $\mathbf{A}$ is symmetric. Using [Remark 2.18](#) again we obtain for all $\vec{x} \in \mathbb{R}^n$

$$\langle \vec{x}, \vec{x} \rangle_{\mathbf{A}} = \vec{x}^T \mathbf{C}^T \mathbf{C} \vec{x} = (\mathbf{C} \vec{x})^T \mathbf{C} \vec{x} = \langle \mathbf{C} \vec{x}, \mathbf{C} \vec{x} \rangle$$

where the bilinear form on the right hand side denotes the standard scalar product on $\mathbb{R}^n$. In particular this implies that $\langle \cdot, \cdot \rangle_{\mathbf{A}}$ is positive. Since the standard scalar product on $\mathbb{R}^n$ is positive definite, the last expression is 0 if and only if $\mathbf{C} \vec{x} = 0_{\mathbb{R}^n}$. Since $\mathbf{C}$ is invertible this condition is equivalent to $\vec{x} = 0_{\mathbb{R}^n}$. It follows that $\langle \cdot, \cdot \rangle_{\mathbf{A}}$ is positive definite as well.

Finally, we observe that the coordinate representation of a vector with respect to an orthonormal basis can be computed easily:

Remark 10.28 (Coordinate representation with respect to an orthonormal basis) Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space equipped with an ordered orthonormal basis $\mathbf{b} = (v_1, \dots, v_n)$ with corresponding linear coordinate system β . Then for all $v \in V$ we have

$$\beta(v) = \begin{pmatrix} \langle v, v_1 \rangle \\ \vdots \\ \langle v, v_n \rangle \end{pmatrix} \iff v = \sum_{i=1}^n \langle v, v_i \rangle v_i$$

Indeed, since $\mathbf{b}$ is a basis we can write $v = \sum_{i=1}^n s_i v_i$ for unique real numbers s_i , where $1 \leq i \leq n$. Using that $\langle v_i, v_j \rangle = 0$ for $i \neq j$ and that $\langle v_i, v_i \rangle = 1$, we obtain

$$\langle v, v_j \rangle = \left\langle \sum_{i=1}^n s_i v_i, v_j \right\rangle = \sum_{i=1}^n s_i \langle v_i, v_j \rangle = s_j.$$

Correspondingly, for all $v \in V$ we obtain the following formula for the length of v

$$\begin{aligned} \|v\| &= \left\| \sum_{i=1}^n \langle v, v_i \rangle v_i \right\| = \sqrt{\left\langle \sum_{i=1}^n \langle v, v_i \rangle v_i, \sum_{j=1}^n \langle v, v_j \rangle v_j \right\rangle} \\ &= \sqrt{\sum_{i=1}^n \sum_{j=1}^n \langle v, v_i \rangle \langle v, v_j \rangle \langle v_i, v_j \rangle} = \sqrt{\sum_{i=1}^n \langle v, v_i \rangle^2}. \end{aligned}$$

Remark 10.29 (Linear independence of orthogonal vectors) Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space and $\{u_1, \dots, u_k\}$ be non-zero orthogonal vectors so that $\langle u_i, u_j \rangle = 0$ for all $1 \leq i, j \leq k$ with $i \neq j$. Suppose we have scalars $s_1, \dots, s_k \in \mathbb{R}$ such that $\sum_{j=1}^k s_j u_j = 0_V$. Then, taking the inner product with u_i gives

$$0 = \langle 0_V, u_i \rangle = \sum_{j=1}^k s_j \langle u_j, u_i \rangle = s_i \langle u_i, u_i \rangle.$$

Since by assumption $u_i \neq 0_V$, we have $\langle u_i, u_i \rangle \neq 0$ and hence $s_i = 0$. It follows that $\{u_1, \dots, u_k\}$ is linearly independent.

Example 10.30 (Example 10.21 continued) If we want to compute $\mathbf{C}(\mathbf{b}, \mathbf{b}')$ we need to compute $\beta'(v_i)$ for $i = 1, 2, 3$ and write the resulting vectors into the columns of $\mathbf{C}(\mathbf{b}, \mathbf{b}')$. Since $\mathbf{b}'$ is orthonormal, the preceding remark gives

$$\beta'(v_1) = \begin{pmatrix} \langle v_1, u_1 \rangle \\ \langle v_1, u_2 \rangle \\ \langle v_1, u_3 \rangle \end{pmatrix} = \begin{pmatrix} \sqrt{3} \\ 0 \\ 0 \end{pmatrix}.$$

Likewise we have

$$\beta'(v_2) = \begin{pmatrix} \langle v_2, u_1 \rangle \\ \langle v_2, u_2 \rangle \\ \langle v_2, u_3 \rangle \end{pmatrix} = \begin{pmatrix} \frac{2}{\sqrt{3}} \\ \frac{\sqrt{2}}{\sqrt{3}} \\ 0 \end{pmatrix}$$

and

$$\beta'(v_3) = \begin{pmatrix} \langle v_3, u_1 \rangle \\ \langle v_3, u_2 \rangle \\ \langle v_3, u_3 \rangle \end{pmatrix} = \begin{pmatrix} \frac{2}{\sqrt{3}} \\ -\frac{1}{\sqrt{6}} \\ \frac{1}{\sqrt{2}} \end{pmatrix}$$

so that

$$\mathbf{C}(\mathbf{b}, \mathbf{b}') = \begin{pmatrix} \sqrt{3} & \frac{2}{\sqrt{3}} & \frac{2}{\sqrt{3}} \\ 0 & \frac{\sqrt{2}}{\sqrt{3}} & -\frac{1}{\sqrt{6}} \\ 0 & 0 & \frac{1}{\sqrt{2}} \end{pmatrix}.$$

The proof of Theorem 10.26 implies that $\mathbf{C}(\mathbf{b}', \mathbf{b}) = \mathbf{C}(\mathbf{b}, \mathbf{b}')^{-1}$ is an upper triangular matrix with positive diagonal entries as well, as predicted by Theorem 10.22.

Exercises

Exercise 10.31 Compute the Cholesky decomposition of the positive definite symmetric matrix

$$\mathbf{A} = \begin{pmatrix} 3 & 0 & -1 \\ 0 & 8 & 4 \\ -1 & 4 & 3 \end{pmatrix}.$$

10.4 The orthogonal group

Recall that an isomorphism of vector spaces V and W is a bijective linear map $f : V \rightarrow W$. In the case where both V and W are equipped with an inner product, we may ask that f preserves the inner products in the following sense:

Definition 10.32 (Orthogonal transformation) Let $(V, \langle \cdot, \cdot \rangle)$ and $(W, \langle \cdot, \cdot \rangle)$ be Euclidean spaces. An isomorphism $f : V \rightarrow W$ is called an *orthogonal transformation* if

$$\langle u, v \rangle = \langle f(u), f(v) \rangle$$

for all $u, v \in V$.

Recall that in a Euclidean space $(V, \langle \cdot, \cdot \rangle)$ both the notion of angle between two vectors (Definition 10.9) and the notion of distance between two vectors (Definition 10.7) only depends on the inner product $\langle \cdot, \cdot \rangle$. Orthogonal transformations thus preserve both angles between vectors and distances between vectors.

We can also consider the set of orthogonal transformations from a Euclidean space to itself:

Definition 10.33 (Orthogonal group & orthogonal matrices)

- Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space. The set of orthogonal transformations from $(V, \langle \cdot, \cdot \rangle)$ to itself is called the *orthogonal group* of $(V, \langle \cdot, \cdot \rangle)$ and denoted by $O(V, \langle \cdot, \cdot \rangle)$.
- A matrix $\mathbf{R} \in M_{n,n}(\mathbb{R})$ is called *orthogonal* if $f_{\mathbf{R}} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ is an orthogonal transformation of $(\mathbb{R}^n, \langle \cdot, \cdot \rangle)$, where $\langle \cdot, \cdot \rangle$ denotes the standard scalar product of $\mathbb{R}^n$. The set of orthogonal $n \times n$ -matrices is denoted by $O(n)$ and called the *orthogonal group*.

The use of the term group in the above definition is indeed justified:

Proposition 10.34 Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space. Then the set $O(V, \langle \cdot, \cdot \rangle)$ is a group in the sense of Definition 8.4 when the group operation is taken to be the composition of mappings. In particular, $O(n)$ is a group when the group operation is taken to be matrix multiplication.

Proof Let $G = O(V, \langle \cdot, \cdot \rangle)$. As the group identity element we take $e_G = \text{Id}_V$, where Id_V denotes the identity mapping on V , so that $\text{Id}_V(v) = v$ for all $v \in V$. Clearly $\text{Id}_V \in G$ and $f \circ \text{Id}_V = \text{Id}_V \circ f = f$ for all $f \in G$. Likewise, if $f \in G$, then the inverse mapping f^{-1} is an element of G as well. Indeed, for all $u, v \in V$ we obtain

$$\begin{aligned} \langle u, v \rangle &= \langle \text{Id}_V(u), \text{Id}_V(v) \rangle = \langle (f \circ f^{-1})(u), (f \circ f^{-1})(v) \rangle = \langle f(f^{-1}(u)), f(f^{-1}(v)) \rangle \\ &= \langle f^{-1}(u), f^{-1}(v) \rangle, \end{aligned}$$

where we use that $f \in G$. Therefore, for all $f \in G$ there exists a group element b , namely f^{-1} such that $f \circ b = b \circ f = e_G = \text{Id}_V$. Since the composition of mappings is associative, it follows that $O(V, \langle \cdot, \cdot \rangle)$ is a group with respect to the composition of mappings.

The second claim follows since for matrices $\mathbf{A}, \mathbf{B} \in M_{n,n}(\mathbb{R})$, we have $f_{\mathbf{A}} \circ f_{\mathbf{B}} = f_{\mathbf{AB}}$, where $\mathbf{AB}$ denotes the matrix multiplication of $\mathbf{A}$ and $\mathbf{B}$, see Theorem 2.21. $\square$

Lemma 10.35 For all $n \in \mathbb{N}$ we have

$$O(n) = \{\mathbf{R} \in M_{n,n}(\mathbb{R}) \mid \mathbf{R}^T \mathbf{R} = \mathbf{1}_n\} = \{\mathbf{R} \in GL(n, \mathbb{R}) \mid \mathbf{R}^T = \mathbf{R}^{-1}\}.$$

Proof By definition, $\mathbf{R} \in M_{n,n}(\mathbb{R})$ is an element of $O(n)$ if and only if

$$\langle \vec{x}, \vec{y} \rangle = \vec{x}^T \vec{y} = \langle \vec{x}, \vec{y} \rangle_{\mathbf{1}_n} = \langle \mathbf{R}\vec{x}, \mathbf{R}\vec{y} \rangle = (\mathbf{R}\vec{x})^T \mathbf{R}\vec{y} = \vec{x}^T \mathbf{R}^T \mathbf{R} \vec{y} = \langle \vec{x}, \vec{y} \rangle_{\mathbf{R}^T \mathbf{R}}$$

for all vectors $\vec{x}, \vec{y} \in \mathbb{R}^n$. From the exercises we know that this condition is equivalent to $\mathbf{R}^T \mathbf{R} = \mathbf{1}_n$, as claimed.

In order to show the second equality sign in the lemma, recall that $GL(n, \mathbb{R})$ consists of the matrices $\mathbf{A} \in M_{n,n}(\mathbb{R})$ that are invertible. If $\mathbf{R} \in GL(n, \mathbb{R})$ satisfies $\mathbf{R}^{-1} = \mathbf{R}^T$, then $\mathbf{R}^T \mathbf{R} = \mathbf{R}^{-1} \mathbf{R} = \mathbf{1}_n$ hence we have

$$\{\mathbf{R} \in GL(n, \mathbb{R}) \mid \mathbf{R}^T = \mathbf{R}^{-1}\} \subset \{\mathbf{R} \in M_{n,n}(\mathbb{R}) \mid \mathbf{R}^T \mathbf{R} = \mathbf{1}_n\}.$$

The converse inclusion of sets follows from the observation that a matrix $\mathbf{R} \in O(n)$ satisfies $\det \mathbf{R} = \pm 1$. Indeed, the product rule for the determinant [Proposition 5.21](#) gives

$$\det(\mathbf{R}^T \mathbf{R}) = \det(\mathbf{R}^T) \det(\mathbf{R}) = (\det(\mathbf{R}))^2 = \det(\mathbf{1}_n) = 1,$$

where we also use that $\det(\mathbf{A}^T) = \det(\mathbf{A})$ for all $\mathbf{A} \in M_{n,n}(\mathbb{R})$. Since $\det \mathbf{R} = \pm 1$, the matrix $\mathbf{R}$ is invertible and hence $\mathbf{R}^T \mathbf{R} = \mathbf{1}_n$ implies that $\mathbf{R}^T = \mathbf{R}^{-1}$. $\square$

The orthogonal transformations in a finite dimensional Euclidean space $(V, \langle \cdot, \cdot \rangle)$ can similarly be characterised in terms of their matrix representation with respect to an orthonormal basis:

Proposition 10.36 Let $n \in \mathbb{N}$ and $(V, \langle \cdot, \cdot \rangle)$ be an n -dimensional Euclidean space equipped with an orthonormal ordered basis $\mathbf{b}$. Then an endomorphism $f : V \rightarrow V$ is an orthogonal transformation if and only if its matrix representation $\mathbf{R} = \mathbf{M}(f, \mathbf{b}, \mathbf{b})$ with respect to $\mathbf{b}$ is an orthogonal matrix.

Proof By definition an endomorphism $f : V \rightarrow V$ is an orthogonal transformation of $(V, \langle \cdot, \cdot \rangle)$ if and only if $\langle u, v \rangle = \langle f(u), f(v) \rangle$ for all vectors $u, v \in V$. Writing $\vec{x} = \beta(u)$, $\vec{y} = \beta(v)$ and $\mathbf{R} = \mathbf{M}(f, \mathbf{b}, \mathbf{b})$, this gives

$$\begin{aligned} \langle u, v \rangle &= \vec{x}^T \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) \vec{y} = \langle \vec{x}, \vec{y} \rangle_{\mathbf{1}_n} = \langle f(u), f(v) \rangle \\ &= (\beta(f(u)))^T \beta(f(v)) = (\mathbf{R}\vec{x})^T \mathbf{R}\vec{y} = \vec{x}^T \mathbf{R}^T \mathbf{R} \vec{y} = \langle \vec{x}, \vec{y} \rangle_{\mathbf{R}^T \mathbf{R}}, \end{aligned}$$

where we use that $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) = \mathbf{1}_n$, [Proposition 9.6](#) and [Proposition 3.98](#). Since every vector $\vec{x} \in \mathbb{R}^n$ can be written as $\vec{x} = \beta(u)$ for some vector $u \in V$, the claim follows as in the proof of [Lemma 10.35](#). $\square$

Corollary 10.37 Let $n \in \mathbb{N}$ and $(V, \langle \cdot, \cdot \rangle)$ be an n -dimensional Euclidean space equipped with an orthonormal ordered basis $\mathbf{b}$. Then an ordered basis $\mathbf{b}'$ of V is orthonormal with respect to $\langle \cdot, \cdot \rangle$ if and only if the change of basis matrix $\mathbf{C}(\mathbf{b}', \mathbf{b})$ is orthogonal.

Proof Write $\mathbf{b} = (v_1, \dots, v_n)$, $\mathbf{b}' = (v'_1, \dots, v'_n)$ and let β, β' denote the corresponding linear coordinate systems. Consider the endomorphism $g = (\beta')^{-1} \circ \beta : V \rightarrow V$ satisfying $\mathbf{M}(g, \mathbf{b}, \mathbf{b}) = \mathbf{C}(\mathbf{b}', \mathbf{b})$. Using [Proposition 10.36](#) it is sufficient to show that g

is orthogonal if and only if $\mathbf{b}'$ is orthonormal. By definition, g satisfies $g(v_i) = v'_i$ for all $1 \leq i \leq n$. Suppose the endomorphism g is orthogonal, then

$$\langle v'_i, v'_j \rangle = \langle g(v_i), g(v_j) \rangle = \langle v_i, v_j \rangle = \delta_{ij},$$

where the last equality uses that $\mathbf{b}$ is orthonormal. We conclude that $\mathbf{b}'$ is orthonormal as well. Conversely, suppose that $\mathbf{b}'$ is orthonormal. Let $u, v \in V$ and write $u = \sum_{i=1}^n s_i v_i$ and $v = \sum_{j=1}^n t_j v_j$ for scalars $s_i, t_i, i = 1, \dots, n$. Then, using the bilinearity of $\langle \cdot, \cdot \rangle$, we compute

$$\begin{aligned} \langle g(u), g(v) \rangle &= \left\langle g\left(\sum_{i=1}^n s_i v_i\right), g\left(\sum_{j=1}^n t_j v_j\right) \right\rangle = \sum_{i=1}^n \sum_{j=1}^n s_i t_j \langle g(v_i), g(v_j) \rangle \\ &= \sum_{i=1}^n \sum_{j=1}^n s_i t_j \langle v'_i, v'_j \rangle = \sum_{i=1}^n \sum_{j=1}^n s_i t_j \delta_{ij} = \sum_{i=1}^n s_i t_i \langle v_i, v_i \rangle \\ &= \left\langle \sum_{i=1}^n s_i v_i, \sum_{j=1}^n t_j v_j \right\rangle = \langle u, v \rangle, \end{aligned}$$

so that g is orthogonal. □

Example 10.38 A matrix $\mathbf{A} \in M_{n,n}(\mathbb{R})$ is orthogonal if and only if its column vectors form an orthonormal basis of $\mathbb{R}^n$ with respect to the standard scalar product $\langle \cdot, \cdot \rangle$. To this end let $\hat{\Omega} : (\mathbb{R}^n)^n \rightarrow M_{n,n}(\mathbb{K})$ denote the map which forms an $n \times n$ matrix from n column vectors of length n . That is, $\hat{\Omega}$ satisfies

$$[\hat{\Omega}(\vec{a}_1, \dots, \vec{a}_n)]_{ij} = [\vec{a}_j]_i$$

for all $1 \leq i, j \leq n$ and where $[\vec{a}_j]_i$ denotes the i -th entry of the vector $\vec{a}_j$. Then, by the definition of matrix multiplication, we have for all $1 \leq i, j \leq n$

$$\begin{aligned} [\hat{\Omega}(\vec{a}_1, \dots, \vec{a}_n)^T \hat{\Omega}(\vec{a}_1, \dots, \vec{a}_n)]_{ij} &= \sum_{k=1}^n [\hat{\Omega}(\vec{a}_1, \dots, \vec{a}_n)^T]_{ik} [\hat{\Omega}(\vec{a}_1, \dots, \vec{a}_n)]_{kj} \\ &= \sum_{k=1}^n [\hat{\Omega}(\vec{a}_1, \dots, \vec{a}_n)]_{ki} [\hat{\Omega}(\vec{a}_1, \dots, \vec{a}_n)]_{kj} \\ &= \sum_{k=1}^n [\vec{a}_i]_k [\vec{a}_j]_k = \langle \vec{a}_i, \vec{a}_j \rangle_{1_n} = \delta_{ij}, \end{aligned}$$

as claimed.

The reader is invited to check that a corresponding statement also holds for the rows of an orthogonal matrix.

Example 10.39 (Permutation matrices) Let $n \in \mathbb{N}$ and $\sigma \in S_n$ be a permutation. Recall that for $1 \leq i \leq n$, the i -th column of the permutation matrix $\mathbf{P}_\sigma$ of σ is given by $\vec{e}_{\sigma(i)}$, where $\mathbf{e} = (\vec{e}_1, \dots, \vec{e}_n)$ denotes the standard ordered basis of $\mathbb{R}^n$. Therefore, the columns of a permutation matrix form an ordered orthonormal basis of $\mathbb{R}^n$ and hence permutation matrices are orthogonal by the previous remark.

Example 10.40 (Reflection along a hyperplane) A plane in $\mathbb{R}^3$ is a subspace U of dimension $2 = 3 - 1$. More generally, a *hyperplane* in an n -dimensional vector space V is a subspace U of dimension $n - 1$.

We can reflect a vector orthogonally along a plane in $\mathbb{E}^3$, see Figure 10.5. This map generalises to hyperplanes as follows: Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space and $U \subset V$ a hyperplane. Then the *orthogonal reflection along U* is the map $r_U : V \rightarrow V$ defined by the rule

$$r_U(v) = v - 2(v - \Pi_U^\perp(v)) = 2\Pi_U^\perp(v) - v$$

for all $v \in V$. This mapping is indeed an orthogonal transformation. To see this consider an orthonormal basis $\{u_1, \dots, u_{n-1}\}$ of U . By Corollary 10.23 we can extend this to an orthonormal basis $\{u_1, \dots, u_{n-1}, u_n\}$ of V . Notice that $\dim U^\perp = 1$ and that $U^\perp = \text{span}\{u_n\}$. For $v \in V$ we write $v = \sum_{i=1}^n s_i u_i$ for unique real numbers s_i , $1 \leq i \leq n$. Then we obtain

$$\begin{aligned} r_U(v) + v &= 2\Pi_U^\perp\left(\sum_{i=1}^n s_i u_i\right) = 2\sum_{j=1}^{n-1} \left\langle u_j, \sum_{i=1}^n s_i u_i \right\rangle u_j \\ &= 2\sum_{i=1}^n \left(\sum_{j=1}^{n-1} s_i \langle u_j, u_i \rangle u_j - s_i \langle u_n, u_i \rangle u_n \right) \\ &= 2\sum_{i=1}^n s_i u_i - 2\left\langle u_n, \sum_{i=1}^n s_i u_i \right\rangle u_n = 2v - 2\langle u_n, v \rangle u_n, \end{aligned}$$

where we use (10.2). Writing $u_n = e$, we conclude that for the orthogonal reflection along a hyperplane $U \subset V$ we have the formula

$$r_U(v) = v - 2\langle e, v \rangle e,$$

where the vector $e \in V$ satisfies $\langle e, e \rangle = 1$ and $U^\perp = \text{span}\{e\}$.

We can now verify that r_U is an orthogonal transformation. For all vectors $u, v \in V$ we have

$$\begin{aligned} \langle r_U(u), r_U(v) \rangle &= \langle u - 2\langle e, u \rangle e, v - 2\langle e, v \rangle e \rangle \\ &= \langle u, v \rangle - 2\langle u, e \rangle \langle e, v \rangle - 2\langle e, u \rangle \langle e, v \rangle + 4\langle e, u \rangle \langle e, v \rangle \langle e, e \rangle \\ &= \langle u, v \rangle, \end{aligned}$$

where we use that $\langle \cdot, \cdot \rangle$ is bilinear, symmetric and that $\langle e, e \rangle = 1$. We conclude that r_U is an orthogonal transformation.

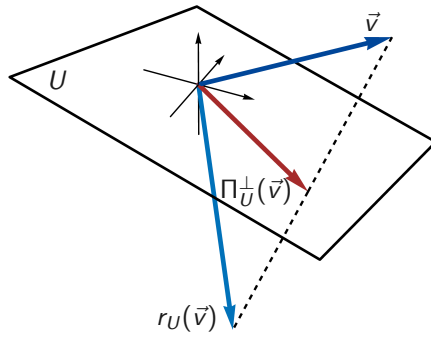
Finally, observe that with respect to the ordered basis $\mathbf{b} = (u_1, \dots, u_{n-1}, u_n)$ of V we have

$$\mathbf{M}(r_U, \mathbf{b}, \mathbf{b}) = \begin{pmatrix} \mathbf{1}_{n-1} & \\ & -1 \end{pmatrix}.$$

Indeed, since $u_i \in U$ for all $1 \leq i \leq n-1$, we obtain $r_U(u_i) = 2\Pi_U^\perp(u_i) - u_i = 2u_i - u_i = u_i$. Furthermore, $r_U(u_n) = u_n - 2\langle u_n, u_n \rangle u_n = -u_n$, as claimed. We conclude that $\det r_U = -1$.

Definition 10.41 (Special orthogonal group & special orthogonal matrices)

- Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space. The subset of $O(V, \langle \cdot, \cdot \rangle)$ consisting of endomorphisms whose determinant is 1 is called the *special orthogonal group* of $(V, \langle \cdot, \cdot \rangle)$ and is denoted by $SO(V, \langle \cdot, \cdot \rangle)$.
- A matrix $\mathbf{R} \in M_{n,n}(\mathbb{R})$ is called *special orthogonal* if $\mathbf{R} \in O(n)$ and $\det \mathbf{R} = 1$. The set of special orthogonal $n \times n$ -matrices is denoted by $SO(n)$ and called the *special orthogonal group*.

FIGURE 10.5. Orthogonal reflection along the plane U in $\mathbb{E}^3$.

Example 10.42 (The group $O(2)$) Recall from the exercises that if a matrix $\mathbf{R} \in M_{2,2}(\mathbb{R})$ satisfies $\mathbf{R}^T \mathbf{R} = \mathbf{1}_2$, then it is either of the form

$$\begin{pmatrix} a & -b \\ b & a \end{pmatrix} \quad \text{or} \quad \begin{pmatrix} a & b \\ b & -a \end{pmatrix}$$

for some real numbers a, b . The condition $\mathbf{R}^T \mathbf{R} = \mathbf{1}_2$ implies that $a^2 + b^2 = 1$, hence we can write $a = \cos \alpha$ and $b = \sin \alpha$ for some $\alpha \in \mathbb{R}$. In the second case we have $\det \mathbf{R} = -a^2 - b^2 = -1$, thus

$$SO(2) = \left\{ \mathbf{R}_\alpha = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \mid \alpha \in \mathbb{R} \right\}.$$

Recall also that the mapping $f_{\mathbf{R}_\alpha} : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ is the counter-clockwise rotation around $0_{\mathbb{R}^2}$ by the angle α . In the second case we obtain

$$\begin{pmatrix} \cos \alpha & \sin \alpha \\ \sin \alpha & -\cos \alpha \end{pmatrix} = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

The matrix

$$\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

is the matrix representation of the orthogonal reflection along the x -axis in $\mathbb{E}^2$ with respect to the standard ordered basis $\mathbf{e}$ of $\mathbb{R}^2$. We thus obtain a complete picture of all orthogonal transformations of $\mathbb{E}^2$. An orthogonal transformation of $\mathbb{E}^2$ is either a special orthogonal transformation in which case it is a rotation around $0_{\mathbb{R}^2}$ or else a composition of the orthogonal reflection along the x -axis and a rotation around $0_{\mathbb{R}^2}$.

We will discuss the structure of $O(n)$ for $n > 2$ below.

10.5 The adjoint mapping

In this section we discuss what one might consider to be the nicest endomorphisms of a Euclidean space $(V, \langle \cdot, \cdot \rangle)$, the so-called self-adjoint endomorphisms. Such endomorphisms are not only diagonalisable, but the basis of eigenvectors can be chosen to consist of orthonormal vectors with respect to $\langle \cdot, \cdot \rangle$.

Lemma 10.43 Let $(V, \langle \cdot, \cdot \rangle)$ and $(W, \langle \cdot, \cdot \rangle)$ be finite dimensional Euclidean spaces and $\mathbf{b} = (v_1, \dots, v_n)$ an orthonormal basis of $(V, \langle \cdot, \cdot \rangle)$ and $\mathbf{c} = (w_1, \dots, w_m)$ an

orthonormal basis of $(W, \langle\langle \cdot, \cdot \rangle\rangle)$. Then the matrix representation of a linear map $f : V \rightarrow W$ satisfies

$$[\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{ij} = \langle\langle f(v_j), w_i \rangle\rangle$$

for all $1 \leq i \leq m$ and for all $1 \leq j \leq n$.

Proof By definition, we have for all $1 \leq j \leq n$

$$f(v_j) = \sum_{k=1}^m [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{kj} w_k.$$

Hence for all $1 \leq i \leq m$, we obtain

$$\begin{aligned} \langle\langle f(v_j), w_i \rangle\rangle &= \left\langle\left\langle \sum_{k=1}^m [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{kj} w_k, w_i \right\rangle\right\rangle = \sum_{k=1}^m [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{kj} \langle\langle w_k, w_i \rangle\rangle \\ &= \sum_{k=1}^m [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{kj} \delta_{ki} = [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{ij}, \end{aligned}$$

as claimed. $\square$

Proposition 10.44 Let $(V, \langle \cdot, \cdot \rangle)$ and $(W, \langle\langle \cdot, \cdot \rangle\rangle)$ be finite dimensional Euclidean spaces and $f : V \rightarrow W$ a linear map. Then there exists a unique linear map $f^* : W \rightarrow V$ such that

$$\langle\langle f(v), w \rangle\rangle = \langle v, f^*(w) \rangle$$

for all $v \in V$ and $w \in W$.

Proof Let $\mathbf{b} = (v_1, \dots, v_n)$ be an orthonormal basis of $(V, \langle \cdot, \cdot \rangle)$ and $\mathbf{c} = (w_1, \dots, w_m)$ be an orthonormal basis of $(W, \langle\langle \cdot, \cdot \rangle\rangle)$. Let $f^* : W \rightarrow V$ be the unique linear map such that

$$\mathbf{M}(f^*, \mathbf{c}, \mathbf{b}) = \mathbf{M}(f, \mathbf{b}, \mathbf{c})^T.$$

Since $\langle \cdot, \cdot \rangle$ and $\langle\langle \cdot, \cdot \rangle\rangle$ are both bilinear it suffices to show that

$$\langle\langle f(v_j), w_i \rangle\rangle = \langle v_j, f^*(w_i) \rangle$$

for all $1 \leq j \leq n$ and all $1 \leq i \leq m$. By the previous lemma we have for all $1 \leq j \leq n$ and all $1 \leq i \leq m$

$$\langle\langle f(v_j), w_i \rangle\rangle = [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{ij} = [\mathbf{M}(f^*, \mathbf{c}, \mathbf{b})]_{ji} = \langle f^*(w_i), v_j \rangle = \langle v_j, f^*(w_i) \rangle.$$

This shows that $f^* : W \rightarrow V$ exists. Let $g : W \rightarrow V$ be another linear map such that

$$\langle\langle f(v), w \rangle\rangle = \langle v, g(w) \rangle$$

for all $v \in V$ and $w \in W$. Then we have for all $v \in V$ and $w \in W$

$$\langle v, f^*(w) - g(w) \rangle = \langle v, f^*(w) \rangle - \langle v, g(w) \rangle = \langle\langle f(v), w \rangle\rangle - \langle\langle f(v), w \rangle\rangle = 0.$$

This shows that for all $w \in W$ the vector $f^*(w) - g(w) \in V$ is orthogonal to all vectors of V . Since $\langle \cdot, \cdot \rangle$ is non-degenerate this implies that $f^*(w) - g(w) = 0_V$, that is, $f^*(w) = g(w)$ for all $w \in W$ and hence $f^* = g$. $\square$

Linear maps for which $f = f^*$ are of particular importance:

Definition 10.45 (Adjoint mapping, self-adjoint mappings and normal mappings)

- Let $(V, \langle \cdot, \cdot \rangle)$ and $(W, \langle\langle \cdot, \cdot \rangle\rangle)$ be finite dimensional Euclidean spaces and $f : V \rightarrow W$ a linear map. The unique mapping $f^* : W \rightarrow V$ guaranteed to exist by Proposition 10.44 is called the *adjoint mapping* of f .

- An endomorphism $f : V \rightarrow V$ of a Euclidean space $(V, \langle \cdot, \cdot \rangle)$ is called *self-adjoint* if $f = f^*$ and *normal* if $f \circ f^* = f^* \circ f$.

Example 10.46

- (i) The proof of [Proposition 10.44](#) implies that an endomorphism $f : V \rightarrow V$ of a finite dimensional Euclidean space $(V, \langle \cdot, \cdot \rangle)$ is self-adjoint if and only if its matrix representation with respect to an orthonormal basis $\mathbf{b}$ of V is symmetric. In particular, in $\mathbb{R}^n$ equipped with the standard scalar product, a mapping $f_{\mathbf{A}} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ for $\mathbf{A} \in M_{n,n}(\mathbb{R})$ is self-adjoint if and only if $\mathbf{A}$ is symmetric.
- (ii) Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space and $f : V \rightarrow V$ an orthogonal transformation. Then f is normal. Indeed, using that f is orthogonal, we obtain for all $u, v \in V$

$$\langle f^{-1}(u), v \rangle = \langle f(f^{-1}(u)), f(v) \rangle = \langle u, f(v) \rangle$$

so that the adjoint mapping of an orthogonal transformation is its inverse mapping, $f^* = f^{-1}$. It follows that $f \circ f^* = f \circ f^{-1} = \text{Id}_V = f^{-1} \circ f = f^* \circ f$ so that f is normal.

Exercises

Exercise 10.47 Verify that $\text{SO}(V, \langle \cdot, \cdot \rangle)$ is a subgroup of $\text{O}(V, \langle \cdot, \cdot \rangle)$ in the sense of [Definition 8.8](#). In particular, $\text{SO}(V, \langle \cdot, \cdot \rangle)$ is indeed a group and hence so is $\text{SO}(n)$.

10.6 The spectral theorem

We now come to one of the core results of the Linear Algebra II module:

Theorem 10.48 (The spectral theorem) *Let $f : V \rightarrow V$ be an endomorphism of the finite dimensional Euclidean space $(V, \langle \cdot, \cdot \rangle)$. Then there exists an orthonormal basis of V consisting of eigenvectors of f if and only if f is self-adjoint.*

For the proof of this statement we need two lemmas.

Lemma 10.49 *Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space of dimension $n \geq 1$ and $f : V \rightarrow V$ a self-adjoint endomorphism. Then f admits an eigenvalue $\lambda \in \mathbb{R}$.*

Proof Let $\mathbf{b}$ be an ordered orthonormal basis of $(V, \langle \cdot, \cdot \rangle)$ and $\mathbf{A} = \mathbf{M}(f, \mathbf{b}, \mathbf{b})$. Since f is self-adjoint, we have that $\mathbf{A} = \mathbf{A}^T$. Recall that the characteristic polynomial $\text{char}_f : \mathbb{R} \rightarrow \mathbb{R}$ of f satisfies $\text{char}_f(x) = \det(x\mathbf{1}_n - \mathbf{A})$ for all $x \in \mathbb{R}$. We may interpret each entry of $\mathbf{A}$ as a complex number and hence the characteristic polynomial as a function $\text{char}_f : \mathbb{C} \rightarrow \mathbb{C}$. In doing so, we can apply the fundamental theorem of algebra and conclude that there exists a complex number w such that $\text{char}_f(w) = 0$. We next argue that w has vanishing imaginary part and hence is a real number. Since $\det(w\mathbf{1}_n - \mathbf{A}) = 0$ we can find a non-zero vector $\vec{z} \in \mathbb{C}^n$ such that $\mathbf{A}\vec{z} = w\vec{z}$. We write $\vec{z} = \vec{x} + i\vec{y}$ for vectors $\vec{x}, \vec{y} \in \mathbb{R}^n$ and $w = s + it$ for real numbers s, t . Decomposing $\mathbf{A}(\vec{x} + i\vec{y}) = (s + it)(\vec{x} + i\vec{y})$ into real and imaginary parts, we obtain the equations

$$\begin{aligned}\mathbf{A}\vec{x} &= s\vec{x} - t\vec{y}, \\ \mathbf{A}\vec{y} &= s\vec{y} + t\vec{x}.\end{aligned}$$

Using the symmetry of $\mathbf{A}$, we compute

$$\langle \mathbf{A}\vec{x}, \vec{y} \rangle_{\mathbf{1}_n} = (\mathbf{A}\vec{x})^T \vec{y} = \vec{x}^T \mathbf{A}\vec{y} = \langle \vec{x}, \mathbf{A}\vec{y} \rangle_{\mathbf{1}_n}.$$

Using the above equations, we obtain

$$\begin{aligned}\langle \mathbf{A}\vec{x}, \vec{y} \rangle_{\mathbf{1}_n} &= \langle s\vec{x} - t\vec{y}, \vec{y} \rangle_{\mathbf{1}_n} = s\langle \vec{x}, \vec{y} \rangle_{\mathbf{1}_n} - t\|\vec{y}\|^2 = \langle \vec{x}, \mathbf{A}\vec{y} \rangle_{\mathbf{1}_n} = \langle \vec{x}, s\vec{y} + t\vec{x} \rangle_{\mathbf{1}_n} \\ &= s\langle \vec{x}, \vec{y} \rangle_{\mathbf{1}_n} + t\|\vec{x}\|^2,\end{aligned}$$

where $\|\cdot\|$ denotes the norm induced by the standard scalar product $\langle \cdot, \cdot \rangle_{\mathbf{1}_n}$ on $\mathbb{R}^n$. The last equation is equivalent to

$$0 = t(\|\vec{x}\|^2 + \|\vec{y}\|^2).$$

Since $\vec{z} \neq 0_{\mathbb{C}^n}$, the properties of the norm $\|\cdot\|$ – see [Proposition 10.12](#) – imply that $(\|\vec{x}\|^2 + \|\vec{y}\|^2) > 0$ and hence we must have $t = 0$. □

Recall that a subspace $U \subset V$ is said to be *stable under an endomorphism* $f : V \rightarrow V$ if $f(u) \in U$ for all $u \in U$.

Lemma 10.50 *Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space, $f : V \rightarrow V$ a self adjoint endomorphism and λ an eigenvalue of f . Then $(\text{Eig}_f(\lambda))^\perp$ is stable under f .*

Proof Write $U = \text{Eig}_f(\lambda)$ and let $w \in U^\perp$. Then, for all $u \in U$ we obtain

$$\langle u, f(w) \rangle = \langle u, f^*(w) \rangle = \langle f(u), w \rangle = \lambda \langle u, w \rangle,$$

where we use the self-adjointness of f and that u is an eigenvector of f . Since $w \in U^\perp$, we have $\langle u, w \rangle = 0$ and hence $\langle u, f(w) \rangle = 0$ for all $u \in U$. This shows that $f(w) \in U^\perp$, hence $U^\perp$ is stable under f . $\square$

Proof of Theorem 10.48 We first show that if f admits an orthonormal basis consisting of eigenvectors of f , then f must be self-adjoint. Let $\mathbf{b} = (u_1, \dots, u_n)$ be such a basis. We need to show that for all $v, w \in V$ we have

$$\langle f(v), w \rangle = \langle v, f(w) \rangle$$

There exist unique scalars $s_1, \dots, s_n \in \mathbb{R}$ and $t_1, \dots, t_n \in \mathbb{R}$ such that $v = \sum_{i=1}^n s_i u_i$ and $w = \sum_{j=1}^n t_j u_j$. From this we compute

$$\begin{aligned} \langle f(v), w \rangle &= \left\langle f \left(\sum_{i=1}^n s_i u_i \right), \sum_{j=1}^n t_j u_j \right\rangle = \sum_{i=1}^n \sum_{j=1}^n s_i t_j \langle f(u_i), u_j \rangle \\ &= \sum_{i=1}^n \sum_{j=1}^n s_i t_j \lambda_i \langle u_i, u_j \rangle = \sum_{i=1}^n \sum_{j=1}^n s_i t_j \lambda_i \delta_{ij} = \sum_{i=1}^n s_i t_i \lambda_i, \end{aligned}$$

where $\lambda_i \in \mathbb{R}$ denotes the eigenvalue of the eigenvector u_i for $i = 1, \dots, n$. Likewise we have

$$\langle v, f(w) \rangle = \sum_{i=1}^n \sum_{j=1}^n s_i t_j \langle u_i, f(u_j) \rangle = \sum_{i=1}^n \sum_{j=1}^n s_i t_j \lambda_j \langle u_i, u_j \rangle = \sum_{i=1}^n s_i t_i \lambda_i,$$

as claimed.

Conversely, assume that f is self-adjoint. We will use induction on the dimension n of V to show that $(V, \langle \cdot, \cdot \rangle)$ admits an orthonormal basis consisting of eigenvector of f . For $n = 1$ every endomorphism is diagonal, hence there is nothing to show and the statement is anchored.

Inductive step: Assume $n \geq 2$ and that the statement is true for all Euclidean spaces of dimension at most $n - 1$. By Lemma 10.49 the self-adjoint endomorphism $f : V \rightarrow V$ admits an eigenvalue $\lambda \in \mathbb{R}$. Write $U = \text{Eig}_f(\lambda)$. By Remark 10.16 we have $V = U \oplus U^\perp$ and by Lemma 10.50 we have that $U^\perp$ is stable under f . We thus obtain a linear map $\hat{f} = f|_{U^\perp} : U^\perp \rightarrow U^\perp$ by restricting f to $U^\perp$. Recall that the restriction $\langle \cdot, \cdot \rangle|_{U^\perp}$ of $\langle \cdot, \cdot \rangle$ to $U^\perp$ turns $(U^\perp, \langle \cdot, \cdot \rangle|_{U^\perp})$ into another Euclidean space. Since $\dim U \geq 1$, the dimension of $U^\perp$ is at most $n - 1$. The self-adjointness condition $f(v) = f^*(v)$ must hold for all vectors $v \in V$ and hence in particular also for all vectors of $U^\perp \subset V$. It follows that $\hat{f} : U^\perp \rightarrow U^\perp$ is self-adjoint with respect to $\langle \cdot, \cdot \rangle|_{U^\perp}$. Write $k = \dim U^\perp$. By the induction hypothesis there exists an orthonormal basis $\{u_1, \dots, u_k\}$ consisting of eigenvectors of $\hat{f}$. Since $\hat{f} = f|_{U^\perp}$, the vectors $\{u_1, \dots, u_k\}$ are also eigenvectors of f and since the inner product of vectors in $U^\perp$ is the same as the inner product computed in V , it follows that $\{u_1, \dots, u_k\}$ is orthonormal with respect to $\langle \cdot, \cdot \rangle$. Finally, using Gram-Schmidt orthonormalisation (Theorem 10.22), we can find an orthonormal basis $\{v_1, \dots, v_{n-k}\}$ of $U = \text{Eig}_f(\lambda)$ consisting of eigenvectors with eigenvalue λ . It follows that $\{u_1, \dots, u_k, v_1, \dots, v_{n-k}\}$ is an orthonormal basis of V consisting of eigenvectors of f . $\square$

Again, there is a matrix version of Theorem 10.48:

Theorem 10.51 (Matrix version of the spectral theorem) *Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{R})$ be a matrix. Then there exists an orthogonal matrix $\mathbf{R} \in M_{n,n}(\mathbb{R})$ such that $\mathbf{R}\mathbf{A}\mathbf{R}^T$ is a diagonal matrix if and only if $\mathbf{A}$ is symmetric.*

Proof We first show that if there exists an orthogonal matrix $\mathbf{R} \in M_{n,n}(\mathbb{R})$ such that $\mathbf{R}\mathbf{A}\mathbf{R}^T = \mathbf{D}$ for some diagonal matrix $\mathbf{D} \in M_{n,n}(\mathbb{R})$, then $\mathbf{A}$ must be symmetric. Since $\mathbf{A} = \mathbf{R}^T\mathbf{D}\mathbf{R}$ we obtain

$$\mathbf{A}^T = (\mathbf{R}^T\mathbf{D}\mathbf{R})^T = \mathbf{R}^T\mathbf{D}^T\mathbf{R} = \mathbf{R}^T\mathbf{D}\mathbf{R} = \mathbf{A},$$

where we use $\mathbf{D}^T = \mathbf{D}$ and [Remark 2.18](#).

For the converse direction consider $V = \mathbb{R}^n$ equipped with its standard scalar product $\langle \cdot, \cdot \rangle$. Since $\mathbf{A}$ is symmetric, the endomorphism $f_{\mathbf{A}} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ is self-adjoint with respect to $\langle \cdot, \cdot \rangle$. Applying [Theorem 10.48](#) we can thus find an ordered orthonormal basis $\mathbf{b}$ of $\mathbb{R}^n$ consisting of eigenvectors of $f_{\mathbf{A}}$. Denoting by $\mathbf{e}$ the standard ordered basis of $\mathbb{R}^n$, we have by [Theorem 3.107](#)

$$\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \mathbf{C}(\mathbf{e}, \mathbf{b})\mathbf{M}(f_{\mathbf{A}}, \mathbf{e}, \mathbf{e})\mathbf{C}(\mathbf{e}, \mathbf{b})^{-1}.$$

The basis $\mathbf{b}$ consists of eigenvectors of $f_{\mathbf{A}}$, hence $\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b})$ is a diagonal matrix by [Remark 6.30](#). Now recall from [Example 3.96](#) that $\mathbf{M}(f_{\mathbf{A}}, \mathbf{e}, \mathbf{e}) = \mathbf{A}$, thus writing $\mathbf{R} = \mathbf{C}(\mathbf{e}, \mathbf{b})$, we conclude that $\mathbf{R}\mathbf{A}\mathbf{R}^{-1}$ is diagonal. The standard ordered basis $\mathbf{e}$ of $\mathbb{R}^n$ is orthonormal with respect to the standard scalar product of $\mathbb{R}^n$, hence [Corollary 10.37](#) implies that $\mathbf{R}$ is orthogonal, $\mathbf{R}^{-1} = \mathbf{R}^T$. We have thus found an orthogonal matrix $\mathbf{R}$ so that $\mathbf{R}\mathbf{A}\mathbf{R}^T$ is diagonal. $\square$

10.6.1 Geometric description of self-adjoint endomorphisms

The spectral theorem tells us that self-adjoint endomorphisms can be diagonalised with an orthonormal basis. As a consequence one can give a precise geometric description of self-adjoint mappings. A first key observation towards this end is the following:

Lemma 10.52 *Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space and $f : V \rightarrow V$ a self-adjoint endomorphism. Then the eigenspaces of f are orthogonal. That is, for eigenvalues $\lambda \neq \mu$ of f we have $\langle u, v \rangle = 0$ for all $u \in \text{Eig}_f(\lambda)$ and for all $v \in \text{Eig}_f(\mu)$.*

Proof Let $u \in \text{Eig}_f(\lambda)$ and $v \in \text{Eig}_f(\mu)$. Then

$$\lambda \langle u, v \rangle = \langle f(u), v \rangle = \langle u, f(v) \rangle = \mu \langle u, v \rangle$$

and hence $0 = (\lambda - \mu) \langle u, v \rangle$. It follows that $\langle u, v \rangle = 0$ since $\lambda - \mu \neq 0$. $\square$

Recall that a vector space V is the direct sum of vector subspaces $U_1, \dots, U_k$ of V if every vector $v \in V$ can be written uniquely as a sum $v = u_1 + u_2 + \dots + u_k$ with $u_i \in U_i$ for $1 \leq i \leq k$. In this case we write $V = \bigoplus_{i=1}^k U_i$. In the presence of an inner product on V , we may ask that the subspaces U_i are all orthogonal:

Definition 10.53 (Orthogonal direct sum) Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space and $U_1, \dots, U_k$ be subspaces of V such that $V = \bigoplus_{i=1}^k U_i$. We say V is the *orthogonal direct sum* of the subspaces $U_1, \dots, U_k$ if for all $i \neq j$, we have $\langle u_i, u_j \rangle = 0$ for all $u_i \in U_i$ and for all $u_j \in U_j$. In this case we write

$$V = \bigoplus_{i=1}^k U_i^\perp.$$

Example 10.54

- (i) Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space and $U \subset V$ a subspace. Then V is the orthogonal direct sum of U and $U^\perp$.
- (ii) Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space and $\{u_1, \dots, u_n\}$ an orthogonal basis of V . Then V is the orthogonal direct sum of the subspaces $U_i = \text{span}\{u_i\}$ for $1 \leq i \leq n$.

Proposition 10.55 Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space and $f : V \rightarrow V$ a self-adjoint endomorphism. Let $\{\lambda_1, \dots, \lambda_k\}$ denote the eigenvalues of f . Then

$$V = \bigoplus_{i=1}^k \text{Eig}_f(\lambda_i)^\perp.$$

Proof By [Proposition 6.46](#) the eigenspaces of f are in direct sum and by [Lemma 10.52](#) this direct sum is orthogonal with respect to $\langle \cdot, \cdot \rangle$. By [Theorem 10.48](#) f is diagonalisable, hence

$$V = \bigoplus_{i=1}^k \text{Eig}_f(\lambda_i). \quad \square$$

We now obtain the aforementioned geometric description: A self adjoint endomorphism of a finite dimensional vector space is a linear combination of orthogonal projections.

Proposition 10.56 Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space and $f : V \rightarrow V$ a self-adjoint endomorphism with eigenvalues $\{\lambda_1, \dots, \lambda_k\}$. Then we have for all $v \in V$

$$f(v) = \sum_{i=1}^k \lambda_i \Pi_{U_i}^\perp(v),$$

where we write $U_i = \text{Eig}_f(\lambda_i)$.

Proof Let $g : V \rightarrow V$ be the endomorphism defined by the rule $g(v) = \sum_{i=1}^k \lambda_i \Pi_{U_i}^\perp(v)$ for all $v \in V$. We want to show that $f(v) = g(v)$ for all $v \in V$. Recall that for an orthogonal projection onto a subspace $U \subset V$ we have

$$\Pi_U^\perp(v) = \begin{cases} v & v \in U, \\ 0_V & v \in U^\perp. \end{cases}$$

Let $j \in \{1, \dots, k\}$ and $v \in U_j = \text{Eig}_f(\lambda_j)$. By [Lemma 10.52](#) we have $U_j \subset U_i^\perp$ for all $i \in \{1, \dots, k\}$ with $j \neq i$. Therefore,

$$g(v) = \sum_{i=1}^k \lambda_i \Pi_{U_i}^\perp(v) = \lambda_j \Pi_{U_j}^\perp(v) = \lambda_j v = f(v)$$

and the two mappings agree on all eigenspaces. Since $V = \bigoplus_{i=1}^k \text{Eig}_f(\lambda_i)$, the claim follows. $\square$

10.7 Quadratic forms

Closely related to the notion of a symmetric bilinear form is that of a quadratic form.

Definition 10.57 (Quadratic form) A function $q : V \rightarrow \mathbb{R}$ is called a *quadratic form* on V if there exists a symmetric bilinear form $\langle \cdot, \cdot \rangle$ on V such that

$$q(v) = \langle v, v \rangle$$

for all $v \in V$.

Remark 10.58

- The adjective quadratic is used since a quadratic form $q : V \rightarrow \mathbb{R}$ is so-called 2-homogeneous, that is, it satisfies

$$q(sv) = s^2 q(v)$$

for all $s \in \mathbb{R}$ and $v \in V$.

- By definition, every symmetric bilinear form $\langle \cdot, \cdot \rangle$ on V gives rise to a quadratic form q . The mapping $\langle \cdot, \cdot \rangle \mapsto q$ from the set of symmetric bilinear forms into the set of quadratic forms is thus surjective. That this mapping is also injective is a consequence of the so-called *polarisation identity*

$$4\langle v_1, v_2 \rangle = \langle v_1 + v_2, v_1 + v_2 \rangle - \langle v_1 - v_2, v_1 - v_2 \rangle$$

which holds for all $v_1, v_2 \in V$. Written in terms of the quadratic form associated to $\langle \cdot, \cdot \rangle$, it becomes

$$4\langle v_1, v_2 \rangle = q(v_1 + v_2) - q(v_1 - v_2).$$

Therefore, if two symmetric bilinear forms define the same quadratic form, then they must agree.

Example 10.59

- (i) Consider $V = \mathbb{R}^2$. The function

$$q : \mathbb{R}^2 \rightarrow \mathbb{R}, \quad \vec{v} = \begin{pmatrix} x \\ y \end{pmatrix} \mapsto q(\vec{v}) = 2x^2 - 4xy + 5y^2$$

is a quadratic form. Indeed, we have $q(\vec{v}) = \langle \vec{v}, \vec{v} \rangle_{\mathbf{A}}$, where

$$\mathbf{A} = \begin{pmatrix} 2 & -2 \\ -2 & 5 \end{pmatrix}.$$

- (ii) Likewise, the function

$$q : \mathbb{R}^3 \rightarrow \mathbb{R}, \quad \vec{v} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mapsto q(\vec{v}) = 4xy - 6yz + z^2$$

is a quadratic form. Indeed, we have $q(\vec{v}) = \langle \vec{v}, \vec{v} \rangle_{\mathbf{A}}$, where

$$\mathbf{A} = \begin{pmatrix} 0 & 2 & 0 \\ 2 & 0 & -3 \\ 0 & -3 & 1 \end{pmatrix}.$$

Applying the spectral theorem [Theorem 10.48](#), we see that we can "diagonalise" quadratic forms.

Theorem 10.60 (Principal axes theorem) Let $(V, \langle \cdot, \cdot \rangle)$ be a Euclidean space of dimension $n \in \mathbb{N}$ and $q : V \rightarrow \mathbb{R}$ a quadratic form. Then there exists an orthonormal ordered basis $\mathbf{b} = (v_1, \dots, v_n)$ of V with corresponding linear coordinate system $\beta : V \rightarrow \mathbb{R}^n$ and a diagonal matrix $\mathbf{D} \in M_{n,n}(\mathbb{R})$ such that for all $v \in V$

$$q(v) = \beta(v)^T \mathbf{D} \beta(v).$$

Remark 10.61 The lines spanned by the vectors v_i for $1 \leq i \leq n$ of the orthonormal basis are known as the *principal axes of the quadratic form q* . We will explain this terminology below.

Proof of Theorem 10.60 Fix an orthonormal ordered basis $\mathbf{b}'$ of $(V, \langle \cdot, \cdot \rangle)$ and let $\langle\langle \cdot, \cdot \rangle\rangle$ denote the symmetric bilinear form on V such that $q(v) = \langle\langle v, v \rangle\rangle$ for all $v \in V$. Let $\mathbf{A} = \mathbf{M}(\langle\langle \cdot, \cdot \rangle\rangle, \mathbf{b}')$ and $f : V \rightarrow V$ denote the endomorphism whose matrix representation is $\mathbf{A}$ with respect to the ordered basis $\mathbf{b}'$ of V . Since $\langle\langle \cdot, \cdot \rangle\rangle$ is a symmetric bilinear form, the matrix $\mathbf{A}$ is symmetric and hence f is self-adjoint with respect to $\langle \cdot, \cdot \rangle$ by [Example 10.46](#). [Theorem 10.48](#) implies that there exists an orthonormal ordered basis $\mathbf{b}$ of $(V, \langle \cdot, \cdot \rangle)$ consisting of eigenvectors of f . Let $\mathbf{D} = \mathbf{M}(f, \mathbf{b}, \mathbf{b})$ be the diagonal matrix representation of f with respect to $\mathbf{b}$. From [Proposition 9.6](#) we have for all $v \in V$

$$(10.4) \quad q(v) = \langle\langle v, v \rangle\rangle = \beta(v)^T \mathbf{M}(\langle\langle \cdot, \cdot \rangle\rangle, \mathbf{b}) \beta(v).$$

By construction we have $\mathbf{M}(\langle\langle \cdot, \cdot \rangle\rangle, \mathbf{b}') = \mathbf{M}(f, \mathbf{b}', \mathbf{b}')$, hence [Proposition 9.6](#) gives

$$\mathbf{M}(\langle\langle \cdot, \cdot \rangle\rangle, \mathbf{b}) = \mathbf{C}^T \mathbf{M}(\langle\langle \cdot, \cdot \rangle\rangle, \mathbf{b}') \mathbf{C} = \mathbf{C}^T \mathbf{M}(f, \mathbf{b}', \mathbf{b}') \mathbf{C},$$

where $\mathbf{C} = \mathbf{C}(\mathbf{b}, \mathbf{b}')$. Since both $\mathbf{b}'$ and $\mathbf{b}$ are ordered basis that are orthonormal with respect to $\langle \cdot, \cdot \rangle$, [Proposition 10.36](#) implies that $\mathbf{C}$ is orthogonal, $\mathbf{C}^T = \mathbf{C}^{-1}$. Finally, using [Theorem 3.107](#), we thus obtain

$$(10.5) \quad \mathbf{M}(\langle\langle \cdot, \cdot \rangle\rangle, \mathbf{b}) = \mathbf{C}^{-1} \mathbf{M}(f, \mathbf{b}', \mathbf{b}') \mathbf{C} = \mathbf{M}(f, \mathbf{b}, \mathbf{b}) = \mathbf{D}.$$

Combining [\(10.4\)](#) and [\(10.5\)](#), we get

$$q(v) = \beta(v)^T \mathbf{D} \beta(v),$$

as claimed. □

Example 10.62 ([Example 10.59](#) (i) continued) Here we are in the case where $V = \mathbb{R}^2$ and $\langle \cdot, \cdot \rangle$ is the standard scalar product. We have $q(\vec{v}) = \langle\langle \vec{v}, \vec{v} \rangle\rangle = \langle \vec{v}, \vec{v} \rangle_{\mathbf{A}}$. Taking $\mathbf{b}' = \mathbf{e}$ to be the orthonormal standard ordered basis of $\mathbb{R}^2$, we get

$$\mathbf{M}(\langle\langle \cdot, \cdot \rangle\rangle, \mathbf{b}') = \mathbf{A} = \begin{pmatrix} 2 & -2 \\ -2 & 5 \end{pmatrix}.$$

Orthonormal eigenvectors of $\mathbf{A}$ can be computed to be

$$\mathbf{b} = (v_1, v_2) = \left(\begin{pmatrix} -\frac{1}{\sqrt{5}} \\ \frac{2}{\sqrt{5}} \end{pmatrix}, -\begin{pmatrix} \frac{2}{\sqrt{5}} \\ \frac{1}{\sqrt{5}} \end{pmatrix} \right)$$

so that

$$\mathbf{C}(\mathbf{b}, \mathbf{b}') = \begin{pmatrix} -\frac{1}{\sqrt{5}} & -\frac{2}{\sqrt{5}} \\ \frac{2}{\sqrt{5}} & -\frac{1}{\sqrt{5}} \end{pmatrix}$$

and

$$\mathbf{C}^T \mathbf{A} \mathbf{C} = \begin{pmatrix} 6 & 0 \\ 0 & 1 \end{pmatrix} = \mathbf{D}.$$

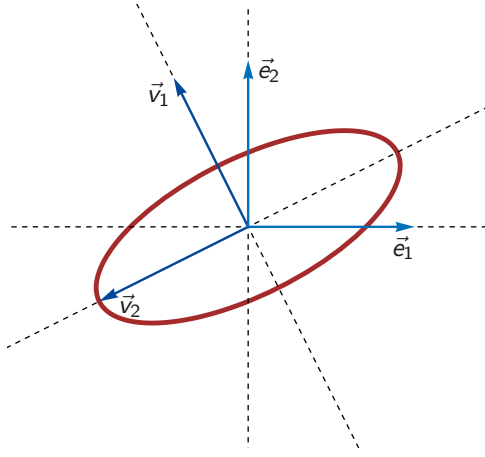


FIGURE 10.6. The ellipse defined by the equation $2x^2 - 4xy + 5y^2 = 1$ and its principal axes spanned by the orthonormal vectors $\vec{v}_1$ and $\vec{v}_2$.

Writing

$$\vec{v} = \begin{pmatrix} x \\ y \end{pmatrix} \quad \text{and} \quad \beta(\vec{v}) = \begin{pmatrix} X(\vec{v}) \\ Y(\vec{v}) \end{pmatrix},$$

we obtain

$$X(\vec{v}) = \begin{pmatrix} x \\ y \end{pmatrix}^T \begin{pmatrix} -\frac{1}{\sqrt{5}} \\ \frac{2}{\sqrt{5}} \end{pmatrix} = -\frac{x}{\sqrt{5}} + \frac{2y}{\sqrt{5}}$$

and

$$Y(\vec{v}) = -\begin{pmatrix} x \\ y \end{pmatrix}^T \begin{pmatrix} \frac{2}{\sqrt{5}} \\ \frac{1}{\sqrt{5}} \end{pmatrix} = -\frac{2x}{\sqrt{5}} - \frac{y}{\sqrt{5}},$$

so that

$$q(\vec{v}) = 2x^2 - 4xy + 5y^2 = 6X(\vec{v})^2 + Y(\vec{v})^2.$$

Remark 10.63 Especially in the physics literature it is customary to also use the letters x, y to denote functions from $\mathbb{R}^2 \rightarrow \mathbb{R}$ (and likewise for higher dimensions). The function x returns the first component of a vector $\vec{v} \in \mathbb{R}^2$ and y returns the second component, so that for instance

$$x\left(\begin{pmatrix} 2 \\ -4 \end{pmatrix}\right) = 2 \quad \text{and} \quad y\left(\begin{pmatrix} 3 \\ 5 \end{pmatrix}\right) = 5.$$

Thinking of x, y as functions – and doing the same for X, Y , the quadratic form from the previous example can then be written as (notice that we write q and not $q(\vec{v})$)

$$q = 2x^2 - 4xy + 5y^2 = 6X^2 + Y^2.$$

Definition 10.64 (Quadric) Let $q : V \rightarrow \mathbb{R}$ be a quadratic form and $c \in \mathbb{R}$. A *quadric* Q in V is the set of solutions $v \in V$ to an equation of the form $q(v) = c$.

Example 10.65 The set

$$Q = \{(x, y) \in \mathbb{R}^2 \mid 2x^2 - 4xy + 5y^2 = 1\}$$

is a quadric in $\mathbb{R}^2$. Written this way it is not immediately clear how the set of solution looks like. With respect to our new orthonormal basis $\mathbf{b} = (v_1, v_2)$ provided by the example above, we can however write Q as

$$Q = \{\vec{v} \in \mathbb{R}^2 \mid 6X(\vec{v})^2 + Y(\vec{v})^2 = 1\}$$

and we recognise Q as an ellipse. The X -axis spanned by v_1 and the Y -axis spanned by v_2 are symmetry axes for the ellipse and are known as its *principal axes*, see Figure 10.6.

Remark 10.66 (♥ - not examinable) Quadratic forms also play an important role in calculus. Let $f : \mathbb{R}^n \rightarrow \mathbb{R}$ be a twice continuously differentiable function. The Hessian matrix of f at $\vec{x} = (x_i)_{1 \leq i \leq n} \in \mathbb{R}^n$ is given by

$$[\mathbf{H}_f(\vec{x})]_{ij} = \frac{\partial^2 f}{\partial x_i \partial x_j}$$

where $1 \leq i, j \leq n$. By the Schwartz theorem, this matrix is symmetric and hence for each $\vec{x} \in \mathbb{R}^n$ we obtain a quadratic form on $\mathbb{R}^n$ defined by the rule

$$q(\vec{h}) = \frac{1}{2} \vec{h}^T \mathbf{H}_f(\vec{x}) \vec{h} = \frac{1}{2} \langle \vec{h}, \vec{h} \rangle_{\mathbf{H}_f(\vec{x})}.$$

for all $\vec{h} \in \mathbb{R}^n$ and where $\langle \cdot, \cdot \rangle$ denotes the standard scalar product of $\mathbb{R}^n$. The significance of this quadratic form arises from the Taylor approximation of f . For vectors $\vec{h} \in \mathbb{R}^n$ of small length we have the approximation

$$f(\vec{x} + \vec{h}) \approx f(\vec{x}) + \langle \nabla f(\vec{x}), \vec{h} \rangle + \frac{1}{2} \langle \vec{h}, \vec{h} \rangle_{\mathbf{H}_f(\vec{x})},$$

where $\nabla f(\vec{x})$ denotes the gradient of f at $\vec{x}$. Recall that at a critical point $\vec{x}$ of f we have $\nabla f(\vec{x}) = 0_{\mathbb{R}^n}$ and hence

$$f(\vec{x} + \vec{h}) \approx f(\vec{x}) + q(\vec{h}).$$

In order to decide whether f admits a local maximum / a local minimum at a critical point, one thus needs to investigate the sign of $q(\vec{h})$ for all $\vec{h}$.

The previous remark is one motivation for the following definition:

Definition 10.67 Let $q : V \rightarrow \mathbb{R}$ be a quadratic form on the $\mathbb{R}$ -vector space V . Then q is called

- *positive or positive semi-definite* if $q(v) \geq 0$ for all $v \in V$;
- *positive definite* if $q(v) \geq 0$ and $q(v) = 0$ if and only if $v = 0_V$;
- *negative or negative semi-definite* if $q(v) \leq 0$ for all $v \in V$;
- *negative definite* if $q(v) \leq 0$ and $q(v) = 0$ if and only if $v = 0_V$;
- *indefinite* if there exists $v \in V$ and $w \in V$ such that $q(v) < 0$ and $q(w) > 0$.

By the principal axes theorem (Theorem 10.60), we can write a quadratic form $q : V \rightarrow \mathbb{R}$ on a Euclidean space $(V, \langle \cdot, \cdot \rangle)$ as $q(v) = \beta(v)^T \mathbf{D} \beta(v)$, where $\mathbf{b}$ is an ordered orthonormal basis of V and $\mathbf{D}$ a diagonal matrix.

Exercises

Exercise 10.68 Show the following characterisations:

- (i) q is positive if and only if all diagonal entries of $\mathbf{D}$ are greater than or equal to zero;
- (ii) q is positive definite if and only if all diagonal entries of $\mathbf{D}$ are positive;
- (iii) q is negative if and only if all diagonal entries of $\mathbf{D}$ are less than or equal to zero;
- (iv) q is negative definite if and only if all diagonal entries of $\mathbf{D}$ are negative;
- (v) q is indefinite if and only if $\mathbf{D}$ has positive and negative diagonal entries.

Unitary spaces

WEEK 8

Unitary spaces are the complex companions of Euclidean spaces. Much of the theory of Euclidean spaces also holds over to the complex numbers when we suitably adapt the notion of an inner product. In addition, almost all proofs carry over from the real case, hence we will only provide proofs when the arguments from the real case do not work.

11.1 Hermitian inner products

Naively one might define a “standard scalar product” on $\mathbb{C}^n$ as in the case of $\mathbb{R}^n$, that is, for $\vec{z} = (z_i)_{1 \leq i \leq n}$ and $\vec{w} = (w_i)_{1 \leq i \leq n} \in \mathbb{C}^n$ we put $\vec{z} \cdot \vec{w} = \sum_{i=1}^n z_i w_i$. However, doing so, it is not true any more that $\vec{z} \cdot \vec{z} = 0$ only for the zero vector in $\mathbb{C}^n$. For instance, the vector

$$\vec{z} = \begin{pmatrix} 1 \\ i \end{pmatrix}$$

satisfies $\vec{z} \cdot \vec{z} = 0$, but $\vec{z} \neq 0_{\mathbb{C}^2}$. Instead of the above definition we define the *Hermitian standard scalar product* on $\mathbb{C}^n$ by the rule

$$\langle \vec{z}, \vec{w} \rangle = \sum_{i=1}^n \bar{z}_i w_i,$$

where $\bar{z}$ denotes the complex conjugate of the complex number $z \in \mathbb{C}$. Recall that $z\bar{z} = \operatorname{Re}(z)^2 + \operatorname{Im}(z)^2 \geq 0$ so that $z\bar{z} = 0$ if and only if $z = 0$. The Hermitian standard scalar product is an example of a *sesquilinear form*:

Definition 11.1 (Sesquilinear form) Let V be a complex vector space. A *sesquilinear form* on V is a map $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{C}$ such that

(i) $\langle \cdot, \cdot \rangle$ is linear in the second variable, that is,

$$\langle v, s_1 w_1 + s_2 w_2 \rangle = s_1 \langle v, w_1 \rangle + s_2 \langle v, w_2 \rangle$$

for all $s_1, s_2 \in \mathbb{C}$ and all $v, w_1, w_2 \in V$;

(ii) $\langle \cdot, \cdot \rangle$ is *conjugate linear* in the first variable, that is,

$$\langle s_1 w_1 + s_2 w_2, v \rangle = \bar{s}_1 \langle w_1, v \rangle + \bar{s}_2 \langle w_2, v \rangle$$

for all $s_1, s_2 \in \mathbb{C}$ and all $v, w_1, w_2 \in V$;

Moreover, a sesquilinear form is called *Hermitian* if

$$\langle v, w \rangle = \overline{\langle w, v \rangle}$$

for all $v, w \in V$.

Remark 11.2

- Sesquilinear forms correspond to bilinear forms in the real setting and Hermitian forms correspond to symmetric bilinear forms.

- In our convention a sesquilinear form is conjugate linear in the first variable and linear in the second variable. The reader is warned that some authors use the opposite convention so that a sesquilinear form is linear in the first variable and conjugate linear in the second variable.

Let V be a finite dimensional $\mathbb{C}$ -vector space and $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V . As in the case of bilinear forms over real vector spaces, we define the matrix representation of a sesquilinear form $\langle \cdot, \cdot \rangle$ on V with respect to $\mathbf{b}$

$$\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) = (\langle v_i, v_j \rangle)_{1 \leq i, j \leq n}.$$

Recall that in the real setting symmetric bilinear forms are represented by symmetric matrices. Similarly, sesquilinear Hermitian forms – usually just called Hermitian forms – are represented by so-called Hermitian matrices. For a precise definition, we need:

Definition 11.3 (Conjugate matrix) Let $\mathbf{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n} \in M_{m,n}(\mathbb{C})$. The *conjugate matrix* of $\mathbf{A}$ is the matrix $\overline{\mathbf{A}} = (\overline{A_{ij}})_{1 \leq i \leq m, 1 \leq j \leq n} \in M_{m,n}(\mathbb{C})$ whose entries are the complex conjugates of the entries of $\mathbf{A}$.

Lemma 11.4 (Properties of the conjugate matrix)

(i) For all $\mathbf{A}, \mathbf{B} \in M_{m,n}(\mathbb{C})$ and all $s, t \in \mathbb{C}$, we have

$$\overline{s\mathbf{A} + t\mathbf{B}} = \overline{s}\overline{\mathbf{A}} + \overline{t}\overline{\mathbf{B}}, \quad \overline{\overline{\mathbf{A}}} = \mathbf{A}, \quad \overline{\mathbf{A}^T} = \overline{\mathbf{A}}^T.$$

(ii) For all $\mathbf{A} \in M_{m,n}(\mathbb{C})$ and $\mathbf{B} \in M_{n,p}(\mathbb{C})$, we have

$$\overline{\mathbf{AB}} = \overline{\mathbf{A}}\overline{\mathbf{B}}.$$

In particular, $\mathbf{A} \in M_{n,n}(\mathbb{C})$ is invertible if and only if $\overline{\mathbf{A}}$ is invertible and $(\overline{\mathbf{A}})^{-1} = \overline{\mathbf{A}^{-1}}$.

(iii) For all $\mathbf{A} \in M_{n,n}(\mathbb{C})$ we have

$$\overline{\det \mathbf{A}} = \det(\overline{\mathbf{A}}).$$

Proof (i) and (ii) follow from the definitions of matrix operations and from $\overline{\overline{z}} = z$ and $\overline{zw} = \overline{z}\overline{w}$ for all complex numbers z, w . (iii) follows from the Leibniz formula [Proposition 5.39](#). $\square$

Hermitian matrices have the property that their transpose equals their conjugate matrix.

Definition 11.5 (Hermitian matrix) A matrix $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n} \in M_{n,n}(\mathbb{C})$ is called *Hermitian* if

$$\mathbf{A}^T = \overline{\mathbf{A}} \iff \mathbf{A} = \overline{\mathbf{A}}^T \iff A_{ji} = \overline{A_{ij}}, \quad 1 \leq i, j \leq n.$$

Remark 11.6

- Notice that the diagonal entries of a Hermitian matrix $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n} \in M_{n,n}(\mathbb{C})$ satisfy $A_{ii} = \overline{A_{ii}}$ for all $1 \leq i \leq n$ and hence must be real.

- If we write $\mathbf{A} \in M_{n,n}(\mathbb{C})$ as $\mathbf{A} = \mathbf{B} + i\mathbf{C}$ for $\mathbf{B}, \mathbf{C} \in M_{n,n}(\mathbb{R})$, then $\mathbf{A}$ is Hermitian if and only if

$$\mathbf{A}^T = (\mathbf{B} + i\mathbf{C})^T = \mathbf{B}^T + i\mathbf{C}^T = \overline{\mathbf{A}} = \mathbf{B} - i\mathbf{C}$$

which is equivalent to $\mathbf{B}$ being symmetric and $\mathbf{C}$ being anti-symmetric.

Example 11.7 2×2 and 3×3 Hermitian matrices are of the form

$$\begin{pmatrix} a & z \\ \bar{z} & b \end{pmatrix}, \quad \begin{pmatrix} a & z & w \\ \bar{z} & b & u \\ \bar{w} & \bar{u} & c \end{pmatrix}$$

for $a, b, c \in \mathbb{R}$ and $u, z, w \in \mathbb{C}$.

In analogy to [Proposition 9.6](#) we obtain:

Proposition 11.8 Let V be a finite dimensional $\mathbb{C}$ -vector space and $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V with associated linear coordinate system $\beta : V \rightarrow \mathbb{K}^n$. Suppose $\langle \cdot, \cdot \rangle$ is a sesquilinear form on V , then

- (i) for all $v, w \in V$ we have

$$\langle v, w \rangle = \overline{\beta(v)}^T \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) \beta(w);$$

- (ii) $\langle \cdot, \cdot \rangle$ is Hermitian if and only if $\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b})$ is a Hermitian matrix;

- (iii) if $\mathbf{b}'$ is another ordered basis of V , then

$$\mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}') = \overline{\mathbf{C}}^T \mathbf{M}(\langle \cdot, \cdot \rangle, \mathbf{b}) \mathbf{C},$$

where $\mathbf{C} = \mathbf{C}(\mathbf{b}', \mathbf{b})$ denotes the change of basis matrix.

Proof Exercise. □

Non-degenerateness of a sesquilinear form is defined exactly as in the real case and correspondingly, a sesquilinear form on a finite dimensional complex vector space is non-degenerate if and only if its matrix representation with respect to some (and hence any) basis has non-vanishing determinant (c.f. [Proposition 9.10](#)).

Again, in analogy to the real case we call a sesquilinear form $\langle \cdot, \cdot \rangle$ on V *positive* if $\langle v, v \rangle \geq 0$ for all $v \in V$ and *positive definite* if $\langle \cdot, \cdot \rangle$ is positive and $\langle v, v \rangle = 0$ if and only if $v = 0_V$. Also, in analogy to [Definition 10.2](#), we define:

Definition 11.9 (Hermitian inner product) Let V be a $\mathbb{C}$ -vector space. A sesquilinear form on V that is positive definite and Hermitian is called a *Hermitian inner product*.

Example 11.10 (Hermitian forms and Hermitian inner products)

- (i) Suppose $\mathbf{A} \in M_{n,n}(\mathbb{C})$ is a Hermitian matrix, then the map $\langle \cdot, \cdot \rangle : \mathbb{C}^n \times \mathbb{C}^n \rightarrow \mathbb{C}$ defined by the rule

$$\langle \vec{z}, \vec{w} \rangle_{\mathbf{A}} = (\overline{\vec{z}})^T \mathbf{A} \vec{w}$$

for all $\vec{z}, \vec{w} \in \mathbb{C}^n$ defines a Hermitian form on $\mathbb{C}^n$.

- (ii) Let $a < b$ be real numbers and consider $V = C([a, b], \mathbb{C})$, the complex-vector space of continuous complex-valued functions on the interval $[a, b]$. We define $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{C}$ by the rule

$$\langle f, g \rangle = \int_a^b \overline{f(x)} g(x) dx.$$

Then the properties of integration from the Analysis module show that $\langle \cdot, \cdot \rangle$ is a Hermitian inner product on V .

- (iii) Let $V = M_{n,n}(\mathbb{C})$ denote the $\mathbb{C}$ -vector space of $n \times n$ -matrices with complex entries. We define a map $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{C}$ defined by the rule

$$\langle \mathbf{A}, \mathbf{B} \rangle = \text{Tr}(\overline{\mathbf{A}}^T \mathbf{B})$$

for all $\mathbf{A}, \mathbf{B} \in M_{n,n}(\mathbb{C})$. Since the trace is a linear map $\text{Tr} : M_{n,n}(\mathbb{C}) \rightarrow \mathbb{C}$ satisfying $\text{Tr}(\overline{\mathbf{A}}) = \overline{\text{Tr}(\mathbf{A})}$ for all $\mathbf{A} \in M_{n,n}(\mathbb{C})$, it follows that $\langle \cdot, \cdot \rangle$ is a Hermitian form on $M_{n,n}(\mathbb{C})$. Writing $\mathbf{A} = (A_{ij})_{1 \leq i, j \leq n}$, we obtain

$$\langle \mathbf{A}, \mathbf{A} \rangle = \sum_{i=1}^n \sum_{j=1}^n \overline{A_{ji}} A_{ji} = \sum_{i=1}^n \sum_{j=1}^n |A_{ji}|^2$$

so that $\langle \mathbf{A}, \mathbf{A} \rangle \geq 0$ and $\langle \mathbf{A}, \mathbf{A} \rangle = 0$ if and only if all entries of $\mathbf{A}$ are zero, that is, $\mathbf{A} = 0$. We conclude that $\langle \cdot, \cdot \rangle$ defines a Hermitian inner product on $M_{n,n}(\mathbb{C})$.

The complex companions of Euclidean spaces (c.f. [Definition 10.7](#)) are the so-called unitary spaces:

Definition 11.11 (Unitary space) A pair $(V, \langle \cdot, \cdot \rangle)$ consisting of an $\mathbb{C}$ -vector space V and a Hermitian inner product $\langle \cdot, \cdot \rangle$ on V is called a *unitary space*.

As in the case of Euclidean spaces, a Hermitian inner product $\langle \cdot, \cdot \rangle$ on a complex vector space V allows to define a norm $\| \cdot \| = \sqrt{\langle \cdot, \cdot \rangle}$ on V . Since $\langle \cdot, \cdot \rangle$ is a Hermitian form, we have that $\langle v, v \rangle = \overline{\langle v, v \rangle}$ for all $v \in V$. Therefore, $\langle v, v \rangle$ is a non-negative real number for all $v \in V$ and hence $\| \cdot \|$ is well defined. Although we will not prove it here, the Cauchy–Schwarz inequality holds as well in the setting of unitary spaces. That is, as in [Proposition 10.8](#), we have again that for all $v_1, v_2 \in V$

$$|\langle v_1, v_2 \rangle| \leq \|v_1\| \|v_2\|$$

with equality if and only if $\{v_1, v_2\}$ are linearly dependent. Here $|\cdot|$ on the left denotes the absolute value.

The distance function is also defined analogously and again we have the triangle inequality. Again, we will not prove this.

The notions of orthogonality, orthonormality, the orthogonal complement, the orthogonal projection onto a subspace are again defined analogously to the Euclidean case.

Example 11.12 Consider $V = C([0, 2\pi], \mathbb{C})$, the $\mathbb{C}$ -vector space of continuous complex-valued functions defined on the interval $[0, 2\pi]$. We equip V with the Hermitian inner product $\langle \cdot, \cdot \rangle$ as defined in [Example 11.10](#) above. For $n \in \mathbb{Z}$ let $f_n : [0, 2\pi] \rightarrow \mathbb{C}$ be defined by the rule

$$f_n(t) = \frac{1}{\sqrt{2\pi}} e^{int}$$

for all $t \in [0, 2\pi]$. Then for $n \neq m$, we obtain

$$\langle f_n, f_m \rangle = \frac{1}{2\pi} \int_0^{2\pi} \overline{e^{int}} e^{imt} dt = \frac{1}{2\pi} \int_0^{2\pi} e^{i(m-n)t} dt = \frac{1}{2\pi i(m-n)} e^{i(m-n)t} \Big|_0^{2\pi} = 0$$

and for all $n \in \mathbb{Z}$ we have that $\langle f_n, f_n \rangle = 1$. It follows that $\{f_n | n \in \mathbb{Z}\}$ is an orthonormal subset of V . This observation is at the heart of the theory of *Fourier series*.

Again, [Theorem 10.22](#) also has a complex version:

Theorem 11.13 (Gram–Schmidt orthonormalisation for unitary spaces) *Let $(V, \langle \cdot, \cdot \rangle)$ be an n -dimensional unitary space and $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V . For $2 \leq i \leq n$ we define recursively*

$$w_i = v_i - \Pi_{U_{i-1}}^\perp(v_i) \quad \text{and} \quad u_i = \frac{w_i}{\|w_i\|},$$

where $U_{i-1} = \text{span}\{u_1, \dots, u_{i-1}\}$ and $u_1 = v_1/\|v_1\|$. Then $\mathbf{b}' = (u_1, \dots, u_n)$ is well defined and an orthonormal ordered basis of V . Moreover, $\mathbf{b}'$ is the unique orthonormal ordered basis of V so that the change of basis matrix $\mathbf{C}(\mathbf{b}', \mathbf{b})$ is an upper triangular matrix whose diagonal entries are real and positive.

As in [Definition 10.25](#), we have:

Definition 11.14 (Positive definite matrix) *Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{C})$. The matrix $\mathbf{A}$ is called *positive definite* if the sesquilinear form $\langle \cdot, \cdot \rangle_{\mathbf{A}}$ on $\mathbb{C}^n$ is positive definite.*

As in [Theorem 10.26](#), we obtain:

Theorem 11.15 (Cholesky decomposition over $\mathbb{C}$) *Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{C})$ be a positive definite Hermitian matrix. Then there exists a unique upper triangular matrix $\mathbf{C} \in M_{n,n}(\mathbb{C})$ with real and positive diagonal entries such that $\mathbf{A} = \overline{\mathbf{C}}^T \mathbf{C}$.*

Remark 11.16 Similar to the real case (c.f. [Remark 10.27](#)), for an invertible complex matrix $\mathbf{C} \in M_{n,n}(\mathbb{C})$, the matrix $\overline{\mathbf{C}}^T \mathbf{C}$ is Hermitian and positive definite.

Remark 11.17 As in [Remark 10.28](#), in a finite dimensional unitary space $(V, \langle \cdot, \cdot \rangle)$ equipped with an ordered orthonormal basis $\mathbf{b} = (v_1, \dots, v_n)$, we have the following identities for all $v \in V$

$$v = \sum_{i=1}^n \langle v, v_i \rangle v_i \quad \text{and} \quad \|v\| = \sqrt{\sum_{i=1}^n \langle v, v_i \rangle^2}.$$

Exercises

Exercise 11.18 Compute the Cholesky decomposition of the positive definite Hermitian matrix

$$\mathbf{A} = \begin{pmatrix} 6 & -1+i & -2 \\ -1-i & 3 & -2+i \\ -2 & -2-i & 3 \end{pmatrix}.$$

11.2 The unitary group

Orthogonal transformations between Euclidean spaces correspond to so-called unitary transformations between unitary spaces:

Definition 11.19 (Unitary transformation) Let $(V, \langle \cdot, \cdot \rangle)$ and $(W, \langle \cdot, \cdot \rangle)$ be unitary spaces. An isomorphism $f : V \rightarrow W$ is called a *unitary transformation* if

$$\langle u, v \rangle = \langle f(u), f(v) \rangle$$

for all $u, v \in V$.

With this definition, all the statements about orthogonal transformations have corresponding statements for unitary transformations.

Definition 11.20 (Unitary group & unitary matrices)

- Let $(V, \langle \cdot, \cdot \rangle)$ be a unitary space. The set of unitary transformations from $(V, \langle \cdot, \cdot \rangle)$ to itself is called the *unitary group* of $(V, \langle \cdot, \cdot \rangle)$ and denoted by $U(V, \langle \cdot, \cdot \rangle)$.
- A matrix $\mathbf{R} \in M_{n,n}(\mathbb{C})$ is called *unitary* if $f_{\mathbf{R}} : \mathbb{C}^n \rightarrow \mathbb{C}^n$ is an unitary transformation of $(\mathbb{C}^n, \langle \cdot, \cdot \rangle)$, where $\langle \cdot, \cdot \rangle$ denotes the standard Hermitian scalar product of $\mathbb{C}^n$. The set of unitary $n \times n$ -matrices is denoted by $U(n)$ and called the *unitary group*.

Like the orthogonal group, the unitary group is indeed a group:

Proposition 11.21 Let $(V, \langle \cdot, \cdot \rangle)$ be a unitary space. Then the set $U(V, \langle \cdot, \cdot \rangle)$ is a group in the sense of [Definition 8.4](#) when the group operation is taken to be the composition of mappings. In particular, $U(n)$ is a group when the group operation is taken to be matrix multiplication.

We have the characterisation:

Lemma 11.22 For all $n \in \mathbb{N}$ we have

$$U(n) = \left\{ \mathbf{R} \in M_{n,n}(\mathbb{C}) \mid \overline{\mathbf{R}^T} \mathbf{R} = \mathbf{1}_n \right\} = \left\{ \mathbf{R} \in GL(n, \mathbb{C}) \mid \overline{\mathbf{R}^T} = \mathbf{R}^{-1} \right\}.$$

A unitary transformation has a unitary matrix representation with respect to an orthonormal basis:

Proposition 11.23 Let $n \in \mathbb{N}$ and $(V, \langle \cdot, \cdot \rangle)$ be an n -dimensional unitary space equipped with an orthonormal ordered basis $\mathbf{b}$. Then an endomorphism $f : V \rightarrow V$ is a unitary transformation if and only if its matrix representation $\mathbf{R} = \mathbf{M}(f, \mathbf{b}, \mathbf{b})$ with respect to $\mathbf{b}$ is a unitary matrix.

We also have:

Corollary 11.24 Let $n \in \mathbb{N}$ and $(V, \langle \cdot, \cdot \rangle)$ be an n -dimensional unitary space equipped with an orthonormal ordered basis $\mathbf{b}$. Then an ordered basis $\mathbf{b}'$ of V is orthonormal with respect to $\langle \cdot, \cdot \rangle$ if and only if the change of basis matrix $\mathbf{C}(\mathbf{b}', \mathbf{b})$ is unitary.

The special unitary transformations are those with determinant one:

Definition 11.25 (Special unitary group & special unitary matrices)

- Let $(V, \langle \cdot, \cdot \rangle)$ be a unitary space. The subset of $U(V, \langle \cdot, \cdot \rangle)$ consisting of endomorphism whose determinant is 1 is called the *special unitary group* of $(V, \langle \cdot, \cdot \rangle)$ and is denoted by $SU(V, \langle \cdot, \cdot \rangle)$.
- A matrix $\mathbf{R} \in M_{n,n}(\mathbb{C})$ is called *special unitary* if $\mathbf{R} \in U(n)$ and $\det \mathbf{R} = 1$. The set of special unitary $n \times n$ -matrices is denoted by $SU(n)$ and called the *special unitary group*.

Again, we have indeed groups:

Example 11.26 While $O(1)$ just consists of the matrices $\pm(1)$. The group $U(1)$ has infinitely many elements. Indeed $(z) \in U(1)$ if and only if $|z|^2 = 1$ so that

$$U(1) = \{ (e^{i\vartheta}) | \vartheta \in \mathbb{R} \}.$$

11.3 Adjoint and normal endomorphisms

The notion of the adjoint for maps between unitary spaces is defined as in the case of Euclidean spaces. Given finite dimensional unitary spaces $(V, \langle \cdot, \cdot \rangle)$ and $(W, \langle \cdot, \cdot \rangle)$ and a linear map $f : V \rightarrow W$, the adjoint of f is the unique map $f^* : W \rightarrow V$ such that

$$\langle f(v), w \rangle = \langle v, f^*(w) \rangle$$

for all $v \in V$ and $w \in W$. The adjoint f^* is constructed by choosing an orthonormal basis $\mathbf{b}$ of V and an orthonormal basis $\mathbf{c}$ of W and by requesting that

$$(11.1) \quad \mathbf{M}(f^*, \mathbf{c}, \mathbf{b}) = \overline{\mathbf{M}(f, \mathbf{b}, \mathbf{c})}^T.$$

The self-adjoint mappings of a unitary space $(V, \langle \cdot, \cdot \rangle)$ are then the linear maps $f : V \rightarrow V$ satisfying $f^* = f$. If we equip V with an ordered orthonormal basis $\mathbf{b}$, then a linear map $f : V \rightarrow V$ is self-adjoint if and only if $\mathbf{M}(f, \mathbf{b}, \mathbf{b})$ is a Hermitian matrix.

Let $\mathbf{A} \in M_{n,n}(\mathbb{C})$ and equip $\mathbb{C}^n$ with the standard Hermitian scalar product $\langle \cdot, \cdot \rangle$. Then (11.1) implies that $(f_{\mathbf{A}})^* = f_{\overline{\mathbf{A}}^T}$. This motivates the following definition:

Definition 11.27 (Adjoint matrix) For a matrix $\mathbf{A} \in M_{n,n}(\mathbb{C})$ we define

$$\mathbf{A}^* = \overline{\mathbf{A}}^T$$

and call $\mathbf{A}^*$ the *adjoint matrix* of $\mathbf{A}$.

The spectral theorem also holds in the unitary setting:

Theorem 11.28 (The spectral theorem for unitary spaces) *Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional unitary space and $f : V \rightarrow V$ a self-adjoint endomorphism. Then there exists an orthonormal basis of V consisting of eigenvectors of f . In particular, f is diagonalisable.*

Again we have a matrix version:

Theorem 11.29 *Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{C})$ be a Hermitian matrix. Then there exists a unitary matrix $\mathbf{R} \in M_{n,n}(\mathbb{C})$ such that $\mathbf{R}\mathbf{A}\mathbf{R}^*$ is a diagonal matrix.*

As in the real case we call an endomorphism $f : V \rightarrow V$ of a unitary space $(V, \langle \cdot, \cdot \rangle)$ *normal* if $f \circ f^* = f^* \circ f$. Normal endomorphisms can be characterised in terms of the following lemma:

Lemma 11.30 *Let $(V, \langle \cdot, \cdot \rangle)$ be a unitary space and $f : V \rightarrow V$ an endomorphism. Then f is normal if and only if*

$$\|f(v)\| = \|f^*(v)\|$$

for all $v \in V$.

Before we give a proof, we remark:

Remark 11.31 Let V be a finite dimensional $\mathbb{C}$ -vector space and $\langle \cdot, \cdot \rangle$ and Hermitian form on V . Similar to the real case, writing $q(v) = \langle v, v \rangle$, we obtain for all $v, w \in V$

$$\begin{aligned} 4 \operatorname{Re} \langle v, w \rangle &= 2(\langle v, w \rangle + \overline{\langle v, w \rangle}) = 2(\langle v, w \rangle + \langle w, v \rangle) \\ &= \langle v + w, v + w \rangle - \langle v - w, v - w \rangle = q(v + w) - q(v - w), \end{aligned}$$

so that the real part of $\langle \cdot, \cdot \rangle$ is determined by q . On the other hand, we have for all $v, w \in V$

$$\operatorname{Re}(\langle iv, w \rangle) = -\operatorname{Re}(i\langle v, w \rangle) = \operatorname{Im}(\langle v, w \rangle),$$

so that the imaginary part of $\langle \cdot, \cdot \rangle$ is determined by q as well. It follows that two Hermitian forms $\langle \cdot, \cdot \rangle$ and $\langle \cdot, \cdot \rangle'$ on V satisfy $\langle \cdot, \cdot \rangle = \langle \cdot, \cdot \rangle'$ if and only if $\langle v, v \rangle = \langle v, v \rangle'$ for all $v \in V$.

Proof Suppose f is normal, then we have for all $v \in V$

$$\begin{aligned} \|f(v)\|^2 &= \langle f(v), f(v) \rangle = \langle v, f^*(f(v)) \rangle = \langle v, (f^* \circ f)(v) \rangle = \langle v, (f \circ f^*)(v) \rangle \\ &= \langle f^*(v), f^*(v) \rangle = \|f^*(v)\|^2. \end{aligned}$$

Taking the square root implies that $\|f(v)\| = \|f^*(v)\|$ for all $v \in V$.

Conversely, suppose $\|f(v)\| = \|f^*(v)\|$ for all $v \in V$, then the previous calculation implies that

$$\langle v, (f^* \circ f)(v) \rangle = \langle v, (f \circ f^*)(v) \rangle$$

for all $v \in V$. We define Hermitian forms φ_1 and φ_2 on V by the rule

$$\varphi_1(v, w) = \langle w, (f^* \circ f)(v) \rangle \quad \text{and} \quad \varphi_2(v, w) = \langle w, (f \circ f^*)(v) \rangle$$

for all $v, w \in V$. We have $\varphi_1(v, v) = \varphi_2(v, v)$ for all $v \in V$. By the previous remark this implies that $\varphi_1 = \varphi_2$. Hence for all $v, w \in V$, we have

$$\langle w, (f^* \circ f - f \circ f^*)(v) \rangle = 0.$$

Taking $w = (f^* \circ f - f \circ f^*)(v)$, we conclude that $\|(f^* \circ f - f \circ f^*)(v)\| = 0$ for all $v \in V$. It follows that f is normal. $\square$

Similar to the real case, every unitary endomorphism is normal. In addition, we mention the following properties of normal endomorphisms:

Proposition 11.32 (Properties of normal endomorphisms) *Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional unitary space and $f : V \rightarrow V$ a normal endomorphism. Then*

- (i) $\text{Ker } f = \text{Ker } f^*$;
- (ii) λ is an eigenvalue of f if and only if $\bar{\lambda}$ is an eigenvalue of f^* ;
- (iii) the eigenspaces of f are orthogonal. That is, for eigenvalues $\lambda \neq \mu$ of f we have $\langle u, v \rangle = 0$ for all $u \in \text{Eig}_f(\lambda)$ and for all $v \in \text{Eig}_f(\mu)$;
- (iv) if f is self-adjoint, then the eigenvalues of f are real;
- (v) if f is unitary, then the eigenvalues of f are complex numbers of modulus 1.

Proof (i) By definition, $v \in \text{Ker } f$ if and only if $f(v) = 0_V$. This condition is equivalent to $\|f(v)\| = 0$, by the property (i) of norms, see [Proposition 10.12](#). Since $\|f(v)\| = \|f^*(v)\|$ for all $v \in V$ by the previous lemma, we conclude that $\text{Ker } f = \text{Ker } f^*$.

(ii) Observe that if f is normal then $f - \text{sld}_V$ is normal as well for all $s \in \mathbb{C}$. Indeed, using the normality of f , we compute

$$\begin{aligned} (\text{sld}_V - f) \circ (\text{sld}_V - f)^* &= (\text{sld}_V - f) \circ (\bar{s}\text{Id}_V - f^*) = f \circ f^* - \bar{s}f - sf^* + |s|^2\text{Id}_V \\ &= f^* \circ f - sf^* - \bar{s}f + |s|^2\text{Id}_V = (\bar{s}\text{Id}_V - f^*) \circ (\text{sld}_V - f) \\ &= (\text{sld}_V - f)^* \circ (\text{sld}_V - f). \end{aligned}$$

Using (i), we conclude that for all $s \in \mathbb{C}$ we have

$$\text{Eig}_f(s) = \text{Ker}(\text{sld}_V - f) = \text{Ker}(\text{sld}_V - f)^* = \text{Ker}(\bar{s}\text{Id}_V - f^*) = \text{Eig}_{f^*}(\bar{s}).$$

(iii) Let λ be an eigenvalue of f with eigenvector u and μ be an eigenvalue of f with eigenvector v . Then, using (ii) and the conjugate linearity of $\langle \cdot, \cdot \rangle$ in the first argument, we obtain

$$\lambda \langle u, v \rangle = \langle \bar{\lambda}u, v \rangle = \langle f^*(u), v \rangle = \langle u, f(v) \rangle = \langle u, \mu v \rangle = \mu \langle u, v \rangle$$

If $\lambda \neq \mu$, it follows that $\langle u, v \rangle = 0$, as claimed.

(iv) if there exists a non-zero vector $v \in V$ and a scalar $\lambda \in \mathbb{C}$ such that $f(v) = \lambda v$, then we obtain

$$\langle v, f(v) \rangle = \langle v, \lambda v \rangle = \lambda \langle v, v \rangle = \langle f(v), v \rangle = \langle \lambda v, v \rangle = \bar{\lambda} \langle v, v \rangle.$$

Since $\langle v, v \rangle \neq 0$, this implies that $\lambda = \bar{\lambda}$ and hence λ is real.

(v) Suppose λ is an eigenvalue with non-zero eigenvector v of the unitary endomorphism f , then

$$|\lambda|^2 \langle v, v \rangle = \bar{\lambda} \lambda \langle v, v \rangle = \langle \lambda v, \lambda v \rangle = \langle f(v), f(v) \rangle = \langle v, f^*(f(v)) \rangle = \langle v, v \rangle,$$

where we use the conjugate linearity of $\langle \cdot, \cdot \rangle$ in the first argument and that $f^* = f^{-1}$ for a unitary endomorphism. Since $\langle v, v \rangle \neq 0$, it follows that $|\lambda|^2 = 1$. $\square$

It turns out that an endomorphism of a unitary space is diagonalisable with a orthonormal basis if and only if it is normal. This is a statement which is not true in the real setting. For instance, a rotation around the origin in $\mathbb{R}^2$ is a normal endomorphism with respect to the standard scalar product of $\mathbb{R}^2$, but rotations have in general no eigenvectors.

Theorem 11.33 (Spectral theorem for normal endomorphisms) *Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional unitary space and $f : V \rightarrow V$ an endomorphism. Then there exists a basis of V consisting of orthonormal eigenvectors of f if and only if f is normal.*

We need the following lemma in order to prove [Theorem 11.33](#).

Lemma 11.34 *Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional unitary space equipped with an orthonormal ordered basis $\mathbf{b}$ and $f : V \rightarrow V$ an endomorphism. Then f is normal if and only if $\mathbf{A}\mathbf{A}^* = \mathbf{A}^*\mathbf{A}$, where $\mathbf{A} = \mathbf{M}(f, \mathbf{b}, \mathbf{b})$.*

Proof Let $f : V \rightarrow V$ be an endomorphism, then

$$\mathbf{M}(f \circ f^*, \mathbf{b}, \mathbf{b}) = \mathbf{M}(f, \mathbf{b}, \mathbf{b})\mathbf{M}(f^*, \mathbf{b}, \mathbf{b}) = \mathbf{A}\mathbf{A}^*$$

and likewise

$$\mathbf{M}(f^* \circ f, \mathbf{b}, \mathbf{b}) = \mathbf{M}(f^*, \mathbf{b}, \mathbf{b})\mathbf{M}(f, \mathbf{b}, \mathbf{b}) = \mathbf{A}^*\mathbf{A},$$

where we use [Corollary 3.101](#) and that $\mathbf{M}(f^*, \mathbf{b}, \mathbf{b}) = \mathbf{M}(f, \mathbf{b}, \mathbf{b})^*$ by [\(11.1\)](#). Applying [Proposition 2.20](#), we conclude that $f \circ f^* = f^* \circ f$ if and only if $\mathbf{A}\mathbf{A}^* = \mathbf{A}^*\mathbf{A}$. $\square$

Proof of Theorem 11.33 $\Rightarrow$ Suppose there exists an ordered orthonormal basis $\mathbf{b}$ of $(V, \langle \cdot, \cdot \rangle)$ consisting of eigenvectors of f . Hence $\mathbf{A} = \mathbf{M}(f, \mathbf{b}, \mathbf{b})$ is diagonal, that is, $\mathbf{A} = \sum_{i=1}^n \lambda_i \mathbf{E}_{i,i}$, where $\lambda_1, \dots, \lambda_n$ denote the eigenvalues of f and $\{\mathbf{E}_{i,j}\}_{1 \leq i,j \leq n}$ the standard basis of $M_{n,n}(\mathbb{C})$. We thus have that $\mathbf{A}^* = \sum_{j=1}^n \bar{\lambda}_j \mathbf{E}_{j,j}$ and

$$\mathbf{A}\mathbf{A}^* = \sum_{i=1}^n \lambda_i \mathbf{E}_{i,i} \sum_{j=1}^n \bar{\lambda}_j \mathbf{E}_{j,j} = \sum_{i=1}^n \sum_{j=1}^n \lambda_i \bar{\lambda}_j \mathbf{E}_{i,i} \mathbf{E}_{j,j} = \sum_{i=1}^n |\lambda_i|^2 \mathbf{E}_{i,i},$$

where we use [Lemma 4.4](#). Likewise we compute that $\mathbf{A}^*\mathbf{A} = \sum_{i=1}^n |\lambda_i|^2 \mathbf{E}_{i,i}$ and applying [Lemma 11.34](#) we conclude that f is normal.

$\Leftarrow$ We use induction. For $n = 1$ every endomorphism is diagonal, hence there is nothing to show and the statement is anchored.

Inductive Step: Assume that $n \geq 2$ and that the statement is true for all unitary spaces of dimension at most $n - 1$. Since we work over the complex numbers, we can apply [Theorem 6.49](#) to conclude that $f : V \rightarrow V$ admits an eigenvalue $\lambda \in \mathbb{C}$. Let $W = \text{Eig}_f(\lambda)$. We will argue next that the orthogonal complement $W^\perp$ of W is stable under f . Let $w_1 \in W^\perp$ and $w_2 \in W = \text{Eig}_f(\lambda) = \text{Eig}_{f^*}(\bar{\lambda})$. Then, we have

$$\langle f(w_1), w_2 \rangle = \langle w_1, f^*(w_2) \rangle = \langle w_1, \bar{\lambda} w_2 \rangle = \bar{\lambda} \langle w_1, w_2 \rangle = 0,$$

where the last equality follows since $w_1 \in W^\perp$ and $w_2 \in W$. It follows that $f(w_1) \in W^\perp$, hence $W^\perp$ is stable under f . Let $g = f|_{W^\perp} : W^\perp \rightarrow W^\perp$ denote the restriction of f to $W^\perp$. We want to show that g is normal with respect to the restriction of $\langle \cdot, \cdot \rangle$ to $W^\perp$. Using [Lemma 11.30](#), we have for all $w \in W^\perp$

$$\|g(w)\| = \|f(w)\| = \|f^*(w)\| = \|g^*(w)\|$$

and hence g is normal. By the induction hypothesis, there exists an orthonormal basis of $W^\perp$ consisting of eigenvectors of g . As in the real case, we can complement this

basis with an orthonormal basis of $W = \text{Eig}_f(\lambda)$ to obtain an orthonormal basis of $V = W \oplus W^\perp$ consisting of eigenvectors of f . $\square$

Exercises

Exercise 11.35 Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{C})$. Show that $\mathbf{A}$ is unitary if and only if its column vectors form an orthonormal basis of $\mathbb{C}^n$ with respect to the standard Hermitian scalar product $\langle \cdot, \cdot \rangle$.

Exercise 11.36 Verify that $\text{SU}(V, \langle \cdot, \cdot \rangle)$ is a subgroup of $\text{U}(V, \langle \cdot, \cdot \rangle)$ in the sense of Definition 8.8. In particular, $\text{SU}(V, \langle \cdot, \cdot \rangle)$ is indeed a group and hence so is $\text{SU}(n)$.

Exercise 11.37 Show that

$$\text{SU}(2) = \left\{ \begin{pmatrix} z & -\overline{w} \\ w & \overline{z} \end{pmatrix} \mid z, w \in \mathbb{C}, |z|^2 + |w|^2 = 1 \right\}.$$

The Jordan normal form

12.1 Generalised eigenvectors and eigenspaces

WEEK 10

Let $f : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{K}$ -vector space V . Recall from [Proposition 6.46](#) that the eigenspaces of f are in direct sum. Denoting by $\lambda_1, \dots, \lambda_k$ the eigenvalues of f , we have

$$(12.1) \quad \text{Eig}_f(\lambda_1) \oplus \text{Eig}_f(\lambda_2) \oplus \dots \oplus \text{Eig}_f(\lambda_k) = V \iff f \text{ is diagonalisable.}$$

Not every endomorphism is diagonalisable, therefore the left hand side of (12.1) does not hold in general. We would like to remedy this by replacing each eigenspace in (12.1) with a suitable notion of generalised eigenspace. The idea is to consider “eigenvectors of higher rank”. For an endomorphism $f : V \rightarrow V$ and $k \in \mathbb{N}$, we write

$$f^k = \underbrace{f \circ f \circ \dots \circ f}_{k\text{-times}} \quad \text{and define} \quad f^0 = \text{Id}_V.$$

Definition 12.1 (Generalised eigenvector) Let $f : V \rightarrow V$ be an endomorphism of a $\mathbb{K}$ -vector space V . A non-zero vector $v \in V$ is called a *generalised eigenvector of f with eigenvalue $\lambda \in \mathbb{K}$* if

$$(f - \lambda \text{Id}_V)^m(v) = 0_V$$

for some integer $m \in \mathbb{N}$. If a generalised eigenvector v satisfies $(f - \lambda \text{Id}_V)^m(v) = 0_V$ and $(f - \lambda \text{Id}_V)^{m-1}(v) \neq 0_V$, then v is said to have *rank m* .

Remark 12.2 Notice that a generalised eigenvector of $f : V \rightarrow V$ of rank 1 and with eigenvalue λ satisfies

$$(f - \lambda \text{Id}_V)(v) = 0_V \quad \text{and} \quad \text{Id}_V(v) \neq 0_V.$$

Equivalently,

$$f(v) = \lambda v \quad \text{and} \quad v \neq 0_V.$$

Generalised eigenvectors of rank 1 are thus precisely the usual eigenvectors.

The good definition of a generalised eigenspace is a bit trickier.

Definition 12.3 (Generalised eigenspace) Let $f : V \rightarrow V$ be an endomorphism of a $\mathbb{K}$ -vector space V . For all $\lambda \in \mathbb{K}$ we define the *generalised λ -eigenspace of f* to be the set

$$\mathcal{E}_f(\lambda) = \bigcup_{k=0}^{\infty} \text{Ker}((f - \lambda \text{Id}_V)^k)$$

The previous definition, while convenient for proofs, is not particularly handy for computations. Observe however that if $g : V \rightarrow V$ is an endomorphism of a $\mathbb{K}$ -vector space V , then

$$\{0_V\} = \text{Ker}(g^0) \subset \text{Ker}(g^1) \subset \text{Ker}(g^2) \subset \text{Ker}(g^3) \subset \dots$$

and correspondingly we have

$$0 \leq \dim \text{Ker}(g) \leq \dim \text{Ker}(g^2) \leq \dim \text{Ker}(g^3) \leq \dots$$

If V is finite dimensional, then $\dim \text{Ker}((f - \lambda \text{Id}_V)^k)$ can be at most $\dim V$ for all $k \in \mathbb{N}$ and therefore there exists an integer $m \in \mathbb{N}$ so that the generalised λ -eigenspace of f satisfies

$$\mathcal{E}_f(\lambda) = \text{Ker}((f - \lambda \text{Id}_V)^m).$$

Lemma 12.4 *Let $f : V \rightarrow V$ be an endomorphism of a $\mathbb{K}$ -vector space V . Then $\mathcal{E}_f(\lambda) \neq \{0_V\}$ if and only if λ is an eigenvalue of f .*

Proof If λ is an eigenvalue of f then there exists a non-zero vector $v \in \text{Ker}(f - \lambda \text{Id}_V)$ and hence $\dim \mathcal{E}_f(\lambda) > 0$ so that $\mathcal{E}_f(\lambda) \neq \{0_V\}$. Conversely, suppose $\mathcal{E}_f(\lambda) \neq \{0_V\}$ so that there exists an integer m and a non-zero vector $v \in V$ such that $(f - \lambda \text{Id}_V)^m(v) = 0_V$. We may assume m to be the smallest such integer. Then, by assumption, $w = (f - \lambda \text{Id}_V)^{m-1}(v) \neq 0_V$ and w satisfies $f(w) = \lambda w$ and hence is an eigenvector of f with eigenvalue λ . $\square$

By a generalised eigenvector or generalised eigenspace of a matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ we mean those of $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$.

Example 12.5 Consider

$$\mathbf{A} = \begin{pmatrix} 3 & 1 \\ 0 & 3 \end{pmatrix}$$

The characteristic polynomial of $\mathbf{A}$ is $\text{char}_{\mathbf{A}}(\lambda) = (\lambda - 3)^2$, hence we have a single eigenvalue 3 of algebraic multiplicity 2. A simple calculation gives that $\text{Eig}_{\mathbf{A}}(3) = \text{span}\{\vec{e}_1\}$. Now

$$(\mathbf{A} - 3 \cdot \mathbf{1}_2)^2 = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}^2 = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix},$$

hence $\vec{e}_2$ satisfies $(\mathbf{A} - 3 \cdot \mathbf{1}_2)^2 \vec{e}_2 = 0_{\mathbb{K}^2}$ and $(\mathbf{A} - 3 \cdot \mathbf{1}_2) \vec{e}_2 \neq 0_{\mathbb{K}^2}$. Therefore, $\vec{e}_2$ is a generalised eigenvector of $\mathbf{A}$ of rank 2 with eigenvalue 3. We thus have $\mathcal{E}_{\mathbf{A}}(3) = \text{span}\{\vec{e}_1, \vec{e}_2\}$.

Recall that an eigenspace of an endomorphism $f : V \rightarrow V$ is a subspace of V that is stable under f . The same holds true for generalised eigenspaces.

Lemma 12.6 *Let $f : V \rightarrow V$ be an endomorphism of a $\mathbb{K}$ -vector space V and $\lambda \in \mathbb{K}$. Then $\mathcal{E}_f(\lambda)$ is a subspace of V that is stable under f .*

Proof By definition, the zero vector 0_V is an element of $\mathcal{E}_f(\lambda)$, hence $\mathcal{E}_f(\lambda)$ is non-empty. Let $t_1, t_2 \in \mathbb{K}$ and $v_1, v_2 \in \mathcal{E}_f(\lambda)$. Then there exist k_1, k_2 such that $(f - \lambda \text{Id}_V)^{k_1}(v_1) = 0_V$ and $(f - \lambda \text{Id}_V)^{k_2}(v_2) = 0_V$. Take k to be the maximum of $\{k_1, k_2\}$. Then, using the

linearity of $f - \lambda \text{Id}_V$ and its powers, we compute

$$\begin{aligned} 0_V &= t_1(f - \lambda \text{Id}_V)^{k-k_1}(0_V) + t_2(f - \lambda \text{Id}_V)^{k-k_2}(0_V) \\ &= t_1(f - \lambda \text{Id}_V)^{k-k_1}((f - \lambda \text{Id}_V)^{k_1}(v_1)) + t_2(f - \lambda \text{Id}_V)^{k-k_2}((f - \lambda \text{Id}_V)^{k_2}(v_2)) \\ &= t_1(f - \lambda \text{Id}_V)^k(v_1) + t_2(f - \lambda \text{Id}_V)^k(v_2) = (f - \lambda \text{Id}_V)^k(t_1 v_1 + t_2 v_2) \end{aligned}$$

so that $t_1 v_1 + t_2 v_2 \in \text{Ker}((f - \lambda \text{Id}_V)^k) \subset \mathcal{E}_f(\lambda)$ and hence $\mathcal{E}_f(\lambda)$ is a subspace by [Definition 3.21](#).

We now show that $\mathcal{E}_f(\lambda)$ is stable under f . Let $v \in \mathcal{E}_f(\lambda)$ so that there exists $k \geq 0$ with $(f - \lambda \text{Id}_V)^k(v) = 0_V$. Write $w = f(v)$. Then we obtain

$$\begin{aligned} (f - \lambda \text{Id}_V)^k(w) &= (f - \lambda \text{Id}_V)^k(f(v) - \lambda v + \lambda v) \\ &= (f - \lambda \text{Id}_V)^k(f(v) - \lambda v) + \lambda(f - \lambda \text{Id}_V)^k(v) \\ &= (f - \lambda \text{Id}_V)^{k+1}(v) + \lambda(f - \lambda \text{Id}_V)^k(v) = 0_V. \end{aligned}$$

Therefore $w = f(v) \in \mathcal{E}_f(\lambda)$ and hence $\mathcal{E}_f(\lambda)$ is stable under f . $\square$

As for usual eigenspaces, generalised eigenspaces are also in direct sum:

Lemma 12.7 *Let $f : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{K}$ -vector space V . Then the generalised eigenspaces of f are in direct sum.*

Proof Let $\lambda_1, \dots, \lambda_k$ be distinct eigenvalues of f and let n_i for $1 \leq i \leq k$ be such that $\mathcal{E}_f(\lambda_i) = \text{Ker}((f - \lambda_i \text{Id}_V)^{n_i})$. For $1 \leq i \leq k$ let $v_i, \hat{v}_i \in \mathcal{E}_f(\lambda_i)$ be such that

$$(12.2) \quad v_1 + v_2 + \dots + v_k = \hat{v}_1 + \hat{v}_2 + \dots + \hat{v}_k$$

We want to show that $w_i = v_i - \hat{v}_i = 0_V$ for all $1 \leq i \leq k$. For $1 \leq i \leq k$ consider the endomorphism

$$g_i = (f - \lambda_1 \text{Id}_V)^{n_1} \circ \dots \circ (f - \lambda_{i-1} \text{Id}_V)^{n_{i-1}} \circ (f - \lambda_{i+1} \text{Id}_V)^{n_{i+1}} \circ \dots \circ (f - \lambda_k \text{Id}_V)^{n_k}.$$

Notice that g_i does not contain the mapping $(f - \lambda_i \text{Id}_V)^{n_i}$. For $i \neq j$ the mapping g_i contains $(f - \lambda_j \text{Id}_V)^{n_j}$. Rearranging the mappings in g_i if necessary, we can assume that $g_i = h \circ (f - \lambda_j \text{Id}_V)^{n_j}$ for some endomorphism h . Rearranging does not change g_i since for all $\mu_1, \mu_2 \in \mathbb{K}$ we have

$$(f - \mu_1 \text{Id}_V) \circ (f - \mu_2 \text{Id}_V) = (f - \mu_2 \text{Id}_V) \circ (f - \mu_1 \text{Id}_V).$$

Since $w_j \in \mathcal{E}_f(\lambda_j) = \text{Ker}((f - \lambda_j \text{Id}_V)^{n_j})$ we thus conclude that $g_i(w_j) = 0_V$.

By [Lemma 12.6](#) the subspace $\mathcal{E}_f(\lambda_i)$ is stable under f and hence it is also stable under $f - \mu \text{Id}_V$ for all $\mu \in \mathbb{K}$. This implies that $\mathcal{E}_f(\lambda_i)$ is also stable under g_i . Write (12.2) as

$$w_1 + w_2 + \dots + w_k = 0_V.$$

Applying the endomorphism g_i to the previous equation and using that $g_i(w_j) = 0_V$ for $i \neq j$, we obtain that $g_i(w_i) = 0_V$. Since for $j \neq i$ none of the λ_j is a generalised eigenvalue of $f|_{\mathcal{E}_f(\lambda_i)}$, the restriction of g_i to $\mathcal{E}_f(\lambda_i)$ is invertible as an endomorphism of $\mathcal{E}_f(\lambda_i)$. Since $g_i(w_i) = 0_V$, this implies that $w_i = 0$. Since i is arbitrary, we have $w_1 = w_2 = \dots = w_k = 0_V$, as desired. $\square$

We now obtain the desired improvement of (12.1) which holds true without the diagonalisability assumption of f .

Proposition 12.8 Let $f : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{C}$ -vector space V of dimension $n \geq 1$ and let $\lambda_1, \dots, \lambda_k$ denote the distinct eigenvalues of f . Then we have

$$\mathcal{E}_f(\lambda_1) \oplus \mathcal{E}_f(\lambda_2) \oplus \dots \oplus \mathcal{E}_f(\lambda_k) = V.$$

Proof Let $U = \mathcal{E}_f(\lambda_1) \oplus \mathcal{E}_f(\lambda_2) \oplus \dots \oplus \mathcal{E}_f(\lambda_k)$ and suppose that $U \neq V$. Then, by Corollary 6.11 there exists a complement U' of U with $\dim U' \geq 1$. Let $\Pi : V \rightarrow U'$ denote the projection onto U' with kernel U and consider the endomorphism $\hat{f} = \Pi \circ f|_{U'} : U' \rightarrow U'$. Since we work over the complex numbers and since $\dim U' \geq 1$, Theorem 6.49 implies that $\hat{f}$ admits an eigenvalue μ . Let $v \in U'$ be a corresponding eigenvector of $\hat{f}$. Since $U = \text{Ker } \Pi$ is a complement of U' , we obtain

$$f(v) = \mu v + u$$

for some vector $u \in U$. We can write $u = \sum_{i=1}^k u_i$ with $u_i \in \mathcal{E}_f(\lambda_i)$. Now define $g = f - \mu \text{Id}_V : V \rightarrow V$ so that

$$g(v) = \sum_{i=1}^k u_i.$$

Suppose $1 \leq i \leq k$ is such that $\lambda_i \neq \mu$. By definition, $\text{Eig}_f(\lambda_i) \subset \mathcal{E}_f(\lambda_i)$, hence the restriction of $g = f - \mu \text{Id}_V$ to $\mathcal{E}_f(\lambda_i)$ is invertible as an endomorphism of $\mathcal{E}_f(\lambda_i)$, so there exists a vector $v_i \in \mathcal{E}_f(\lambda_i)$ such that $g(v_i) = u_i$. If $\lambda_i \neq \mu$ for all $1 \leq i \leq k$, then we obtain

$$g\left(v - \sum_{i=1}^k v_i\right) = 0_V$$

so that $v - \sum_{i=1}^k v_i$ is an element of $\text{Ker } g = \text{Ker}(f - \mu \text{Id}_V) = \{0_V\}$, where the last equality follows since μ is not an eigenvalue of f . We can therefore write $v = \sum_{i=1}^k v_i \in U$, but this contradicts the assumption that $v \in U'$.

We conclude that we can find an integer i with $1 \leq i \leq k$ such that $\lambda_i = \mu$. After possibly renumbering the eigenvalues we can assume that $\lambda_1 = \mu$ and hence that $\lambda_i \neq \mu$ for $2 \leq i \leq k$, since the eigenvalues are distinct. So again for $2 \leq i \leq k$ we have vectors $v_i \in \mathcal{E}_f(\lambda_i)$ such that $g(v_i) = u_i$. We thus have

$$g\left(v - \sum_{i=2}^k v_i\right) = u_1.$$

Since $\mathcal{E}_f(\lambda_1) = \text{Ker}((f - \lambda_1 \text{Id}_V)^{n_1})$ for some integer n_1 and $g = f - \lambda_1 \text{Id}_V$, applying g^{n_1} , we obtain

$$g^{n_1+1}\left(v - \sum_{i=2}^k v_i\right) = g^{n_1}(u_1) = 0_V,$$

where the last equality uses that $u_1 \in \mathcal{E}_f(\lambda_1)$. It follows that $v - \sum_{i=2}^k v_i \in \mathcal{E}_f(\lambda_1)$ and hence that $v \in U$ which is again a contradiction to the assumption that $v \in U'$. $\square$

Each generalised eigenspace $\mathcal{E}_f(\lambda_i)$ is stable under f . Therefore, if we fix an ordered basis $\mathbf{b}_i$ of $\mathcal{E}_f(\lambda_i)$, then we obtain matrices $\mathbf{A}_i = \mathbf{M}(f|_{\mathcal{E}_f(\lambda_i)}, \mathbf{b}_i, \mathbf{b}_i)$ and the matrix representation of $f : V \rightarrow V$ with respect to the ordered basis $\mathbf{b}$ of V obtained by joining the ordered bases $\mathbf{b}_1, \dots, \mathbf{b}_k$ takes the *block diagonal form* (where unprinted entries are understood to be zero)

$$\begin{pmatrix} \mathbf{A}_1 & & & \\ & \mathbf{A}_2 & & \\ & & \ddots & \\ & & & \mathbf{A}_k \end{pmatrix}$$

We write $\text{diag}(\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_k)$ for such a block diagonal matrix.

Example 12.9 Let

$$\mathbf{A}_1 = \begin{pmatrix} 1 & -3 \\ 4 & 8 \end{pmatrix}, \quad \mathbf{A}_2 = (2), \quad \mathbf{A}_3 = \begin{pmatrix} 7 & -5 & 2 \\ 0 & 1 & -1 \\ 9 & 2 & 0 \end{pmatrix},$$

then we have

$$\text{diag}(\mathbf{A}_1, \mathbf{A}_2, \mathbf{A}_3) = \begin{pmatrix} 1 & -3 & 0 & 0 & 0 & 0 \\ 4 & 8 & 0 & 0 & 0 & 0 \\ 0 & 0 & 2 & 0 & 0 & 0 \\ 0 & 0 & 0 & 7 & -5 & 2 \\ 0 & 0 & 0 & 0 & 1 & -1 \\ 0 & 0 & 0 & 9 & 2 & 0 \end{pmatrix}.$$

12.2 Jordan blocks

The [Proposition 12.8](#) thus tells us that for an endomorphism $f : V \rightarrow V$ of a finite dimensional $\mathbb{C}$ -vector space V , we can always find an ordered basis of V so that the matrix representation of f takes block diagonal form $\text{diag}(\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_k)$. This is already a nice statement, but it turns out that we can say more about how the individual blocks $\mathbf{A}_i$ look like. For a precise statement, we need the notion of a Jordan block. For $m \in \mathbb{N}$ and $\lambda \in \mathbb{K}$ let $\mathbf{J}_m(\lambda) \in M_{m,m}(\mathbb{K})$ denote the $m \times m$ -matrix

$$\mathbf{J}_m(\lambda) = \begin{pmatrix} \lambda & 1 & & & \\ & \lambda & 1 & & \\ & & \ddots & \ddots & \\ & & & \ddots & 1 \\ & & & & \lambda \end{pmatrix} = \begin{cases} \sum_{i=1}^m (\lambda \mathbf{E}_{i,i}) + \sum_{i=1}^{m-1} \mathbf{E}_{i,i+1} & m > 1 \\ (\lambda) & m = 1 \end{cases},$$

where $\{\mathbf{E}_{i,j}\}_{1 \leq i,j \leq m}$ denotes the standard basis of $M_{m,m}(\mathbb{K})$. A matrix of the form $\mathbf{J}_m(\lambda)$ is known as a *Jordan block* of size m .

Example 12.10 (Jordan blocks)

$$\mathbf{J}_1(\lambda) = (\lambda), \quad \mathbf{J}_2(3) = \begin{pmatrix} 3 & 1 \\ 0 & 3 \end{pmatrix}, \quad \mathbf{J}_3(0) = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix},$$

We can now state precisely how the individual matrix blocks look like:

Proposition 12.11 *Let $f : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{K}$ -vector space V and $\lambda \in \mathbb{K}$ an eigenvalue of f . Then there exists an integer $\ell \in \mathbb{N}$, integers $m_1, \dots, m_\ell$ and an ordered basis $\mathbf{b}$ of $\mathcal{E}_f(\lambda)$ such that*

$$\mathbf{M}(f|_{\mathcal{E}_f(\lambda)}, \mathbf{b}, \mathbf{b}) = \text{diag}(\mathbf{J}_{m_1}(\lambda), \mathbf{J}_{m_2}(\lambda), \dots, \mathbf{J}_{m_\ell}(\lambda)).$$

By [Proposition 12.8](#), the vector space V is a direct sum of the generalised eigenspaces of f and by the previous proposition we can find an ordered basis of each eigenspace so that the matrix representation of the restriction of f onto each eigenspace is a sum of Jordan blocks. Combining these two statements, we have thus shown:

Theorem 12.12 (Jordan normal form) *Let $f : V \rightarrow V$ be an endomorphism of a finite dimensional $\mathbb{C}$ -vector space V of dimension $n \geq 1$. Then there exists an ordered basis $\mathbf{b}$ of V , an integer $k \geq 1$, integers $n_1, \dots, n_k$ with $n = n_1 + n_2 + \dots + n_k$ and complex numbers $\lambda_1, \dots, \lambda_k$ such that $\mathbf{M}(f, \mathbf{b}, \mathbf{b}) = \text{diag}(\mathbf{J}_{n_1}(\lambda_1), \mathbf{J}_{n_2}(\lambda_2), \dots, \mathbf{J}_{n_k}(\lambda_k))$, that is,*

$$\mathbf{M}(f, \mathbf{b}, \mathbf{b}) = \begin{pmatrix} \mathbf{J}_{n_1}(\lambda_1) & & & \\ & \mathbf{J}_{n_2}(\lambda_2) & & \\ & & \ddots & \\ & & & \mathbf{J}_{n_k}(\lambda_k) \end{pmatrix}.$$

Remark 12.13 The ordered basis $\mathbf{b}$ of V provided by the Jordan normal form theorem is called a *Jordan basis* for f .

Before we prove [Proposition 12.11](#), we first relate Jordan blocks to the notion of generalised eigenvectors. To this end we first show:

Lemma 12.14 Let $m \in \mathbb{N}$ and $\lambda \in \mathbb{K}$. The only eigenvalue of $\mathbf{J}_m(\lambda)$ is λ . Its algebraic multiplicity is m and its geometric multiplicity is 1.

Proof Recall from [Proposition 5.24](#) that the determinant of an upper triangular matrix is the product of its diagonal entries, hence the characteristic polynomial of the Jordan block $\mathbf{J}_m(\lambda)$ is

$$\text{char}_{\mathbf{J}_m(\lambda)}(x) = (x - \lambda)^m,$$

where here we denote the variable of the characteristic polynomial by x . It follows that λ is the only eigenvalue of $\mathbf{J}_m(\lambda)$ and that its algebraic multiplicity is m . An eigenvector $\vec{v} = (v_i)_{1 \leq i \leq m}$ of $\mathbf{J}_m(\lambda)$ with eigenvalue λ satisfies $\mathbf{J}_m(\lambda)\vec{v} = \lambda\vec{v}$, that is,

$$\lambda v_1 + v_2 = \lambda v_1, \quad \lambda v_2 + v_3 = \lambda v_2, \quad \dots \quad \lambda v_{m-1} + v_m = \lambda v_{m-1}, \quad \lambda v_m = \lambda v_m.$$

Hence $v_2 = v_3 = \dots = v_m = 0$ while v_1 is arbitrary. It follows that the geometric multiplicity of λ is 1. $\square$

The relation between generalised eigenvectors and Jordan blocks is explained by the following two lemmas:

Lemma 12.15 Let $m \in \mathbb{N}$ and $\lambda \in \mathbb{K}$. Then $\vec{e}_m$ is a generalised eigenvector of rank m and with eigenvalue λ of the endomorphism $f_{\mathbf{J}_m(\lambda)} : \mathbb{K}^m \rightarrow \mathbb{K}^m$.

Proof We assume $m > 1$ since for $m = 1$ the statement is trivial. By definition, we need to show that

$$(f_{\mathbf{J}_m(\lambda)} - \lambda \text{Id}_{\mathbb{K}^m})^m(\vec{e}_m) = 0_{\mathbb{K}^m} \quad \text{and} \quad (f_{\mathbf{J}_m(\lambda)} - \lambda \text{Id}_{\mathbb{K}^m})^{m-1}(\vec{e}_m) \neq 0_{\mathbb{K}^m}.$$

By definition, we have $\mathbf{J}_m(\lambda) - \lambda \mathbf{1}_m = \mathbf{J}_m(0) = \sum_{i=1}^{m-1} \mathbf{E}_{i,i+1}$. We use induction to show that for $1 \leq k \leq m-1$, we have

$$(12.3) \quad (\mathbf{J}_m(0))^k = \sum_{i=1}^{m-k} \mathbf{E}_{i,i+k}.$$

For $k = 1$ the statement is obviously correct and hence anchored.

Inductive step: Suppose the statement is correct for $k \geq 1$. We want to show that it is correct for $k + 1 \leq m - 1$. Using the induction hypothesis, we compute

$$(\mathbf{J}_m(0))^{k+1} = \mathbf{J}_m(0)(\mathbf{J}_m(0))^k = \sum_{j=1}^{m-1} \mathbf{E}_{j,j+1} \sum_{i=1}^{m-k} \mathbf{E}_{i,i+k} = \sum_{i=2}^{m-k} \mathbf{E}_{i-1,i+k},$$

where the last equality uses [Lemma 4.4](#). Since

$$\sum_{i=2}^{m-k} \mathbf{E}_{i-1,i+k} = \sum_{i=1}^{m-k-1} \mathbf{E}_{i,i+k+1},$$

(12.3) follows. Now we obtain

$$(f_{\mathbf{J}_m(\lambda)} - \lambda \text{Id}_{\mathbb{K}^m})^{m-1}(\vec{e}_m) = (\mathbf{J}_m(0))^{m-1} \vec{e}_m = \mathbf{E}_{1,m} \vec{e}_m = \vec{e}_1 \neq 0_{\mathbb{K}^m},$$

where the last equality uses that

$$(12.4) \quad \mathbf{E}_{i,j} \vec{e}_k = \delta_{jk} \vec{e}_i,$$

for all $1 \leq i, j, k \leq m$, as can be verified by direct computation. Moreover, using Lemma 4.4 again, we have

$$(12.5) \quad (\mathbf{J}_m(0))^m = (\mathbf{J}_m(0))^{m-1} \mathbf{J}_m(0) = \mathbf{E}_{1,m} \sum_{i=1}^{m-1} \mathbf{E}_{i,i+1} = \mathbf{0}_{m,m}$$

and hence $(f - \lambda \text{Id}_{\mathbb{K}^m})^m(v) = 0_{\mathbb{K}^m}$ for all $v \in V$. In particular $\vec{e}_m$ is a generalised eigenvector of rank m and with eigenvalue λ . $\square$

Using the identities (12.3) and (12.4), we compute for $1 \leq k \leq m-1$

$$(\mathbf{J}_m(0))^k \vec{e}_m = \sum_{i=1}^{m-k} \mathbf{E}_{i,i+k} \vec{e}_m = \sum_{i=1}^{m-k} \delta_{i+k,m} \vec{e}_i = \vec{e}_{m-k}$$

so that

$$((\mathbf{J}_m(0))^{m-1} \vec{e}_m, (\mathbf{J}_m(0))^{m-2} \vec{e}_m, \dots, \mathbf{J}_m(0) \vec{e}_m, \vec{e}_m) = (\vec{e}_1, \vec{e}_2, \dots, \vec{e}_{m-1}, \vec{e}_m).$$

Applying $\mathbf{J}_m(\lambda) - \lambda \mathbf{1}_m$ repeatedly to the generalised eigenvector $\vec{e}_m$ thus gives an ordered basis of V . In general we have:

Lemma 12.16 *Let V be a $\mathbb{K}$ -vector space and $f : V \rightarrow V$ an endomorphism. Suppose $v \in V$ is a generalised eigenvector of f of rank $m \in \mathbb{N}$ with eigenvalue $\lambda \in \mathbb{K}$ and define $u_i = (f - \lambda \text{Id}_V)^{m-i}(v)$ for $1 \leq i \leq m$. Then*

- (i) $\mathbf{b} = (u_1, \dots, u_m)$ is an ordered basis of the subspace $Z(g_\lambda, v) = \text{span}\{u_1, \dots, u_m\}$;
- (ii) the subspace $Z(g_\lambda, v)$ is stable under f ;
- (iii) let $\hat{f}$ denote the restriction of f to $Z(g_\lambda, v)$, then we have $\mathbf{M}(\hat{f}, \mathbf{b}, \mathbf{b}) = \mathbf{J}_m(\lambda)$.

Proof (i) We only need to show that the vectors $\{u_1, \dots, u_m\}$ are linearly independent as by definition, $\{u_1, \dots, u_m\}$ is a generating set for $Z(g_\lambda, v)$. Write $g_\lambda = f - \lambda \text{Id}_V$ then

$$(u_1, \dots, u_m) = (g_\lambda^{m-1}(v), g_\lambda^{m-2}(v), \dots, g_\lambda(v), v).$$

Suppose we have scalars $\mu_1, \dots, \mu_m$ such that

$$(12.6) \quad 0_V = \mu_1 u_1 + \dots + \mu_m u_m = \mu_1 g_\lambda^{m-1}(v) + \mu_2 g_\lambda^{m-2}(v) + \dots + \mu_{m-1} g_\lambda(v) + \mu_m v.$$

Since by assumption $g_\lambda^m(v) = 0_V$ we have $g_\lambda^k(v) = 0_V$ for all $k \geq m$. Applying g_λ $(m-1)$ -times to (12.6) thus gives

$$\mu_1 g_\lambda^{2m-2}(v) + \mu_2 g_\lambda^{2m-3}(v) + \dots + \mu_{m-1} g_\lambda^m(v) + \mu_m g_\lambda^{m-1}(v) = \mu_m g_\lambda^{m-1}(v) = 0_V.$$

By assumption $g_\lambda^{m-1}(v) \neq 0_V$, hence we conclude that $\mu_m = 0$. Therefore, (12.6) becomes

$$\mu_1 u_1 + \dots + \mu_m u_m = \mu_1 g_\lambda^{m-1}(v) + \mu_2 g_\lambda^{m-2}(v) + \dots + \mu_{m-1} g_\lambda(v) = 0_V.$$

Applying g_λ $(m-2)$ -times to the previous equation we conclude that $\mu_{m-1} = 0$ as well. Continuing in this fashion it follows that $\mu_1 = \mu_2 = \dots = \mu_m = 0$, hence the vectors $\{u_1, \dots, u_m\}$ are linearly independent, as claimed.

(ii) Since $\{u_1, \dots, u_m\}$ is a basis of $Z(g_\lambda, v)$, it is sufficient to show that for all $1 \leq i \leq m$ the vector $f(u_i)$ is a linear combination of $\{u_1, \dots, u_m\}$. By construction, we have

$$\begin{aligned}(f - \lambda \text{Id}_V)(u_1) &= g_\lambda^m(v) = 0_V, \\(f - \lambda \text{Id}_V)(u_2) &= g_\lambda^{m-1}(v) = u_1, \\(f - \lambda \text{Id}_V)(u_3) &= g_\lambda^{m-2}(v) = u_2, \\&\vdots \\(f - \lambda \text{Id}_V)(u_m) &= g_\lambda(v) = u_{m-1}\end{aligned}$$

Equivalently, we have

$$f(u_1) = \lambda u_1, \quad f(u_2) = u_1 + \lambda u_2, \quad f(u_3) = u_2 + \lambda u_3, \quad \dots \quad f(u_m) = u_{m-1} + \lambda u_m,$$

which shows the claim.

(iii) Previously we showed that $f(u_1) = \lambda u_1$, hence the first column vector of $\mathbf{M}(\hat{f}, \mathbf{b}, \mathbf{b})$ is $\lambda \vec{e}_1$. For $2 \leq i \leq m$, we have $f(u_i) = u_{i-1} + \lambda u_i$ and hence the i -th column vector of $\mathbf{M}(\hat{f}, \mathbf{b}, \mathbf{b})$ is given by $\vec{e}_{i-1} + \lambda \vec{e}_i$. This shows that $\mathbf{M}(\hat{f}, \mathbf{b}, \mathbf{b}) = \mathbf{J}_m(\lambda)$. $\square$

12.3 Nilpotent endomorphisms

We will prove [Proposition 12.11](#) as a consequence of a statement about so-called nilpotent endomorphisms.

Definition 12.17 (Nilpotent endomorphism) An endomorphism $g : V \rightarrow V$ of a $\mathbb{K}$ -vector space V is called *nilpotent* if there exists an integer $m \in \mathbb{N}$ such that $g^m = o$, where $o : V \rightarrow V$ denotes the zero endomorphism defined by the rule $o(v) = 0_V$ for all $v \in V$. A matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ is called nilpotent if $f_{\mathbf{A}} : \mathbb{K}^n \rightarrow \mathbb{K}^n$ is nilpotent.

Lemma 12.18 Let V be a finite dimensional $\mathbb{K}$ -vector space and $\lambda \in \mathbb{K}$ an eigenvalue of the endomorphism $f : V \rightarrow V$. Then the restriction $g = (f - \lambda \text{Id}_V)|_{\mathcal{E}_f(\lambda)}$ of $f - \lambda \text{Id}_V$ to the generalised eigenspace $\mathcal{E}_f(\lambda)$ is a nilpotent endomorphism.

Proof There exists an integer $m \in \mathbb{N}$ such that $\mathcal{E}_f(\lambda) = \text{Ker}((f - \lambda \text{Id}_V)^m)$. Therefore, for all $v \in \mathcal{E}_f(\lambda)$ we have $(f - \lambda \text{Id}_V)^m(v) = 0_V$ which shows that $g^m = o$, as claimed. $\square$

For nilpotent endomorphisms, we can always find a natural ordered basis of V :

Theorem 12.19 Let V be a finite dimensional $\mathbb{K}$ -vector space and $g : V \rightarrow V$ a nilpotent endomorphism. Then there exists an integer $\ell \in \mathbb{N}$, integers $m_1, \dots, m_\ell \in \mathbb{N}$ and vectors $v_1, \dots, v_\ell \in V$ such that

$$\mathbf{b} = (g^{m_1-1}(v_1), g^{m_1-2}(v_1), \dots, g(v_1), v_1, g^{m_2-1}(v_2), g^{m_2-2}(v_2), \dots, g(v_2), v_2, \dots, g^{m_\ell-1}(v_\ell), g^{m_\ell-2}(v_\ell), \dots, g(v_\ell), v_\ell)$$

is an ordered basis of V and such that $g^{m_1}(v_1) = g^{m_2}(v_2) = \dots = g^{m_\ell}(v_\ell) = 0_V$. In particular, we have

$$\mathbf{M}(g, \mathbf{b}, \mathbf{b}) = \text{diag}(\mathbf{J}_{m_1}(0), \mathbf{J}_{m_2}(0), \dots, \mathbf{J}_{m_\ell}(0)).$$

Proof We use induction on the dimension of the vector space V . For $\dim V = 1$ the only nilpotent endomorphism is the zero endomorphism $o : V \rightarrow V$ and we can take $\ell = 1$, $m_1 = 1$ and $\mathbf{b} = (v)$ for any non-zero vector $v \in V$. The statement is thus anchored.

Inductive step: Suppose $\dim V > 1$ and that the statement is true for all vector spaces of dimension at most $\dim(V) - 1$. Since g is nilpotent, we must have $\det g = 0$ and hence g cannot be surjective by [Proposition 6.22](#). Therefore $U = \text{Im}(g)$ is a subspace of V whose dimension is at most $\dim(V) - 1$. Observe that U is stable under g and hence $h = g|_U : U \rightarrow U$ is a nilpotent endomorphism of U . The induction hypothesis implies that there exists an integer k , integers $n_1, \dots, n_k$ and vectors $u_1, \dots, u_k \in U$ such that

$$\mathbf{c} = (h^{n_1-1}(u_1), h^{n_1-2}(u_1), \dots, h(u_1), u_1, h^{n_2-1}(u_2), h^{n_2-2}(u_2), \dots, h(u_2), u_2, \dots, h^{n_k-1}(u_k), h^{n_k-2}(u_k), \dots, h(u_k), u_k)$$

is an ordered basis of U and such that $h^{n_1}(u_1) = h^{n_2}(u_2) = \dots = h^{n_k}(u_k) = 0_U$.

Since $u_1, \dots, u_k \in U = \text{Im}(g)$, there exist vectors $v_1, \dots, v_k$ such that $u_i = g(v_i)$ for all $1 \leq i \leq k$. Set $m_i = n_i + 1$ for $1 \leq i \leq k$ and consider the set

$$S = \{g^{m_1-1}(v_1), g^{m_1-2}(v_1), \dots, g(v_1), v_1, g^{m_2-1}(v_2), g^{m_2-2}(v_2), \dots, g(v_2), v_2, \dots, g^{m_k-1}(v_k), g^{m_k-2}(v_k), \dots, g(v_k), v_k\}.$$

We claim S is linearly independent. Suppose we can find a linear combination w of the elements of S that gives the zero vector. Applying g to this linear combination, we obtain a linear combination of the elements of

$$\{g^{m_1}(v_1), g^{m_1-1}(v_1), \dots, g^2(v_1), g(v_1), g^{m_2}(v_2), g^{m_2-1}(v_2), \dots, g^2(v_2), g(v_2), \dots, g^{m_k}(v_k), g^{m_k-1}(v_k), \dots, g^2(v_k), g(v_k)\}$$

that gives the zero vector. Equivalently, we obtain a linear combination of the elements of

$$\{g^{m_1-1}(u_1), g^{m_1-2}(u_1), \dots, g(u_1), u_1, g^{m_2-1}(u_2), g^{m_2-2}(u_2), \dots, g(u_2), u_2, \dots, g^{m_k-1}(u_k), g^{m_k-2}(u_k), \dots, g(u_k), u_k\}$$

that gives the zero vector. Equivalently, we obtain a linear combination of the elements of

$$\{h^{n_1}(u_1), h^{n_1-1}(u_1), \dots, h(u_1), u_1, h^{n_2}(u_2), h^{n_2-1}(u_2), \dots, h(u_2), u_2, \dots, h^{n_k}(u_k), h^{n_k-1}(u_k), \dots, h(u_k), u_k\}$$

that gives the zero vector. Here we use that $m_i = n_i + 1$ for $1 \leq i \leq k$ and that $h = g$ on $\text{Im}(g)$. The tuple $\mathbf{c}$ is an ordered basis of U , hence all the coefficients in this linear combination must vanish, except the coefficients before each vector $h^{n_i}(u_i)$, since $h^{n_i}(u_i) = 0_V$ for all $1 \leq i \leq k$. The initial linear combination w thus simplifies to become

$$\mu_1 g^{m_1-1}(v_1) + \mu_2 g^{m_2-1}(v_2) + \dots + \mu_k g^{m_k-1}(v_k) = 0_V.$$

for some scalars $\mu_1, \dots, \mu_k$. It remains to argue that these scalars are all zero. The previous equation is equivalent to

$$\mu_1 h^{n_1-1}(u_1) + \mu_2 h^{n_2-2}(u_2) + \dots + \mu_k h^{n_k-1}(u_k) = 0_V.$$

Using the linear independence of the elements of $\mathbf{c}$ again, we conclude that $\mu_1 = \dots = \mu_k = 0$, as desired.

Observe that by construction, the vectors $v_1, \dots, v_k$ satisfy $g^{m_1}(v_1) = g^{m_2}(v_2) = \dots = g^{m_k}(v_k) = 0_V$.

By [Theorem 3.64](#) we can find an integer $\ell \geq k + 1$ and vectors $T = \{\hat{v}_{k+1}, \dots, \hat{v}_\ell\} \subset V$ such that $S \cup T$ is a basis of V . For each $k + 1 \leq i \leq \ell$, the vector $g(\hat{v}_i)$ is an element of $\text{Im}(g)$ and hence a linear combination of the elements of $\mathbf{c}$. By construction, the elements of $\mathbf{c}$ arise by applying g to the elements of S . It follows that for each $k + 1 \leq i \leq \ell$ there exists a vector $z_i \in \text{span}(S)$ such that $g(z_i) = g(\hat{v}_i)$. For $k + 1 \leq i \leq \ell$, define $v_i = \hat{v}_i - z_i$ and consider the tuple

$$\mathbf{b} = (g^{m_1-1}(v_1), g^{m_1-2}(v_1), \dots, g(v_1), v_1, g^{m_2-1}(v_2), g^{m_2-2}(v_2), \dots, g(v_2), v_2, \dots, g^{m_k-1}(v_k), g^{m_k-2}(v_k), \dots, g(v_k), v_k, v_{k+1}, \dots, v_\ell)$$

Observe that by construction we have $g(v_i) = 0_V$ for $k + 1 \leq i \leq \ell$ so that $m_i = 1$ for $k + 1 \leq i \leq \ell$. Furthermore, the tuple $\mathbf{b}$ has the same number of elements as $S \cup T$ it must thus be the desired ordered basis of V , provided the elements of $\mathbf{b}$ span all of V . Since each v_i arises from $\hat{v}_i$ by subtracting an element in the span of S and since $S \cup T$ generates V , the elements of $\mathbf{b}$ must also generate V .

Finally, the first m_1 vectors of $\mathbf{b}$ are $y_i = g^{m_1-i}(v_1)$ for $1 \leq i \leq m_1$ and we have $g(y_1) = 0_V$ and $g(y_i) = y_{i-1}$ for $2 \leq i \leq m_1$. This contributes the Jordan block $\mathbf{J}_{m_1}(0)$ to the matrix representation of g with respect to $\mathbf{b}$. The remaining blocks arise by considering the vectors $g^{m_k-i}(v_k)$ for $2 \leq k \leq \ell$ and where $1 \leq i \leq m_k$. $\square$

As an application, we obtain:

Proof of Proposition 12.11 Let $f : V \rightarrow V$ be an endomorphism of the finite dimensional $\mathbb{K}$ -vector space V and λ an eigenvalue of f . By [Lemma 12.18](#), the restriction of $g = f - \lambda \text{Id}_V$ to the generalised eigenspace $W = \mathcal{E}_f(\lambda)$ is nilpotent. By [Theorem 12.19](#), there exists an integer $\ell \in \mathbb{N}$, integers $m_1, \dots, m_\ell \in \mathbb{N}$ and vectors $v_1, \dots, v_\ell$ such that

$$\mathbf{b} = (g^{m_1-1}(v_1), g^{m_1-2}(v_1), \dots, g(v_1), v_1, g^{m_2-1}(v_2), g^{m_2-2}(v_2), \dots, g(v_2), v_2, \dots, g^{m_\ell-1}(v_\ell), g^{m_\ell-2}(v_\ell), \dots, g(v_\ell), v_\ell)$$

is an ordered basis of W and such that $g^{m_1}(v_1) = g^{m_2}(v_2) = \dots = g^{m_\ell}(v_\ell) = 0_V$. Notice that this implies that for all $1 \leq i \leq \ell$, the vector v_i is a generalised eigenvector of rank m_i with eigenvalue λ of f and moreover that we have

$$\mathcal{E}_f(\lambda) = \bigoplus_{i=1}^{\ell} Z(g_\lambda, v_i).$$

With respect to this basis we obtain

$$\mathbf{M}(g|_{\mathcal{E}_f(\lambda)}, \mathbf{b}, \mathbf{b}) = \text{diag}(\mathbf{J}_{m_1}(0), \mathbf{J}_{m_2}(0), \dots, \mathbf{J}_{m_\ell}(0))$$

Since $f = g + \lambda \text{Id}_V$, it follows that

$$\mathbf{M}(f|_{\mathcal{E}_f(\lambda)}, \mathbf{b}, \mathbf{b}) = \text{diag}(\mathbf{J}_{m_1}(\lambda), \mathbf{J}_{m_2}(\lambda), \dots, \mathbf{J}_{m_\ell}(\lambda)),$$

as claimed. $\square$

Exercises

Exercise 12.20 Show that $\mathbf{A} \in M_{n,n}(\mathbb{K})$ is nilpotent if and only if there exists an integer $m \in \mathbb{N}$ such that $\mathbf{A}^m = \mathbf{0}_n$.

12.4 Calculations

Let $f : V \rightarrow V$ be an endomorphism of the finite dimensional $\mathbb{K}$ -vector space V . A generalised eigenvector of rank $m \in \mathbb{N}$ with eigenvalue $\lambda \in \mathbb{K}$ of f is an element of $U_m = \text{Ker } g_\lambda^m$, where $g_\lambda = f - \lambda \text{id}_V$. Therefore, we have at most $\dim \text{Ker } g_\lambda^m$ linearly independent generalised eigenvectors of rank m . However the subspace $\text{Ker } g_\lambda^m$ also contains generalised eigenvectors of rank j for $1 \leq j \leq m-1$ and those are elements of $U_{m-1} = \text{Ker } g_\lambda^{m-1} \subset U_m$. The number $\rho_m(\lambda)$ of generalised eigenvectors of rank m with eigenvalue λ of f in a Jordan basis of f is thus given by the dimension of the quotient vector space U_m/U_{m-1} . For $\lambda \in \mathbb{K}$ and $m \in \mathbb{N}$ we define

$$\rho_m(\lambda) = \dim(U_m/U_{m-1}) = \dim \text{Ker}(g_\lambda^m) - \dim \text{Ker}(g_\lambda^{m-1}),$$

where the second equality uses [Proposition 7.10](#). Using the rank-nullity [Theorem 3.76](#), we obtain

$$\rho_m(\lambda) = \dim V - \text{rank } g_\lambda^m - (\dim V - \text{rank } g_\lambda^{m-1}) = \text{rank } g_\lambda^{m-1} - \text{rank } g_\lambda^m.$$

There are only finitely many integers m for which $\rho_m(\lambda) \geq 0$ is non-zero. This follows from the following observation:

Lemma 12.21 *Let $g : V \rightarrow V$ be an endomorphism of the $\mathbb{K}$ -vector space V and suppose there exists $m \in \mathbb{N}$ such that*

$$\text{Ker}(g^{m+1}) = \text{Ker}(g^m).$$

Then we have

$$\text{Ker}(g^m) = \text{Ker}(g^{m+1}) = \text{Ker}(g^{m+2}) = \text{Ker}(g^{m+3}) = \text{Ker}(g^{m+4}) = \dots$$

Proof Let $k \in \mathbb{N}$ be arbitrary. We want to show that $\text{Ker}(g^{m+k}) = \text{Ker}(g^{m+k+1})$. Since $\text{Ker}(g^{m+k}) \subset \text{Ker}(g^{m+k+1})$ we only need to show that $\text{Ker}(g^{m+k+1}) \subset \text{Ker}(g^{m+k})$. Let $v \in \text{Ker}(g^{m+k+1})$. Then

$$g^{m+1}(g^k(v)) = g^{m+k+1}(v) = 0_V$$

and hence $g^k(v) \in \text{Ker}(g^{m+1}) = \text{Ker}(g^m)$. This implies that $g^m(g^k(v)) = g^{m+k}(v) = 0_V$, therefore $v \in \text{Ker}(g^{m+k})$ which shows that $\text{Ker}(g^{m+k+1}) \subset \text{Ker}(g^{m+k})$. $\square$

Example 12.22 Let

$$\mathbf{A} = \begin{pmatrix} 2 & 1 & -1 & 0 & 0 & 0 \\ 0 & 2 & 1 & 0 & 0 & 0 \\ 0 & 0 & 2 & 0 & 0 & 0 \\ 0 & 0 & 0 & 2 & 1 & 0 \\ 0 & 0 & 0 & 0 & 2 & 1 \\ 0 & 0 & 0 & 0 & 0 & 4 \end{pmatrix}.$$

Since $\mathbf{A}$ is an upper triangular matrix we see immediately that its eigenvalues are $\lambda_1 = 2$ and $\lambda_2 = 4$. We compute

$$\mathbf{A} - 2 \cdot \mathbf{1}_6 = \begin{pmatrix} 0 & 1 & -1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 2 \end{pmatrix}$$

and

$$(\mathbf{A} - 2 \cdot \mathbf{1}_6)^2 = \begin{pmatrix} 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 2 \\ 0 & 0 & 0 & 0 & 0 & 4 \end{pmatrix}, \quad (\mathbf{A} - 2 \cdot \mathbf{1}_6)^3 = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 2 \\ 0 & 0 & 0 & 0 & 0 & 4 \\ 0 & 0 & 0 & 0 & 0 & 8 \end{pmatrix}$$

From the expression for $(\mathbf{A} - 2 \cdot \mathbf{1}_6)^3$ we conclude that $\text{rank}((\mathbf{A} - 2 \cdot \mathbf{1}_6)^k) = 1$ for $k \geq 3$. We thus obtain

$$\rho_1(2) = \text{rank } \mathbf{1}_6 - \text{rank}(\mathbf{A} - 2 \cdot \mathbf{1}_6) = 6 - 4 = 2,$$

$$\rho_2(2) = \text{rank}(\mathbf{A} - 2 \cdot \mathbf{1}_6) - \text{rank}((\mathbf{A} - 2 \cdot \mathbf{1}_6)^2) = 4 - 2 = 2,$$

$$\rho_3(2) = \text{rank}((\mathbf{A} - 2 \cdot \mathbf{1}_6)^2) - \text{rank}((\mathbf{A} - 2 \cdot \mathbf{1}_6)^3) = 2 - 1 = 1,$$

$$\rho_k(2) = 0, \quad k \geq 4.$$

A Jordan basis of $f_{\mathbf{A}}$ thus contains $1 = \rho_3(2)$ generalised eigenvector of rank 3 with eigenvalue 2. Since

$$\text{Ker}((\mathbf{A} - 2 \cdot \mathbf{1}_6)^3) = \text{span}\{\vec{e}_1, \vec{e}_2, \vec{e}_3, \vec{e}_4, \vec{e}_5\}$$

and $(\mathbf{A} - 2 \cdot \mathbf{1}_6)^2 \vec{e}_i = 0_{\mathbb{K}^6}$ for $i \neq 3, 6$ we conclude that $\vec{e}_3$ is a generalised eigenvector of rank 3 with eigenvalue 2. The first three vectors of a Jordan basis of $f_{\mathbf{A}}$ are thus given by

$$(\mathbf{A} - 2 \cdot \mathbf{1}_6)^2 \vec{e}_3 = \vec{e}_1, \quad (\mathbf{A} - 2 \cdot \mathbf{1}_6) \vec{e}_3 = -\vec{e}_1 + \vec{e}_2, \quad \vec{e}_3$$

By construction, $-\vec{e}_1 + \vec{e}_2$ is a generalised eigenvector of rank 2 and since $\rho_2(2) = 2$, there must be one more generalised eigenvector of rank 2 in a Jordan basis of $f_{\mathbf{A}}$. We compute

$$\text{Ker}((\mathbf{A} - 2 \cdot \mathbf{1}_6)^2) = \text{span}\{\vec{e}_1, \vec{e}_2, \vec{e}_4, \vec{e}_5\}$$

and that $(\mathbf{A} - 2 \cdot \mathbf{1}_6) \vec{e}_2 \neq 0_{\mathbb{K}^6}$ and $(\mathbf{A} - 2 \cdot \mathbf{1}_6) \vec{e}_5 \neq 0_{\mathbb{K}^6}$. While $\vec{e}_2$ is a generalised eigenvector of rank 2 with eigenvalue 2, it is not linearly independent from our first three Jordan basis vectors $\{\vec{e}_1, -\vec{e}_1 + \vec{e}_2, \vec{e}_3\}$. The vector $\vec{e}_5$ is however linearly independent from the previous Jordan basis vectors and we obtain $(\mathbf{A} - 2 \cdot \mathbf{1}_6) \vec{e}_5 = \vec{e}_4$. The linearly independent vectors $\vec{e}_1, -\vec{e}_1 + \vec{e}_2, \vec{e}_3, \vec{e}_4, \vec{e}_5$ thus span $\mathcal{E}_{\mathbf{A}}(2)$.

The eigenvalue $\lambda_2 = 4$ has algebraic multiplicity 1 and hence also geometric multiplicity 1. We compute

$$\mathcal{E}_{\mathbf{A}}(4) = \text{Eig}_{\mathbf{A}}(4) = \text{span}\{\vec{e}_4 + 2\vec{e}_5 + 4\vec{e}_6\}.$$

Summarising, an ordered Jordan basis of $f_{\mathbf{A}}$ is given by

$$\mathbf{b} = (\vec{e}_1, -\vec{e}_1 + \vec{e}_2, \vec{e}_3, \vec{e}_4, \vec{e}_5, \vec{e}_4 + 2\vec{e}_5 + 4\vec{e}_6).$$

and by construction, we have

$$\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \text{diag}(\mathbf{J}_3(2), \mathbf{J}_2(2), \mathbf{J}_1(4)),$$

as can also be verified by direct computation.

Example 12.23 Let

$$\mathbf{A} = \begin{pmatrix} 1 & 1 & 0 & -1 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

Here we have a single eigenvalue 1 of algebraic multiplicity 4. We obtain

$$\mathbf{A} - 1 \cdot \mathbf{1}_4 = \begin{pmatrix} 0 & 1 & 0 & -1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad \text{and} \quad (\mathbf{A} - 1 \cdot \mathbf{1}_4)^2 = \mathbf{0}_4.$$

Correspondingly, we compute $\rho_2(1) = 2$ and $\rho_1(1) = 2$. A Jordan basis thus contains $2 = \rho_2(1)$ generalised eigenvectors of rank 2 with eigenvalue 1 and those can be chosen to be $\vec{e}_2$ and $\vec{e}_4$. We obtain $(\mathbf{A} - 1 \cdot \mathbf{1}_4)\vec{e}_2 = \vec{e}_1$ and $(\mathbf{A} - 1 \cdot \mathbf{1}_4)\vec{e}_4 = -\vec{e}_1 + \vec{e}_3$. Summarising, an ordered Jordan basis of $f_{\mathbf{A}}$ is given by

$$\mathbf{b} = (\vec{e}_1, \vec{e}_2, -\vec{e}_1 + \vec{e}_3, \vec{e}_4).$$

and by construction, we have

$$\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \text{diag}(\mathbf{J}_2(1), \mathbf{J}_2(1))$$

as can also be verified by direct computation.

Example 12.24 Let

$$\mathbf{A} = \begin{pmatrix} 4 & 0 & 1 & 0 \\ 2 & 2 & 3 & 0 \\ -1 & 0 & 2 & 0 \\ 4 & 0 & 1 & 2 \end{pmatrix}$$

Here the characteristic polynomial is $\text{char}_{\mathbf{A}}(x) = (x - 3)^2(x - 2)^2$ so that we have eigenvalues $\lambda_1 = 3$ and $\lambda_2 = 2$, both with algebraic multiplicity 2. As before, we compute that $\rho_2(3) = 1$ and $\rho_1(3) = 1$ so that a Jordan basis for $f_{\mathbf{A}}$ contains $1 = \rho_2(3)$ generalised eigenvector of rank 2 and $1 = \rho_1(3)$ generalised eigenvector of rank 1, both with eigenvalue 3. The generalised eigenvector of rank 2 can be chosen to be $\vec{e}_1 + 3\vec{e}_2 + \vec{e}_4$ and hence $(\mathbf{A} - 3 \cdot \mathbf{1}_4)(\vec{e}_1 + 3\vec{e}_2 + \vec{e}_4) = \vec{e}_1 - \vec{e}_2 - \vec{e}_3 + 3\vec{e}_4$ is the corresponding generalised eigenvector of rank 1.

Likewise, we obtain $\rho_1(2) = 2$ so that a Jordan basis contains two eigenvectors (of rank 1) with eigenvalue 2. These can be chosen to be $\vec{e}_2$ and $\vec{e}_4$.

Summarising, an ordered Jordan basis of $f_{\mathbf{A}}$ is given by

$$\mathbf{b} = (\vec{e}_1 - \vec{e}_2 - \vec{e}_3 + 3\vec{e}_4, \vec{e}_1 + 3\vec{e}_2 + \vec{e}_4, \vec{e}_2, \vec{e}_4).$$

and by construction, we have

$$\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \text{diag}(\mathbf{J}_2(3), \mathbf{J}_1(2), \mathbf{J}_1(2))$$

as can also be verified by direct computation.

12.5 Applications

12.5.1 The Cayley–Hamilton theorem

Recall that the $\mathbb{K}$ -vector space $M_{n,n}(\mathbb{K})$ of $n \times n$ -matrices has dimension n^2 . Therefore, for a matrix $\mathbf{B} \in M_{n,n}(\mathbb{K})$ the sequence of vectors in $M_{n,n}(\mathbb{K})$ given by the powers of $\mathbf{B}$

$$\mathbf{1}_n, \mathbf{B}, \mathbf{B}^2, \mathbf{B}^3, \mathbf{B}^4, \dots$$

must become linearly dependent. That is, there must exist coefficients $a_i \in \mathbb{K}$ for $0 \leq i \leq n^2$, not all zero such that

$$a_{n^2} \mathbf{B}^{n^2} + a_{n^2-1} \mathbf{B}^{n^2-1} + \cdots + a_2 \mathbf{B}^2 + a_1 \mathbf{B} + a_0 \mathbf{1}_n = \mathbf{0}_n.$$

Remark 12.25 (Notation) Let $a_n, a_{n-1}, \dots, a_1, a_0 \in \mathbb{K}$. For a polynomial $p : \mathbb{K} \rightarrow \mathbb{K}$ defined by the rule $p(x) = a_n x^n + a_{n-1} x^{n-1} + \cdots + a_1 x + a_0$ for all $x \in \mathbb{K}$ and a matrix $\mathbf{B} \in M_{n,n}(\mathbb{K})$, we define

$$p(\mathbf{B}) = a_n \mathbf{B}^n + a_{n-1} \mathbf{B}^{n-1} + \cdots + a_1 \mathbf{B} + a_0 \mathbf{1}_n.$$

We say a matrix $\mathbf{B} \in M_{n,n}(\mathbb{K})$ is a *zero of the polynomial* p if $p(\mathbf{B}) = \mathbf{0}_n$.

Above we have seen that every matrix $\mathbf{B} \in M_{n,n}(\mathbb{K})$ is a zero of a polynomial of degree at most n^2 . One might wonder whether there exists a positive integer d that is strictly smaller than n^2 so that every $n \times n$ -matrix is a zero of a polynomial of degree d .

It turns out that such an integer d must be at least as big as n . For scalars $\lambda_1, \lambda_2, \dots, \lambda_n$ consider the diagonal matrix

$$\mathbf{D} = \begin{pmatrix} \lambda_1 & & & \\ & \lambda_2 & & \\ & & \ddots & \\ & & & \lambda_n \end{pmatrix} \quad \text{with} \quad \mathbf{D}^k = \begin{pmatrix} \lambda_1^k & & & \\ & \lambda_2^k & & \\ & & \ddots & \\ & & & \lambda_n^k \end{pmatrix}$$

for all $k \in \mathbb{N}$. Say we can find a polynomial p of degree $n-1$ such that $p(\mathbf{D}) = \mathbf{0}_n$. Write $p(x) = a_{n-1}x^{n-1} + a_{n-2}x^{n-2} + \cdots + a_1x + a_0$ for coefficients $a_{n-1}, \dots, a_0$. All off-diagonal entries of $p(\mathbf{D})$ are zero and for the i -th diagonal entry of $p(\mathbf{D})$ we obtain $[p(\mathbf{D})]_{ii} = a_{n-1}\lambda_i^{n-1} + a_{n-2}\lambda_i^{n-2} + \cdots + a_1\lambda_i + a_0$. The equation $p(\mathbf{D}) = \mathbf{0}_n$ is equivalent to the linear system of equations $[p(\mathbf{D})]_{11} = [p(\mathbf{D})]_{22} = \cdots = [p(\mathbf{D})]_{nn} = 0$ for the coefficients $a_0, a_1, \dots, a_{n-1}$ and it can be written as

$$\begin{pmatrix} 1 & \lambda_1 & (\lambda_1)^2 & \cdots & (\lambda_1)^{n-1} \\ 1 & \lambda_2 & (\lambda_2)^2 & \cdots & (\lambda_2)^{n-1} \\ 1 & \lambda_3 & (\lambda_3)^2 & \cdots & (\lambda_3)^{n-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & \lambda_n & (\lambda_n)^2 & \cdots & (\lambda_n)^{n-1} \end{pmatrix} \begin{pmatrix} a_0 \\ a_1 \\ a_2 \\ \vdots \\ a_{n-1} \end{pmatrix} = \mathbf{0}_{\mathbb{K}^n}$$

The matrix on the left hand side is the *Vandermonde matrix* $\mathbf{V}_{\vec{\lambda}}$ for the vector $\vec{\lambda} = (\lambda_i)_{1 \leq i \leq n}$. Unless $\det(\mathbf{V}_{\vec{\lambda}}) = 0$, we cannot find a non-zero solution of coefficients $a_{n-1}, \dots, a_0$ such that $p(\mathbf{D}) = \mathbf{0}_n$. By, [Example 5.42](#), we have

$$\det(\mathbf{V}_{\vec{\lambda}}) = \prod_{1 \leq i < j \leq n} (\lambda_j - \lambda_i)$$

and hence if all eigenvalues of $\mathbf{D}$ are distinct, then $\det(\mathbf{V}_{\vec{\lambda}}) \neq 0$. It follows that the smallest positive integer d , so that every $n \times n$ -matrix is a zero of a polynomial of degree d , must be at least n .

For every $n \times n$ -matrix $\mathbf{A}$ we can indeed always find a polynomial p of degree n , so that $p(\mathbf{A}) = \mathbf{0}_n$:

Theorem 12.26 (Cayley–Hamilton theorem) Every matrix $\mathbf{A} \in M_{n,n}(\mathbb{K})$ is a zero of its characteristic polynomial $\text{char}_{\mathbf{A}} : \mathbb{K} \rightarrow \mathbb{K}$

$$\text{char}_{\mathbf{A}}(\mathbf{A}) = \mathbf{0}_n.$$

Example 12.27 Recall from [Remark 6.42](#) that for $\mathbf{A} \in M_{2,2}(\mathbb{K})$ we have $\text{char}_{\mathbf{A}}(\lambda) = \lambda^2 - \text{Tr}(\mathbf{A})\lambda + \det(\mathbf{A})$. Thus, [Theorem 12.26](#) implies that for all $\mathbf{A} \in M_{2,2}(\mathbb{K})$ we have

$$\mathbf{A}^2 - \text{Tr}(\mathbf{A})\mathbf{A} + \det(\mathbf{A})\mathbf{1}_2 = \mathbf{0}_2.$$

For an invertible 2×2 -matrix $\mathbf{A}$ we may write $\mathbf{1}_2 = \mathbf{A}\mathbf{A}^{-1}$ so that

$$\mathbf{A}(\mathbf{A} - \text{Tr}(\mathbf{A})\mathbf{1}_2 + \det(\mathbf{A})\mathbf{A}^{-1}) = \mathbf{0}_2$$

and hence

$$\mathbf{A}^{-1} = \frac{1}{\det \mathbf{A}} (\text{Tr}(\mathbf{A})\mathbf{1}_2 - \mathbf{A})$$

which can of course also be verified by direct computation.

Remark 12.28 It is tempting to argue that

$$\text{char}_{\mathbf{A}}(\mathbf{A}) = \det(\mathbf{A}\mathbf{1}_n - \mathbf{A}) = 0.$$

Notice however that $\text{char}_{\mathbf{A}}(\mathbf{A})$ is an $n \times n$ -matrix, whereas $\det(\mathbf{A}\mathbf{1}_n - \mathbf{A})$ is a scalar, so the previous equation makes no sense if $n > 1$.

That this incorrect calculation gives the correct answer is an accident. To see this observe that for any function $h : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ and to every $\mathbf{A} \in M_{n,n}(\mathbb{K})$ we obtain a function

$$q : \mathbb{K} \rightarrow \mathbb{K}, \quad x \mapsto h(x\mathbf{1}_n - \mathbf{A}).$$

If h is polynomial in the entries of the input matrix, the function q is a polynomial $p_{\mathbf{A}} : \mathbb{K} \rightarrow \mathbb{K}$ depending on $\mathbf{A}$, so that $q(x) = p_{\mathbf{A}}(x)$ for all $x \in \mathbb{K}$. Arguing (wrongly!) as before we would expect that $p_{\mathbf{A}}(\mathbf{A}) = \mathbf{0}_n$. This is however not true in general. Consider for instance

$$h : M_{2,2}(\mathbb{K}) \rightarrow \mathbb{K}, \quad \begin{pmatrix} a & b \\ c & d \end{pmatrix} \mapsto bd$$

so that for

$$\mathbf{A} = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix}$$

we have

$$q(x) = p_{\mathbf{A}}(x) = -A_{12}(x - A_{22})$$

and hence

$$p_{\mathbf{A}}(\mathbf{A}) = -A_{12}\mathbf{A} + A_{12}A_{22}\mathbf{1}_2.$$

For

$$\mathbf{A} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$$

we thus obtain

$$p_{\mathbf{A}}(\mathbf{A}) = -1 \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} + 1 \cdot 0 \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} 0 & -1 \\ 0 & 0 \end{pmatrix} \neq \mathbf{0}_2.$$

Proof of Theorem 12.26 Let $\mathbf{B} \in M_{n,n}(\mathbb{K})$. Recall that $\text{char}_{\mathbf{B}}(x) = \det(x\mathbf{1}_n - \mathbf{B})$ for all $x \in \mathbb{K}$. Using the product rule [Proposition 5.21](#), for an invertible $n \times n$ -matrix $\mathbf{C}$ we thus obtain

$$\begin{aligned} \det(\mathbf{C}(x\mathbf{1}_n - \mathbf{B})\mathbf{C}^{-1}) &= \det(\mathbf{C}) \det((x\mathbf{1}_n - \mathbf{B})\mathbf{C}^{-1}) = \det(\mathbf{C}) \det(x\mathbf{1}_n - \mathbf{B}) \det(\mathbf{C}^{-1}) \\ &= \det(x\mathbf{1}_n - \mathbf{B}) = \det(x\mathbf{1}_n - \mathbf{CBC}^{-1}) \end{aligned}$$

and hence conjugate matrices have the same characteristic polynomial, that is

$$(12.7) \quad \text{char}_{\mathbf{B}}(x) = \text{char}_{\mathbf{CBC}^{-1}}(x)$$

for all $x \in \mathbb{K}$.

We first consider the case $\mathbb{K} = \mathbb{C}$. Let $\mathbf{A} \in M_{n,n}(\mathbb{C})$ be an $n \times n$ -matrix with complex entries. By [Theorem 12.12](#), there exists an ordered basis $\mathbf{b}$ of $\mathbb{C}^n$, an integer $k \geq 1$, integers $n_1, \dots, n_k$ and complex numbers $\lambda_1, \dots, \lambda_k$ such that $\mathbf{M}(f_{\mathbf{A}}, \mathbf{b}, \mathbf{b}) = \mathbf{B}$, where we write

$$\mathbf{B} = \text{diag}(\mathbf{J}_{n_1}(\lambda_1), \mathbf{J}_{n_2}(\lambda_2), \dots, \mathbf{J}_{n_k}(\lambda_k)).$$

Let $\mathbf{e}$ denote the standard ordered basis of $\mathbb{C}^n$ so that $\mathbf{M}(f_{\mathbf{A}}, \mathbf{e}, \mathbf{e}) = \mathbf{A}$ and let $\mathbf{C} = \mathbf{C}(\mathbf{b}, \mathbf{e})$ denote the change of basis matrix. By [Theorem 3.107](#) we have $\mathbf{A} = \mathbf{CBC}^{-1}$.

We want to show that

$$\mathbf{0}_n = \text{char}_{\mathbf{A}}(\mathbf{A}) = \text{char}_{\mathbf{CBC}^{-1}}(\mathbf{CBC}^{-1}) = \text{char}_{\mathbf{B}}(\mathbf{CBC}^{-1}),$$

where the third equality uses [\(12.7\)](#). By induction one shows that $(\mathbf{CBC}^{-1})^k = \mathbf{CB}^k\mathbf{C}^{-1}$ for all $k \in \mathbb{N} \cup \{0\}$. Therefore, we obtain

$$\text{char}_{\mathbf{B}}(\mathbf{CBC}^{-1}) = \mathbf{C} \text{char}_{\mathbf{B}}(\mathbf{B})\mathbf{C}^{-1}$$

and hence – since $\mathbf{C}$ is invertible – we have

$$\mathbf{0}_n = \text{char}_{\mathbf{A}}(\mathbf{A}) \iff \mathbf{0}_n = \text{char}_{\mathbf{B}}(\mathbf{B}).$$

It is thus sufficient to show that $\text{char}_{\mathbf{B}}(\mathbf{B}) = \mathbf{0}_n$. A Jordan block is an upper triangular matrix and hence a block diagonal matrix consisting of Jordan blocks is an upper triangular matrix as well. The [Proposition 5.24](#) thus shows that

$$\text{char}_{\mathbf{B}}(x) = (x - \lambda_1)^{n_1}(x - \lambda_2)^{n_2} \cdots (x - \lambda_k)^{n_k}$$

for all $x \in \mathbb{C}$ and hence

$$\text{char}_{\mathbf{B}}(\mathbf{B}) = (\mathbf{B} - \lambda_1 \mathbf{1}_n)^{n_1} (\mathbf{B} - \lambda_2 \mathbf{1}_n)^{n_2} \cdots (\mathbf{B} - \lambda_k \mathbf{1}_n)^{n_k}.$$

Since $\mathbf{B}\mathbf{1}_n = \mathbf{1}_n\mathbf{B}$, we can rearrange factors in the expression for $\text{char}_{\mathbf{B}}(\mathbf{B})$ so that for each $1 \leq i \leq k$,

$$\begin{aligned} \text{char}_{\mathbf{B}}(\mathbf{B}) &= (\mathbf{B} - \lambda_1 \mathbf{1}_n)^{n_1} \cdots (\mathbf{B} - \lambda_{n_i-1} \mathbf{1}_n)^{n_i-1} (\mathbf{B} - \lambda_{n_i+1} \mathbf{1}_n)^{n_i+1} \cdots \\ &\quad \cdots (\mathbf{B} - \lambda_k \mathbf{1}_n)^{n_k} (\mathbf{B} - \lambda_i \mathbf{1}_n)^{n_i}. \end{aligned}$$

Now observe that

$$\mathbf{B} - \lambda_i \mathbf{1}_n = \text{diag}(\mathbf{J}_{n_1}(\lambda_1 - \lambda_i), \dots, \mathbf{J}_{n_{i-1}}(\lambda_{i-1} - \lambda_i), \mathbf{J}_{n_i}(0), \mathbf{J}_{n_{i+1}}(\lambda_{i+1} - \lambda_i), \dots, \mathbf{J}_{n_k}(\lambda_k - \lambda_i)).$$

By [\(12.5\)](#), we have $(\mathbf{J}_{n_i}(0))^{n_i} = \mathbf{0}_{n_i}$ and hence

$$(\mathbf{B} - \lambda_i \mathbf{1}_n)^{n_i} = \text{diag}(\dots, (\mathbf{J}_{n_i}(0))^{n_i}, \dots) = \text{diag}(\dots, \mathbf{0}_{n_i}, \dots).$$

Therefore, the matrix $(\mathbf{B} - \lambda_i \mathbf{1}_n)^{n_i}$ contains a zero block of size n_i after a diagonal block of size $n_1 + n_2 + \cdots + n_{i-1}$. This shows that $\text{char}_{\mathbf{B}}(\mathbf{B})\vec{e}_j = \mathbf{0}_{\mathbb{C}^n}$ for

$$n_1 + n_2 + \cdots + n_{i-1} < j \leq n_1 + n_2 + \cdots + n_{i-1} + n_i.$$

Since $\text{char}_{\mathbf{B}}(\mathbf{B})\vec{e}_j$ equals the j -th column vector of $\text{char}_{\mathbf{B}}(\mathbf{B})$, it follows that $\text{char}_{\mathbf{B}}(\mathbf{B}) = \mathbf{0}_n$.

Finally, for $\mathbb{K} = \mathbb{R}$ (or in fact any subfield of $\mathbb{C}$) the claim follows by interpreting the entries of $\mathbf{A} \in M_{n,n}(\mathbb{K})$ as complex numbers. □

12.5.2 A matrix is similar to its transpose

Let $\lambda \in \mathbb{K}$ and $n \in \mathbb{N}$. Observe that the matrix representation of $f_{J_n(\lambda)} : \mathbb{K}^n \rightarrow \mathbb{K}^n$ with respect to the ordered basis $\mathbf{b}' = (\vec{e}_n, \vec{e}_{n-1}, \dots, \vec{e}_2, \vec{e}_1)$ of $\mathbb{K}^n$ satisfies $\mathbf{M}(f_{J_n(\lambda)}, \mathbf{b}', \mathbf{b}') = (\mathbf{J}_n(\lambda))^T$. This shows that a Jordan block is similar to its transpose, that is,

$$(\mathbf{J}_n(\lambda))^T = \mathbf{C}(\mathbf{b}, \mathbf{b}')\mathbf{J}_n(\lambda)\mathbf{C}(\mathbf{b}, \mathbf{b}')^{-1}$$

by Theorem 3.107. Using the Jordan normal form, we obtain:

Corollary 12.29 *Let $n \in \mathbb{N}$ and $\mathbf{A} \in M_{n,n}(\mathbb{C})$. Then $\mathbf{A}$ and $\mathbf{A}^T$ are similar, that is, there exists an invertible matrix $\mathbf{X} \in M_{n,n}(\mathbb{C})$ such that $\mathbf{A}^T = \mathbf{X}\mathbf{A}\mathbf{X}^{-1}$.*

Proof By the Jordan normal form theorem there exists an integer $\ell \in \mathbb{N}$, integers $n_1, \dots, n_\ell$ and complex numbers $\lambda_1, \dots, \lambda_\ell$ such that $\mathbf{A}$ is similar to the matrix

$$\mathbf{B} = \text{diag}(\mathbf{J}_{n_1}(\lambda_1), \mathbf{J}_{n_2}(\lambda_2), \dots, \mathbf{J}_{n_\ell}(\lambda_\ell)).$$

That is, there exists an invertible matrix $\mathbf{C} \in M_{n,n}(\mathbb{C})$ such that $\mathbf{A} = \mathbf{C}\mathbf{B}\mathbf{C}^{-1}$. Each Jordan block is similar to its transpose, for $1 \leq i \leq \ell$ we can thus find invertible matrices $\mathbf{Y}_i \in M_{n_i, n_i}(\mathbb{C})$ such that

$$(\mathbf{J}_{n_i}(\lambda_i))^T = \mathbf{Y}_i \mathbf{J}_{n_i}(\lambda_i) \mathbf{Y}_i^{-1}.$$

The invertible block diagonal matrix $\mathbf{Y} = \text{diag}(\mathbf{Y}_1, \dots, \mathbf{Y}_\ell)$ thus satisfies

$$\mathbf{Y}\mathbf{B}\mathbf{Y}^{-1} = \text{diag}((\mathbf{J}_{n_1}(\lambda_1))^T, (\mathbf{J}_{n_2}(\lambda_2))^T, \dots, (\mathbf{J}_{n_\ell}(\lambda_\ell))^T) = \mathbf{B}^T.$$

Since $\mathbf{A} = \mathbf{C}\mathbf{B}\mathbf{C}^{-1}$, we obtain

$$\begin{aligned} \mathbf{A}^T &= (\mathbf{C}^{-1})^T \mathbf{B}^T \mathbf{C}^T = (\mathbf{C}^{-1})^T \mathbf{Y} \mathbf{B} \mathbf{Y}^{-1} \mathbf{C}^T = (\mathbf{C}^{-1})^T \mathbf{Y} \mathbf{C}^{-1} \mathbf{C} \mathbf{B} \mathbf{C}^{-1} \mathbf{C} \mathbf{Y}^{-1} \mathbf{C}^T \\ &= \mathbf{X} \mathbf{A} \mathbf{X}^{-1}, \end{aligned}$$

where $\mathbf{X} = (\mathbf{C}^{-1})^T \mathbf{Y} \mathbf{C}^{-1}$. □

Duality

WEEK 13

13.1 The dual vector space

An important class of vector spaces arises from considering the set of linear maps between two given vector spaces. This set can be turned into a vector space itself in a natural way.

Definition 13.1 (Homomorphism between vector spaces) Let V, W be $\mathbb{K}$ -vector spaces. A linear map $f : V \rightarrow W$ is also called a *homomorphism* between the vector spaces V and W . The set of linear maps between V and W is denoted by $\text{Hom}(V, W)$.

We define addition for $f, g \in \text{Hom}(V, W)$ by the rule

$$(f +_{\text{Hom}(V, W)} g)(v) = f(v) +_W g(v)$$

for all $v \in V$. Here $+_W$ denotes the addition of vectors in W . We define scalar multiplication for $f \in \text{Hom}(V, W)$ and $s \in \mathbb{K}$ by the rule

$$(s \cdot_{\text{Hom}(V, W)} f)(v) = s \cdot_W f(v)$$

for all $v \in V$. Here $\cdot_W$ denotes the scalar multiplication in W . Furthermore, we define the zero vector $0_{\text{Hom}(V, W)}$ to be the function $o : V \rightarrow W$ defined by the rule $o(v) = 0_W$ for all $v \in V$. With these definitions, $\text{Hom}(V, W)$ is a $\mathbb{K}$ -vector space, as can be checked without difficulty.

Proposition 13.2 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $\mathbf{b}$ an ordered basis of V and $\mathbf{c}$ an ordered basis of W . Then the mapping

$$\Theta : \text{Hom}(V, W) \rightarrow M_{m,n}(\mathbb{K}), \quad f \mapsto \mathbf{M}(f, \mathbf{b}, \mathbf{c})$$

is an isomorphism. In particular $\dim \text{Hom}(V, W) = \dim(V) \dim(W)$.

Proof Suppose $\dim V = n, \dim W = m$ and write $\mathbf{b} = (v_1, \dots, v_n)$ and $\mathbf{c} = (w_1, \dots, w_m)$.

We first show that Θ is linear. Let $s_1, s_2 \in \mathbb{K}$ and $f_1, f_2 \in \text{Hom}(V, W)$. By definition

$$\Theta(s_1 f_1 + s_2 f_2) = \mathbf{M}(s_1 f_1 + s_2 f_2, \mathbf{b}, \mathbf{c}),$$

where we omit writing $\cdot_{\text{Hom}(V, W)}$ and where we write $+$ instead of $+_{\text{Hom}(V, W)}$. Linearity means that

$$\Theta(s_1 f_1 + s_2 f_2) = s_1 \mathbf{M}(f_1, \mathbf{b}, \mathbf{c}) + s_2 \mathbf{M}(f_2, \mathbf{b}, \mathbf{c}).$$

Hence we need to show that

$$\mathbf{M}(s_1 f_1 + s_2 f_2, \mathbf{b}, \mathbf{c}) = s_1 \mathbf{M}(f_1, \mathbf{b}, \mathbf{c}) + s_2 \mathbf{M}(f_2, \mathbf{b}, \mathbf{c}).$$

Write

$$\mathbf{M}(f_1, \mathbf{b}, \mathbf{c}) = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n} \quad \text{and} \quad \mathbf{M}(f_2, \mathbf{b}, \mathbf{c}) = (B_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}.$$

Recall from [Proposition 3.93](#) that this means that for all $1 \leq j \leq n$, we have

$$f_1(v_j) = \sum_{i=1}^m A_{ij} w_i \quad \text{and} \quad f_2(v_j) = \sum_{i=1}^m B_{ij} w_i.$$

Therefore, for all $1 \leq j \leq n$, we obtain

$$(s_1 f_1 + s_2 f_2)(v_j) = s_1 f_1(v_j) + s_2 f_2(v_j) = \sum_{i=1}^m (s_1 A_{ij} + s_2 B_{ij}) w_i$$

so that

$$\mathbf{M}(s_1 f_1 + s_2 f_2, \mathbf{b}, \mathbf{c}) = s_1 \mathbf{M}(f_1, \mathbf{b}, \mathbf{c}) + s_2 \mathbf{M}(f_2, \mathbf{b}, \mathbf{c})$$

as claimed.

We next show that Θ is surjective. Let $\mathbf{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n} \in M_{m,n}(\mathbb{K})$ and define $f : V \rightarrow W$ as follows. For $v \in V$ there exist unique scalars $s_1, \dots, s_n$ such that $v = \sum_{i=1}^n s_i v_i$ (since $\mathbf{b}$ is an ordered basis of V). We define

$$f(v) = \sum_{j=1}^n \sum_{i=1}^m A_{ij} s_j w_i.$$

Then f satisfies $f(v_j) = \sum_{i=1}^m A_{ij} w_i$ for all $1 \leq j \leq n$. Hence $\Theta(f) = \mathbf{M}(f, \mathbf{b}, \mathbf{c}) = \mathbf{A}$ and Θ is surjective.

If mappings $f, g \in \text{Hom}(V, W)$ satisfy $\Theta(f) = \mathbf{M}(f, \mathbf{b}, \mathbf{c}) = \Theta(g) = \mathbf{M}(g, \mathbf{b}, \mathbf{c})$, then they agree in particular on the ordered basis $\mathbf{b}$ and hence agree by [Lemma 3.88](#). It follows that Θ is injective as well and hence bijective and thus an isomorphism. Since Θ is an isomorphism we have $\dim \text{Hom}(V, W) = \dim M_{m,n}(\mathbb{K}) = mn = \dim(V) \dim(W)$. $\square$

A case of particular interest is when $W = \mathbb{K}$.

Definition 13.3 (Dual vector space) Let V be a $\mathbb{K}$ -vector space. The $\mathbb{K}$ -vector space $\text{Hom}(V, \mathbb{K})$ is called the *dual vector space* of V and denoted by V^* .

Remark 13.4 Notice that if V is finite dimensional, then

$$\dim(V^*) = \dim(\text{Hom}(V, \mathbb{K})) = \dim(V) \dim(\mathbb{K}) = \dim(V),$$

since $\dim \mathbb{K} = 1$. Therefore, V and V^* have the same dimension and are thus isomorphic vector spaces by [Proposition 3.80](#).

Remark 13.5 (Notation) For $\nu \in V^*$ and $v \in V$ we will sometimes write $v \lrcorner \nu$ for “plugging v into ν ”, that is

$$v \lrcorner \nu = \nu(v).$$

Example 13.6

- (i) For $V = \mathbb{K}^n$ we consider the map which sends a vector $\vec{x} = (x_i)_{1 \leq i \leq n}$ to its i -th entry, $\vec{x} \mapsto x_i$. This map is linear and hence an element of $(\mathbb{K}^n)^*$.
- (ii) Recall that the trace of a matrix is a linear map $\text{Tr} : M_{n,n}(\mathbb{K}) \rightarrow \mathbb{K}$ and hence we may think of the trace as an element of $(M_{n,n}(\mathbb{K}))^*$.

(iii) For $V = P(\mathbb{K})$ and $x_0 \in \mathbb{K}$, we can consider the *evaluation map*

$$\text{ev}_{x_0} : P(\mathbb{K}) \rightarrow \mathbb{K}, \quad p \mapsto p(x_0).$$

The map ev_{x_0} is linear and hence an element of V^* .

(iv) Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean and let $u \in V$. Then we obtain a map

$$\varphi_u : V \rightarrow \mathbb{K}, \quad v \mapsto \langle u, v \rangle.$$

The bilinearity of $\langle \cdot, \cdot \rangle$ implies that φ_u is linear and hence an element of V^* . We thus obtain a map $\Phi_{\langle \cdot, \cdot \rangle} : V \rightarrow V^*$ defined by the rule

$$u \mapsto \varphi_u = \langle u, \cdot \rangle$$

for all $u \in V$. This map is linear and moreover an isomorphism. The linearity is a consequence of the bilinearity of $\langle \cdot, \cdot \rangle$ and since $\dim V = \dim V^*$, it is sufficient to show that $\text{Ker } \Phi_{\langle \cdot, \cdot \rangle} = \{0_V\}$. So suppose that $\varphi_u = 0_{V^*}$ so that $\varphi_u(v) = \langle u, v \rangle = 0$ for all $v \in V$. Since $\langle \cdot, \cdot \rangle$ is non-degenerate, this implies that $u = 0_V$, hence Φ is injective and an isomorphism.

Recall that if V is a $\mathbb{K}$ -vector space of dimension $n \in \mathbb{N}$, then a linear coordinate system on V is an injective (and hence bijective) linear map $\beta : V \rightarrow \mathbb{K}^n$. For a linear coordinate system β and $1 \leq i \leq n$, we may define

$$\nu_i : V \rightarrow \mathbb{K}, \quad v \mapsto [\beta(v)]_i,$$

where $[\beta(v)]_i$ denotes the i -th entry of the vector $\beta(v) \in \mathbb{K}^n$. Both β and taking the i -th entry of a vector in $\mathbb{K}^n$ are linear maps, hence $\nu_i : V \rightarrow \mathbb{K}$ is linear as well and thus an element of V^* . We will argue next that if $\beta : V \rightarrow \mathbb{K}^n$ is a linear coordinate system, then $(\nu_1, \dots, \nu_n)$ is an ordered basis of V^* . Since $\dim V^* = n$, we only need to show that $\{\nu_1, \dots, \nu_n\}$ is linearly independent. Suppose therefore that there are scalars $s_1, \dots, s_n \in \mathbb{K}$ such that

$$(13.1) \quad s_1 \nu_1 + \dots + s_n \nu_n = 0_{V^*} = o,$$

where $o : V \rightarrow \mathbb{K}$ denotes the zero function, that is, $o(v) = 0$ for all $v \in V$. Let $\mathbf{b} = (v_1, \dots, v_n)$ denote the ordered basis of V corresponding to the linear coordinate system β so that $\beta(v_j) = \vec{e}_j$ for all $1 \leq j \leq n$. This is equivalent to

$$\nu_i(v_j) = [\beta(v_j)]_i = [\vec{e}_j]_i = \delta_{ij}$$

for all $1 \leq i, j \leq n$. The Equation (13.1) needs to hold for all choices of $v \in V$, choosing v_k for $1 \leq k \leq n$ gives

$$s_1 \nu_1(v_k) + \dots + s_n \nu_n(v_k) = s_k = o(v_k) = 0$$

so that $s_1 = \dots = s_n = 0$ and $\{\nu_1, \dots, \nu_n\}$ are linearly independent and hence $(\nu_1, \dots, \nu_n)$ is indeed an ordered basis of V^* . We may write

$$\beta = (\nu_1, \dots, \nu_n)$$

and think of a linear coordinate system β on V as an ordered basis $(\nu_1, \dots, \nu_n)$ of V^* .

Definition 13.7 (Dual basis) Let V be a finite dimensional $\mathbb{K}$ -vector space and $\mathbf{b} = (v_1, \dots, v_n)$ an ordered basis of V . The ordered basis $\beta = (\nu_1, \dots, \nu_n)$ of V^* satisfying $\nu_i(v_j) = \delta_{ij}$ for all $1 \leq i, j \leq n$ is called the *ordered dual basis* of $\mathbf{b}$.

13.2 The transpose map

We now come to an important application of the theory of dual vector spaces which leads to a deeper understanding of the matrix transpose.

Definition 13.8 (The transpose map) Let V, W be $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. The map $f^T : W^* \rightarrow V^*$ defined by the rule

$$f^T(\omega) = \omega \circ f$$

for all $\omega \in W^*$ is called the *transpose* of f . Notice that for all $\omega \in W^*$ and for all $v \in V$ we have

$$v \lrcorner f^T(\omega) = f(v) \lrcorner \omega = \omega(f(v)).$$

The transpose map is linear as well.

Lemma 13.9 The transpose $f^T : W^* \rightarrow V^*$ of a linear map $f : V \rightarrow W$ is linear.

Proof We need to show that for all $s_1, s_2 \in \mathbb{K}$ and $\omega_1, \omega_2 \in W^*$, we have

$$f^T(s_1\omega_1 + s_2\omega_2) = s_1f^T(\omega_1) + s_2f^T(\omega_2).$$

This is a condition that needs to hold for all $v \in V$ and indeed, by definition, we have for all $v \in V$

$$\begin{aligned} v \lrcorner f^T(s_1\omega_1 + s_2\omega_2) &= f(v) \lrcorner (s_1\omega_1 + s_2\omega_2) = s_1\omega_1(f(v)) + s_2\omega_2(f(v)) \\ &= s_1(v \lrcorner f^T(\omega_1)) + s_2(v \lrcorner f^T(\omega_2)), \end{aligned}$$

as claimed. □

The relation between the matrix transpose and the transpose mapping is given by the following proposition which states that the matrix representation of the transpose of a linear map is the transpose of the matrix representation of the linear map.

Proposition 13.10 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces equipped with ordered bases $\mathbf{b}, \mathbf{c}$ and corresponding ordered dual bases β, γ of V^*, W^* , respectively. If $f : V \rightarrow W$ is a linear map, then

$$\mathbf{M}(f^T, \gamma, \beta) = \mathbf{M}(f, \mathbf{b}, \mathbf{c})^T.$$

Proof Let $\mathbf{b} = (v_1, \dots, v_n)$, $\mathbf{c} = (w_1, \dots, w_m)$ and $\beta = (\nu_1, \dots, \nu_n)$, $\gamma = (\omega_1, \dots, \omega_m)$. Then, by definition, we have for all $1 \leq j \leq m$

$$f^T(\omega_j) = \sum_{i=1}^n [\mathbf{M}(f^T, \gamma, \beta)]_{ij} \nu_i.$$

Hence for all $1 \leq k \leq n$, we obtain

$$\begin{aligned} v_k \lrcorner f^T(\omega_j) &= v_k \lrcorner \sum_{i=1}^n [\mathbf{M}(f^T, \gamma, \beta)]_{ij} \nu_i = \sum_{i=1}^n [\mathbf{M}(f^T, \gamma, \beta)]_{ij} (v_k \lrcorner \nu_i) \\ &= \sum_{i=1}^n [\mathbf{M}(f^T, \gamma, \beta)]_{ij} \nu_i(v_k) = [\mathbf{M}(f^T, \gamma, \beta)]_{kj}, \end{aligned}$$

where the last equality uses that $\nu_i(v_k) = \delta_{ik}$. By definition, we also have

$$\begin{aligned} v_k \lrcorner f^T(\omega_j) &= f(v_k) \lrcorner \omega_j = \left(\sum_{i=1}^m [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{ik} w_i \right) \lrcorner \omega_j \\ &= \sum_{i=1}^m [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{ik} \omega_j(w_i) = [\mathbf{M}(f, \mathbf{b}, \mathbf{c})]_{jk} = [\mathbf{M}(f, \mathbf{b}, \mathbf{c})^T]_{kj}, \end{aligned}$$

where the second last equality uses $\omega_j(w_i) = \delta_{ji}$. □

Corollary 13.11 Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. Then $\det(f^T) = \det(f)$ and $\text{Tr}(f^T) = \text{Tr}(f)$.

Proof The proof is an exercise. □

Remark 13.12 Recall that for matrices $\mathbf{A} \in M_{m,n}(\mathbb{K})$ and $\mathbf{B} \in M_{n,p}(\mathbb{K})$, we have $(\mathbf{AB})^T = \mathbf{B}^T \mathbf{A}^T$. Correspondingly, let V, W, Z be finite dimensional vector spaces and $f : V \rightarrow W$ and $g : W \rightarrow Z$ be linear maps. Then we obtain $(g \circ f)^T = f^T \circ g^T$. Indeed, for all $\zeta \in Z^*$ we have

$$(g \circ f)^T(\zeta) = \zeta \circ g \circ f = f^T(\zeta \circ g) = f^T(g^T(\zeta)) = (f^T \circ g^T)(\zeta).$$

13.3 Properties of the transpose

For a subspace $U \subset V$ we can consider those elements of V^* that map all vectors of U to 0.

Definition 13.13 (Annihilator) Let V be a $\mathbb{K}$ -vector space and $U \subset V$ a subspace. The *annihilator* of U is the subspace

$$U^0 = \{\nu \in V^* \mid \nu(u) = 0 \forall u \in U\}.$$

Remark 13.14 The annihilator is indeed a subspace. The zero mapping $0 : V \rightarrow \mathbb{K}$ is clearly an element of U^0 , hence U^0 is non-empty. If $\nu_1, \nu_2 \in U$, then we have for all $s_1, s_2 \in \mathbb{K}$ and all $u \in U$

$$(s_1 \nu_1 + s_2 \nu_2)(u) = s_1 \nu_1(u) + s_2 \nu_2(u) = 0,$$

hence by [Definition 3.21](#) it follows that U^0 is a subspace of V^* .

Example 13.15

- (i) Consider $V = \mathcal{P}(\mathbb{R})$ and U to be the subspace of polynomials which contain x^2 as a factor

$$U = \{p \in \mathcal{P}(\mathbb{R}) \mid \text{there exists } q \in \mathcal{P}(\mathbb{R}) \text{ such that } p(x) = x^2 q(x) \forall x \in \mathbb{R}\}.$$

Define a linear map $\varphi : \mathcal{P}(\mathbb{R}) \rightarrow \mathbb{R}$ by the rule

$$\varphi(p) = p'(0)$$

for all $p \in P(\mathbb{R})$ and where p' denotes the derivative of p with respect to x . Then $\varphi \in U^0$.

- (ii) Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional Euclidean space and $U \subset V$ a subspace. Recall that $\langle \cdot, \cdot \rangle$ gives us an isomorphism $\Phi_{\langle \cdot, \cdot \rangle} : V \rightarrow V^*, u \mapsto \langle u, \cdot \rangle$. Observe that $\Phi_{\langle \cdot, \cdot \rangle}(U^\perp) \subset U^0$. Indeed, let $v \in U^\perp$, then

$$\varphi_v(u) = \langle v, u \rangle = 0$$

for all $u \in U$. In fact, $\Phi_{\langle \cdot, \cdot \rangle}(U^\perp) = U^0$. To see this consider an element $\nu \in U^0$. Since $\Phi_{\langle \cdot, \cdot \rangle}$ is surjective it can be written as $\nu = \varphi_v$ for some vector $v \in V$. Now for all $u \in U$ we have

$$\nu(u) = 0 = \langle v, u \rangle$$

which shows that $v \in U^\perp$. The restriction of $\Phi_{\langle \cdot, \cdot \rangle}$ to $U^\perp$ is thus an isomorphism from $U^\perp$ to U^0 .

Previously we saw that for a finite dimensional Euclidean space $(V, \langle \cdot, \cdot \rangle)$ and a subspace $U \subset V$ we have that U^0 is isomorphic to $U^\perp$. Since $V = U \oplus U^\perp$, this implies that $\dim V = \dim U + \dim U^0$. We will give a proof of this fact which also holds over the complex numbers (and in fact over an arbitrary field).

Proposition 13.16 *For a finite dimensional $\mathbb{K}$ -vector space V and a subspace $U \subset V$ we have*

$$\dim V = \dim U + \dim U^0.$$

For the proof we need the following lemma which shows that we can always extend $\mathbb{K}$ -valued linear mappings from subspaces to the whole vector space:

Lemma 13.17 *Let V be a finite dimensional $\mathbb{K}$ -vector space and $U \subset V$ a subspace. Then for every $\omega \in U^*$ there exists an $\Omega \in V^*$ such that $\Omega(u) = \omega(u)$ for all $u \in U$.*

Proof Choose a complement U' of U in V so that $V = U \oplus U'$. Recall that such a complement exists by [Corollary 6.11](#). Consequently, every vector $v \in V$ can be written uniquely as $v = u + u'$. We then define $\Omega(v) = \omega(u)$. $\square$

Proof of Proposition 13.16 We use the rank-nullity [Theorem 3.76](#). Recall that the identity mapping of U is the linear mapping from U to U which returns its input $\text{Id}_U(u) = u$ for all $u \in U$. Since $U \subset V$, we can also think of the identity mapping on U as a mapping into V , $\text{Id}_U : U \rightarrow V$. Applying the rank-nullity theorem to the transpose $\text{Id}_U^T : V^* \rightarrow U^*$, we obtain

$$\dim V = \dim V^* = \dim \text{Ker}(\text{Id}_U^T) + \dim \text{Im}(\text{Id}_U^T),$$

where the first equality uses [Remark 13.4](#). By definition we have

$$\text{Ker}(\text{Id}_U^T) = \left\{ \nu \in V^* \mid \text{Id}_U^T(\nu) = \nu \circ \text{Id}_U = 0_{U^*} \right\}.$$

Again by definition 0_{U^*} is the linear map $o : U \rightarrow \mathbb{K}$ which satisfies $o(u) = 0$ for all $u \in U$. Therefore we have

$$\text{Ker}(\text{Id}_U^T) = \{ \nu \in V^* \mid \nu(u) = 0 \ \forall u \in U \} = U^0.$$

We want to show next that $\text{Id}_U^T : V^* \rightarrow U^*$ is surjective. Let $\omega \in U^*$, by the previous lemma we have $\Omega \in V^*$ so that $\Omega(u) = \omega(u)$ for all $u \in U$. Now notice that for all $u \in U$

we have

$$u \lrcorner \text{Id}_U^T(\Omega) = u \lrcorner (\Omega \circ \text{Id}_U) = \Omega(u) = \omega(u) = u \lrcorner \omega$$

and hence Id_U^T is surjective. It follows that $\dim \text{Im}(\text{Id}_U^T) = \dim U^* = \dim U$. Putting all together, we obtain

$$\dim V = \dim U + \dim U^0,$$

as claimed. $\square$

The kernel of the transpose of a linear map is related to the image of the map:

Proposition 13.18 *Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. Then we have*

- (i) $\text{Ker } f^T = (\text{Im } f)^0$;
- (ii) $\dim \text{Ker } f^T = \dim \text{Ker } f + \dim W - \dim V$.

Proof (i) An element $\omega \in W^*$ lies in the kernel of $f^T : W^* \rightarrow V^*$ if and only if

$$v \lrcorner f^T(\omega) = 0 = f(v) \lrcorner \omega$$

for all $v \in V$. Equivalently, $w \lrcorner \omega = 0$ for all elements w in the image of f , that is, $\omega \in (\text{Im } f)^0$.

(ii) We have

$$\dim \text{Ker } f^T = \dim (\text{Im } f)^0 = \dim W - \dim \text{Im } f = \dim \text{Ker } f + \dim W - \dim V.$$

The first equality uses (i), the second equality uses [Proposition 13.16](#) and the last equality uses the rank-nullity [Theorem 3.76](#). $\square$

Surjectivity of a linear map corresponds to injectivity of its transpose:

Proposition 13.19 *Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. Then f is surjective if and only if f^T is injective.*

Proof The linear map $f : V \rightarrow W$ is surjective if and only if $\text{Im}(f) = W$, equivalently $\text{Im}(f)^0 = \{0_{W^*}\} = \text{Ker}(f^T)$, where the second equality uses the previous proposition. By the characterisation of injectivity of a linear map, [Lemma 3.31](#), we have $\{0_{W^*}\} = \text{Ker}(f^T)$ if and only if f^T is injective. $\square$

Similar to [Proposition 13.18](#) we obtain:

Proposition 13.20 *Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. Then we have*

- (i) $\dim \text{Im}(f^T) = \dim \text{Im}(f)$;
- (ii) $\text{Im}(f^T) = (\text{Ker } f)^0$.

Proof (i) We have

$$\dim \text{Im}(f^T) = \dim W^* - \dim \text{Ker}(f^T) = \dim W^* - \dim \text{Im}(f)^0 = \dim \text{Im}(f),$$

where the first equality uses the rank-nullity [Theorem 3.76](#), the second equality uses [Proposition 13.18](#) and the third equality uses [Proposition 13.16](#).

(ii) First suppose $\nu \in \text{Im}(f^T)$. Then there exists $\omega \in W^*$ with $f^T(\omega) = \nu$. We want to argue that $\nu \in (\text{Ker } f)^0$. By definition

$$(\text{Ker } f)^0 = \{\nu \in V^* \mid \nu(v) = 0 \ \forall v \in \text{Ker } f\}.$$

Let $v \in \text{Ker } f$, then

$$\nu(v) = v \lrcorner \nu = v \lrcorner f^T(\omega) = f(v) \lrcorner \omega = 0_W \lrcorner \omega = 0.$$

It follows that $\text{Im}(f^T) \subset (\text{Ker } f)^0$. We complete the proof by showing that $\text{Im}(f^T)$ and $(\text{Ker } f)^0$ have the same dimension. We compute

$$\dim \text{Im}(f^T) = \dim \text{Im}(f) = \dim V - \dim \text{Ker}(f) = \dim \text{Ker}(f)^0,$$

where the first equality uses (i), the second equality uses the rank-nullity [Theorem 3.76](#) and the last equality uses [Proposition 13.16](#). $\square$

Again, similar to [Proposition 13.19](#) we obtain:

Proposition 13.21 *Let V, W be finite dimensional $\mathbb{K}$ -vector spaces and $f : V \rightarrow W$ a linear map. Then f is injective if and only if f^T is surjective.*

Proof Recall that surjectivity of f^T means that $\text{Im}(f^T) = V^*$. By the characterisation of injectivity, [Lemma 3.31](#), f is injective if and only if $\text{Ker } f = \{0_V\}$, equivalently, $(\text{Ker } f)^0 = V^* = \text{Im}(f^T)$, by the previous proposition. $\square$

13.3.1 The rank of a matrix

Recall that for $\mathbf{A} \in M_{m,n}(\mathbb{K})$ we have defined $\text{rank}(\mathbf{A}) = \dim \text{Im}(f_{\mathbf{A}})$ (c.f. [Definition 3.75](#)). By [Lemma 4.13](#), we have

$$\text{Im}(f_{\mathbf{A}}) = \text{span}\{\mathbf{A}\vec{e}_1, \dots, \mathbf{A}\vec{e}_n\},$$

where $\{\vec{e}_1, \dots, \vec{e}_n\}$ denotes the standard basis of $\mathbb{K}^n$. If we think of the matrix $\mathbf{A}$ as consisting of n column vectors $\vec{a}_1 = \mathbf{A}\vec{e}_1, \dots, \vec{a}_n = \mathbf{A}\vec{e}_n$, then we obtain

$$\text{Im}(f_{\mathbf{A}}) = \text{span}\{\vec{a}_1, \dots, \vec{a}_n\}$$

and hence the rank of $\mathbf{A}$ equals the number of linearly independent column vectors of $\mathbf{A}$, the so-called *column rank* of $\mathbf{A}$. Likewise, we may think of $\mathbf{A}$ as consisting of m row vectors $\vec{\alpha}_1, \dots, \vec{\alpha}_m$ and we can define the *row rank* of $\mathbf{A}$ to be the number of linearly independent row vectors of $\mathbf{A}$. The row rank of a matrix and the column rank are always the same (and hence we simply speak of the rank of the matrix):

Proposition 13.22 *The row rank of every matrix $\mathbf{A} \in M_{m,n}(\mathbb{K})$ equals its column rank.*

Proof The column rank of $\mathbf{A}$ equals $\dim \text{Im}(f_{\mathbf{A}})$. Now

$$\dim \text{Im}(f_{\mathbf{A}}) = \dim \text{Im}((f_{\mathbf{A}})^T) = \dim \text{Im}(f_{\mathbf{A}^T}),$$

where we first use [Proposition 13.20](#) and then [Proposition 13.10](#). Since the matrix transpose interchanges the role of rows and columns, $\dim \text{Im}(f_{\mathbf{A}^T})$ is equal to the number of linearly independent row vectors of $\mathbf{A}$. $\square$

Exercises

Exercise 13.23 Show that the dual basis is indeed uniquely defined by the condition $\nu_i(v_j) = \delta_{ij}$ for all $1 \leq i, j \leq n$.

Exercise 13.24 For a finite dimensional $\mathbb{K}$ -vector space V , we may consider the dual of the dual space, that is $(V^*)^*$. So an element of $(V^*)^*$ is a linear map which takes an element of V^* as its input and produces a scalar as its output. Consider the map $\Xi : V \rightarrow (V^*)^*$ defined by the rule

$$\nu \lrcorner \Xi(v) = v \lrcorner \nu = \nu(v)$$

for all $v \in V$ and all $\nu \in V^*$. That is, the map $\Xi(v) \in (V^*)^*$ applied to $\nu \in V^*$ is given by the application of ν to v . Show that Ξ is an isomorphism.

Exercise 13.25 Consider $V = \mathbb{R}^5$ equipped with the ordered standard basis $\mathbf{e} = (\vec{e}_1, \dots, \vec{e}_5)$ and let $U = \text{span}\{\vec{e}_1, \vec{e}_2\}$. Show that

$$U^0 = \text{span}\{\vec{\varepsilon}_3, \vec{\varepsilon}_4, \vec{\varepsilon}_5\},$$

where $\varepsilon = (\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_5)$ denotes the ordered dual basis of $\mathbf{e}$.