

Advanced Mathematics - Induction

Laura Ballotta

MSc Quantitative Finance

2026

Version 4.1

Contents

I	Advanced Calculus	7
1	Calculus in 1 variable	8
1.1	Functions	8
1.2	A catalogue of useful functions	9
1.2.1	Linear functions	9
1.2.2	Polynomial functions	12
1.2.3	Power functions	15
1.2.4	Inverse functions	17
1.2.5	Rational functions	19
1.2.6	Exponential and logarithmic functions	19
1.2.7	Trigonometric functions	21
1.2.8	Composite functions	23
1.2.9	Other functions	24
1.3	Limits	24
1.4	Derivatives	29
1.4.1	Meaning of the derivative	30
1.4.2	Rules of differentiation	33
1.4.3	Higher order derivatives	38
1.5	Applications of derivatives	40
1.5.1	Linear approximation and Taylor expansion	40
1.5.2	Computing limits: L'Hôpital	46
1.5.3	Implicit differentiation	47
1.5.4	Minima & Maxima of a function	48
1.6	Integrals	54
1.6.1	The Definite Integral	55
1.6.2	How to evaluate integrals: the Indefinite Integral	57
1.6.3	Rules of integration for the Indefinite Integral	60
1.6.4	Applications of integrals - some	65
1.6.5	Improper Integrals	66
2	Linear Algebra	80
2.1	Matrices	81
2.1.1	Special matrices	82
2.1.2	Matrix operations	86
2.2	Rank & nonsingular matrix	94
2.3	Determinant of a matrix	95
2.3.1	Determinant of a 1×1 matrix	95
2.3.2	Determinant of 2×2 matrix: the Leibniz formula	96
2.3.3	Determinant of a 3×3 matrix: Sarrus rule	96
2.3.4	Minors and co-factors	98
2.3.5	Determinant of $n \times n$ matrix: co-factor expansion	100
2.3.6	Properties of the determinant	103
2.3.7	Adjoint Matrix and the Inverse Matrix	108

2.3.8	What is a determinant? A geometric interpretation	110
2.4	Linear Systems	112
2.4.1	Existence and number of solutions	113
2.4.2	Calculation of the solution	115
2.5	Quadratic forms	121
3	Calculus in more than one variable	130
3.1	Real functions of several variables	130
3.2	Partial differentiation	132
3.2.1	First order partial derivatives and the gradient	132
3.2.2	Second order partial derivatives and the Hessian	136
3.3	Minima and Maxima	139
4	Optimization and Applications in Economics and Finance	144
4.1	Solving by substitution	145
4.2	Solving by implicit differentiation	146
4.3	Solving by Lagrange multipliers	147
4.3.1	An intuitive take	147
4.3.2	Method of Lagrange multipliers: the full formulation	150
4.3.3	Meaning of the Lagrange multiplier: the shadow price	156
5	A simple overview of Ordinary Differential Equations (ODEs)	158
5.1	First order linear ODEs	158
5.1.1	Solution of homogenous linear ODEs: Separation of variables	159
5.1.2	Solution of inhomogenous linear ODEs: integrating factor	161
5.2	First order quasi-linear ODEs	164
II	Probability and Statistics	167
6	The basic framework: the probability space. Describing a random experiment	168
7	Random variables	172
7.1	Moments, Cumulants and the Generating Functions - a reminder	174
7.2	Changing measures	181
8	Conditional expectation	183

Foreword

The purpose of this set of Lecture Notes is to offer a common reference for reviewing the fundamental knowledge of Mathematics required to be able to start the modules of the MSc Quantitative Finance with good pace.

Useful Notation

Some important symbols we will use during the course are presented in the following table.

Table 1: Miscellaneous notation

Symbol	Meaning	
$\in$	is a element of (set membership)	$x \in A$: x is an element of set A .
$\iff$	if and only if	$A \iff B$ means A is true if B is true and A is false if B is false.
$\exists$ ($\nexists$)	There exists (does not exist)	$\exists n \in \mathbb{N} : n$ is even.
$\mathbb{N}$	set of the natural numbers	$\{1, 2, 3, \dots\}$.
$\mathbb{R}$	set of real numbers	$3.1416 \in \mathbb{R}$.
$\mathbb{C}$	set of complex numbers	$i = \sqrt{-1} \in \mathbb{C}$.
∂	partial derivative	If $f(x, y) = x^2y$, then $\partial f / \partial x = 2xy$.
$\cup$	the union of (set-theoretic union)	$A \cup B$ means the set of those elements which are either in A , or in B , or in both.
$\cap$	the intersection of (set-theoretic intersection)	$A \cap B$ means the set that contains all those elements that A and B have in common.
∞	infinity	∞ is an element of the extended number line that is greater than all real numbers; it often occurs in limits.
Δ	Difference operator	Greek capital letter ‘delta’ - used in mathematics to denote a difference and/or increment.

Intervals are sections of the real line, for example, ‘all numbers between 0 and 1’. According to whether the endpoints of the interval are included or not, we have closed, semi-closed and open intervals. The notation is set out in Table 2.

Table 2: Interval notation

Standard notation	Interval notation	
$a \leq x \leq b$	$x \in [a, b]$	closed interval
$a < x \leq b$	$x \in (a, b]$	semi-closed interval
$a \leq x < b$	$x \in [a, b)$	semi-closed interval
$a < x < b$	$x \in (a, b)$	open interval
$a \leq x$	$x \in [a, \infty)$	semi-closed interval
$x < b$	$x \in (-\infty, b)$	open interval

Useful Mathematical Constants

1. The constant π (pi) has a natural definition in Euclidean geometry as the ratio between the circumference and diameter of a circle (or, equivalently, the circumference of the circle with unit diameter), but may be found in many areas of mathematics: for example, in the definition of the normal and the Cauchy distributions in statistics.

The numeric value of π is approximately 3.1416. Higher precision is allowed by standard packages such as Excel, Matlab, Stata.

2. Euler’s number e , also known as the exponential growth constant, appears in many areas of mathematics, such as statistics in which it is used to define a number of distributions (for example, the normal, and the exponential one) and useful functions linked to these distributions like the moment generating function and the characteristic functions. These

find important applications in finance too. One possible definition of e is the value of the following expression.

$$e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n.$$

A direct application of e in finance traces back to the Swiss mathematician Jacob Bernoulli, who discovered that e arises in compound interest.

The numeric value of e is approximately 2.7183. Higher precision is allowed by standard packages such as Excel, Matlab, Stata.

3. The imaginary unit or unit imaginary number, denoted as i , is a quantity such that $i^2 = -1$. The term “imaginary” is used because there is no real number having a negative square. It is a mathematical concept which extends the set of real numbers $\mathbb{R}$ to the so-called complex plane $\mathbb{C}$, which in turn provides at least one root for every polynomial $P(x)$. The fundamental application in finance of this constant is in the definition of characteristic function and its counterpart known as the Fourier transform. These are extremely important tools in modern quantitative and computational finance.
4. Pythagoras’ constant is the square root of 2, often known as root 2, radical 2. It is written as $\sqrt{2}$ and is the positive algebraic number that, when multiplied by itself, gives the number 2. $\sqrt{2}$ is the length of a diagonal across a square with sides of one unit of length, and therefore it follows from the Pythagorean theorem. It is fairly ubiquitous and the first number known to be irrational.

The numeric value of $\sqrt{2}$ is approximately 1.4142. Higher precision is allowed by standard packages such as Excel, Matlab, Stata.

There are more important constants in mathematics (for example ϕ , known as the most beautiful number because it represents the ‘golden ratio’ in biology, art and... everywhere), however the ones listed above are the ones we will meet more frequently in finance related applications.

Part I

Advanced Calculus

1 Calculus in 1 variable

1.1 Functions

Functions arise whenever one quantity depends on another.

Example 1 Consider investing £100 in the bank at the rate 2% per annum for two years. The bank offers to pay interest either quarterly or every semester. If it pays quarterly, i.e. four times a year, the value V if this investment (with compound interest) is

$$V = 100 \times \left(1 + \frac{0.02}{4}\right)^{4 \times 2} = 104.0707$$

i.e. £104.0707. If instead the banks pays interest every six months, i.e. twice a year, the value is

$$V = 100 \times \left(1 + \frac{0.02}{2}\right)^{2 \times 2} = 104.0406.$$

The previous example describes a rule by which given a number (payment frequency), another number (value of the investment in 2 years) is assigned.

Definition 1 A function f is a rule that assigns to each element x in a set X exactly one element, called $f(x)$, in a set Y . The set X is called the domain of f , the set Y is called the range of f , and x is the independent variable.

Example 2 1. Consider the function $f(x) = \sqrt{x+2}$. Then, x is the independent variable. Further, the square root is defined only for non-negative arguments, thus $x+2 \geq 0$, and returns non-negative numbers. This implies $X = \{x : x \geq -2\} = [-2, +\infty)$, and for $y = f(x)$, $Y = \{y : y \geq 0\} = [0, +\infty)$.

2. Consider the function $f(x) = \frac{1}{x(x-1)}$. Then, x is the independent variable. Further, the fraction is defined only if the denominator is different from zero, thus $x \neq 0$ and $x \neq 1$. This implies $X = \{x : x \neq 0, x \neq 1\} = (-\infty, 0) \cup (0, 1) \cup (1, +\infty)$, and for $y = f(x)$, $Y = \{y : y \in \mathbb{R}\} = (-\infty, +\infty)$.

3. Let us consider Example 1 again. If we consider the number of times payments are made per year, say m , as our independent variable, then the value V in t years of a principal amount P invested at the rate of interest r is

$$V = P \left(1 + \frac{r}{m}\right)^{mt}.$$

Thus, $V = f(m)$, with f described above. The domain X is the set of the integer numbers, i.e. $X = \{1, 2, 3, \dots\}$; the range is the set of all positive numbers, i.e. $Y = \{y > 0\}$. The quantities P, r, t are the parameters of the function. The function V is plotted in Figure 1.

The remaining of this section is organized as follows. We first study functions commonly used on a day-to-day basis, keeping a keen eye on the ones of particular relevance for financial application. Then, we aim to learn how to get to know these functions better as to understand for example if a given choice of a function works for our purposes, if the plot we are looking at really corresponds the given mathematical expression, and also how to compute some related key quantities such as the maximum value, or to understand the behaviour of the function for very large values.

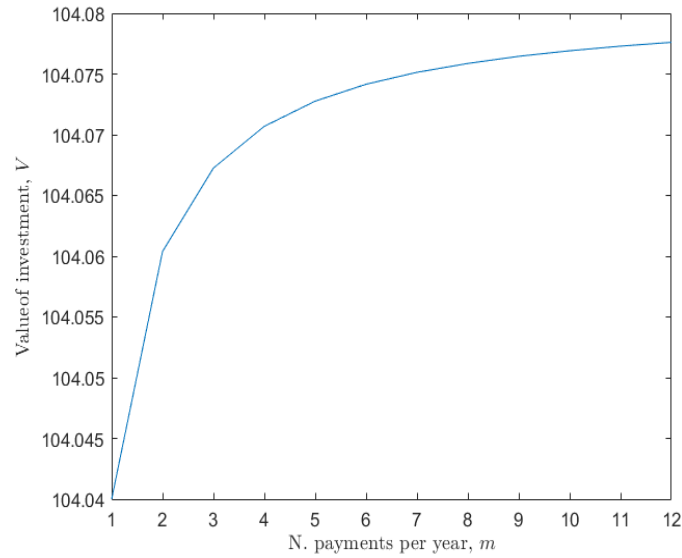


Figure 1: Compound interest from Example 1. Parameter set: $P = 100$; $t = 2$ years; $r = 2\%$ p.a.

1.2 A catalogue of useful functions

1.2.1 Linear functions

Linear functions have the general form

$$f(x) = a + bx$$

with a a real number denoting the so-called intercept, and b a real number denoting the slope of the line, as illustrated in Figure 2.

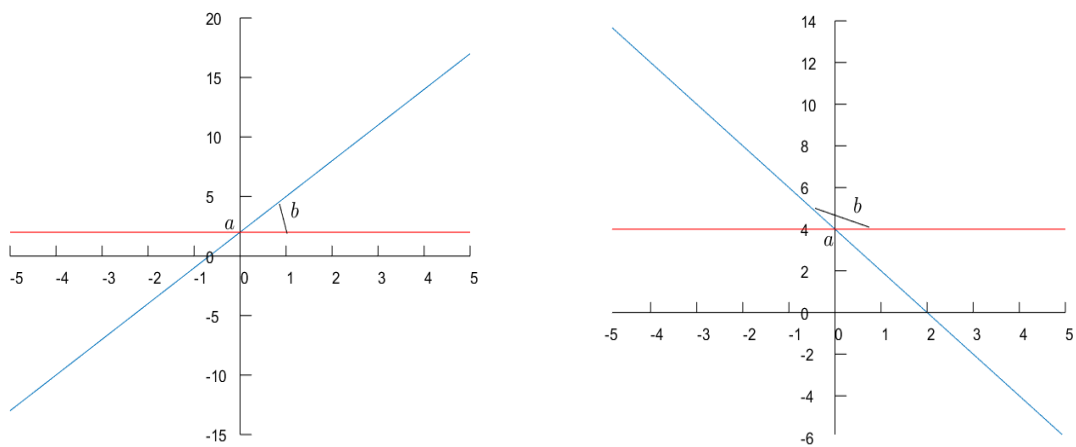


Figure 2: Linear function: $y = f(x)$ with $f(x) = a + bx$. Intercept and slope. Right hand side panel: $a = 2, b = -2$; left hand side panel: $a = 4, b = 3$.

- Example 3**
1. Consider the function $f(x) = x$. The intercept is $a = 0$, which means that the line passes through the origin of the Cartesian axes. The slope is 1. Then, for each value of x that we choose, we obtain a value, say y , which is identical to x . This is why this function is also called the identity function. It is represented in blue colour in Figure 3.
 2. Consider the function $f(x) = 2 + x$. The intercept is $a = 2$, which means that when $x = 0$, $y = f(0) = 2$. The slope is still 1. Then, for each value of x that we choose, we obtain a value, say y , which is equal to x augmented by the value 2. This is a parallel shift for the identity function as also shown in Figure 3 (red colour).
 3. Consider the function $f(x) = 2 + 3x$. The intercept is $a = 2$, which means that, like in the previous case, that when $x = 0$, $y = f(0) = 2$. The slope is now 3. Then, for each value of x that we choose, we obtain a value, say y , which is 3 times x , augmented by 2. This rescaling of the slope leads to a rotation of the line compared to the previous case, as shown in Figure 3 (magenta colour).

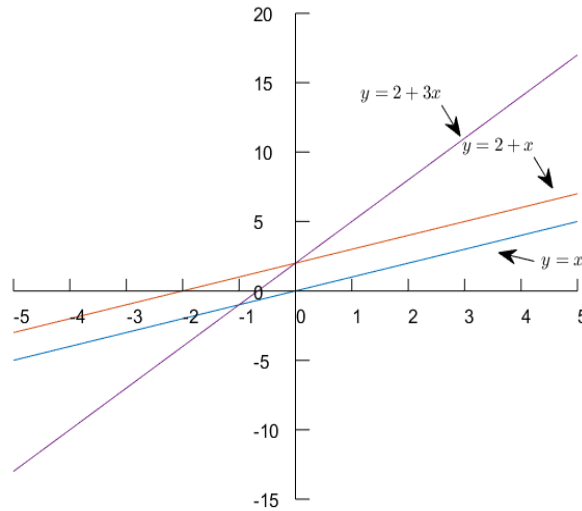


Figure 3: Linear functions: some examples. Blue line: $f(x) = x$ - the identity function. Red line: $f(x) = 2 + x$. Magenta line: $f(x) = 2 + 3x$.

In the previous examples, we are given the linear function, in the sense that both intercept and slope are known. At times, though, all we know is a couple of points that go through the line. How can we determine the intercept and the slope in this situation?

Suppose we know that for $x = 0$, $y = q$ and that for $x = \bar{x}$, $y = \bar{y}$. Consider the general form of the linear function, i.e.

$$y = f(x) = a + bx.$$

Let's use the information available and substitute in the linear function. Thus, the pair $(x = 0, y = q)$ returns

$$q = f(0) = a,$$

therefore $a = q$. Similarly, the pair $(x = \bar{x}, y = \bar{y})$ returns

$$\bar{y} = f(\bar{x}) = q + b\bar{x}.$$

Now let us solve the above in terms of b , i.e.

$$b = \frac{f(\bar{x}) - q}{\bar{x}} = \frac{\bar{y} - q}{\bar{x}}.$$

We can generalize this to any two points on the plane with coordinates $(x = \bar{x}_1, y = \bar{y}_1)$ and $(x = \bar{x}_2, y = \bar{y}_2)$. We obtain

$$\begin{aligned}\bar{y}_1 &= a + b\bar{x}_1 \\ \bar{y}_2 &= a + b\bar{x}_2.\end{aligned}$$

Rearrange the above system of equations to obtain

$$\begin{aligned}a &= \bar{y}_1 - b\bar{x}_1 \\ a &= \bar{y}_2 - b\bar{x}_2,\end{aligned}$$

which implies

$$\bar{y}_1 - b\bar{x}_1 = \bar{y}_2 - b\bar{x}_2$$

and therefore

$$b = \frac{\bar{y}_2 - \bar{y}_1}{\bar{x}_2 - \bar{x}_1}.$$

Substitute back in, for example, $a = \bar{y}_1 - b\bar{x}_1$ to obtain a and then plug in into the linear equation. Then your linear function reads

$$y = \bar{y}_1 + \frac{\bar{y}_2 - \bar{y}_1}{\bar{x}_2 - \bar{x}_1}(x - \bar{x}_1).$$

(The same result is achieved by obtaining a from $a = \bar{y}_2 - b\bar{x}_2$.)

Example 4 Consider the points $(\bar{x}_1, \bar{y}_1) = (1, 5)$ and $(\bar{x}_2, \bar{y}_2) = (3, 11)$. The line passing through these two points then has slope

$$b = \frac{11 - 5}{3 - 1} = \frac{6}{2} = 3.$$

Then the linear function in our example is $y = 5 + 3 \times (x - 1)$, i.e. $y = 2 + 3x$.

The reason why we want to look into this result is to understand the meaning of the slope coefficient b . A graphical illustration can help, so we refer to Figure 4. This shows that the quantity $\bar{y}_2 - \bar{y}_1$ is the distance measured on the y -axis between the vertical coordinates of the two given points, i.e. the segment AB . Let us denote this distance Δy , i.e. $\Delta y = \bar{y}_2 - \bar{y}_1$. Similarly the quantity $\Delta x = \bar{x}_2 - \bar{x}_1$ is the distance measured on the x -axis between the horizontal coordinates of the two given points, i.e. the segment OB . Thus, the formula

$$b = \frac{\bar{y}_2 - \bar{y}_1}{\bar{x}_2 - \bar{x}_1} = \frac{\Delta y}{\Delta x}$$

is telling us that the slope of the linear function is given by the ratio between the distance AB and the distance OB . You can interpret these distances as increments, i.e. AB represents the increment between $\bar{y}_1$ and $\bar{y}_2$, i.e. Δy , and similarly OB represents the increment between $\bar{x}_1$ and $\bar{x}_2$, i.e. Δx . So, the slope coefficient b is a ratio between the increment on the y -axis and the increment on the x -axis. For this reason, from now on we refer to the slope coefficient as the incremental ratio (or differential quotient). This concept will be very important when we define the derivative of a generic function f .

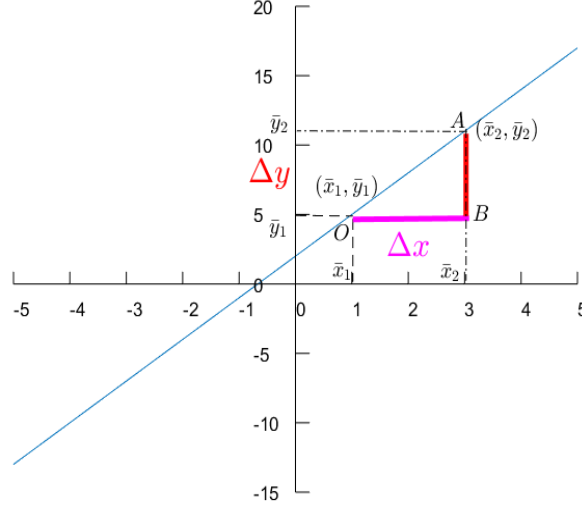


Figure 4: A line across two points $O = (\bar{x}_1, \bar{y}_1)$ and $A = (\bar{x}_2, \bar{y}_2)$: the function $y = 2 + 3x$.

1.2.2 Polynomial functions

Polynomial functions can be considered as a generalization of the linear function and are defined as

$$f(x) = a_n x^n + a_{n-1} x^{n-1} + \cdots + a_2 x^2 + a_1 x + a_0$$

for n integer; the quantities $a_1, \dots, a_n$ are the coefficients of the polynomial, a_0 is the intercept. The coefficient a_n is the so-called leading coefficient as it tells us the degree of the polynomial, i.e. n .

Example 5 1. If $n = 1$, then we have the linear function, say for example that $a_0 = 2$ and $a_1 = 3$ then

$$y = 3x + 2.$$

This is a polynomial of degree 1.

2. If $n = 2$, we have a quadratic function, or parabola, say for example that $a_0 = 2, a_1 = 3$ and $a_2 = 1$, then

$$y = x^2 + 3x + 2.$$

This is a polynomial of degree 2.

3. If $n = 3$, we have a cubic function, say for example that $a_0 = 2, a_1 = 3, a_2 = 1$ and $a_3 = 1/2$, then

$$y = \frac{1}{2}x^3 + x^2 + 3x + 2.$$

This is a polynomial of degree 3.

These functions are shown in Figure 5.

Some special polynomials are

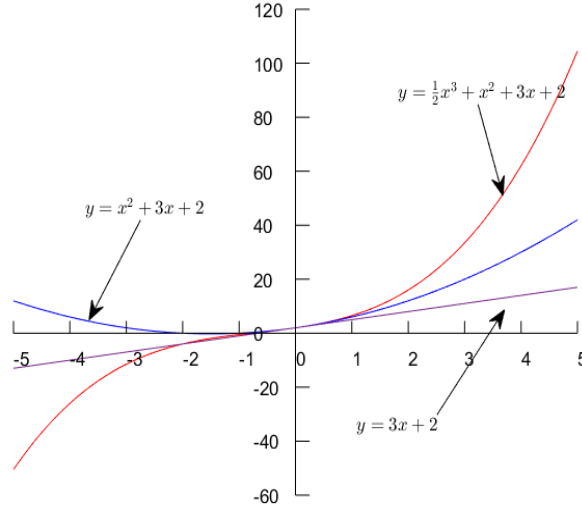


Figure 5: Polynomial functions from Example 5.

1. $y = (a + x)^2 = a^2 + 2ax + x^2$,
2. $y = (a + x)^3 = a^3 + 3a^2x + 3ax^2 + x^3$,
3. and the more general $y = (a + x)^n$, for n integer. The explicit expression is given by the binomial theorem which says that for any n integer

$$(a+x)^n = x^n + \frac{n!}{1!(n-1)!}ax^{n-1} + \frac{n!}{2!(n-2)!}a^2x^{n-2} + \cdots + \frac{n!}{(n-2)!2!}a^{n-2}x^2 + \frac{n!}{(n-1)!1!}a^{n-1}x + a^n.$$

This can be written in a compact form as

$$(a+x)^n = \sum_{k=0}^n \frac{n!}{k!(n-k)!} a^k x^{n-k} \quad (1)$$

with the symbol $n!$ (read ' n factorial') defined as

$$n! = n(n-1) \times (n-2) \times \cdots \times 2 \times 1$$

that is the product of the first n integer numbers, and $0! = 1$ by convention (so, for example $1! = 1$, $2! = 2 \times 1 = 2$, $3! = 3 \times 2 \times 1 = 6$ and so on). All special polynomials listed above can be expressed in this form. For example, let us consider the polynomial $(a+x)^2$, then the general formula states that

$$(a+x)^2 = \sum_{k=0}^2 \frac{2!}{k!(2-k)!} a^k x^{2-k}.$$

To compute each term we follow these steps. First: compute the addendum for each value

of $k = 0, 1, 2$, i.e.

$$\text{Set } k = 0, \text{ then } \frac{2!}{k!(2-k)!} a^k x^{2-k} = \frac{2}{1 \times 2} a^0 x^2 = x^2.$$

$$\text{Set } k = 1, \text{ then } \frac{2!}{k!(2-k)!} a^k x^{2-k} = \frac{2}{1 \times 1} a^1 x^1 = 2ax.$$

$$\text{Set } k = 2, \text{ then } \frac{2!}{k!(2-k)!} a^k x^{2-k} = \frac{2}{2 \times 1} a^2 x^0 = a^2.$$

Then sum all the terms, which returns $a^2 + 2ax + x^2$ as we would expect. What is important is to keep always in mind the so-called binomial coefficient

$$\frac{n!}{k!(n-k)!}$$

which you can also find written as

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

An easy way to remember the terms that compose it is to use the Pascal triangle rule, which is illustrated in Figure 6. The triangle has 1s along the two sides; the numbers inside the triangle are obtained summing up the pairs of numbers in the line above. For example, consider the case $n = 3$. The coefficients are respectively 1, 3, 3, 1; the ‘interior’ coefficients, i.e. 3, are obtained by summing each pair of coefficients of the case $n - 1 = 2$: i.e. $1 + 2 = 3$, $2 + 1 = 3$.

The binomial theorem will help us to calculate some derivatives in Section 1.4.

n	$(a+b)^n$	COEFFICIENTS
0	1	1
1	$a+b$	1 1
2	$a^2+2ab+b^2$	1 2 1
3	$a^3+3a^2b+3ab^2+b^3$	1 3 3 1
4	$a^4+4a^3b+6a^2b^2+4ab^3+b^4$	1 4 6 4 1
⋮	⋮	⋮

Figure 6: The coefficients of $(a+b)^n$: the Pascal triangle.

Remark 1 Let p be such that $0 \leq p \leq 1$, i.e. $p \in [0, 1]$. The binomial formula above allows us to show that

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

Therefore expression

$$\frac{n!}{k!(n-k)!} p^k (1-p)^{n-k}$$

is a legitimate probability mass function, and indeed it is used in Statistics to define the probability mass function of the binomial distribution, which represents the building block for important financial models such as the Binomial Model for Option Pricing by Cox, Ross and Rubinstein (1979).

1.2.3 Power functions

Power functions the have form

$$f(x) = x^a$$

for any constant a (any real number). Let us consider a few interesting cases.

1. $a = n$ for n integer. Then we obtain special cases of polynomial functions with only one term, i.e. the leading term,

$$y = x^n.$$

Note that if $x \geq 0$, y is always positive; for $x \leq 0$, y is either positive or negative according to whether n is even or odd (respectively). This is shown in Figure 7 for the cases $y = x^2$ and $y = x^3$.

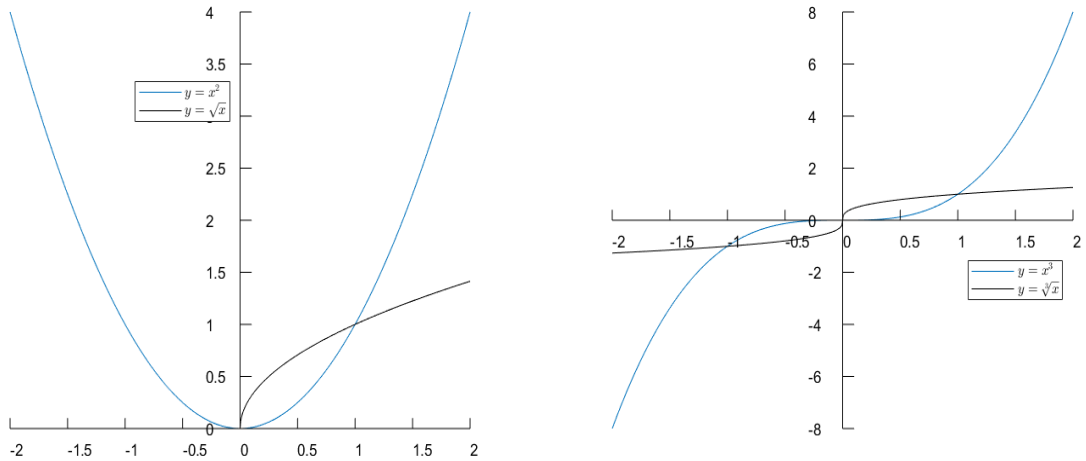


Figure 7: Power functions. Left hand side panel: $y = x^2$ and $y = \sqrt{x}$. Right hand side panel: $y = x^3$ and $y = \sqrt[3]{x}$.

2. $a = \frac{1}{n}$ for n integer. Then we obtain a function of the form

$$y = x^{1/n} = \sqrt[n]{x}$$

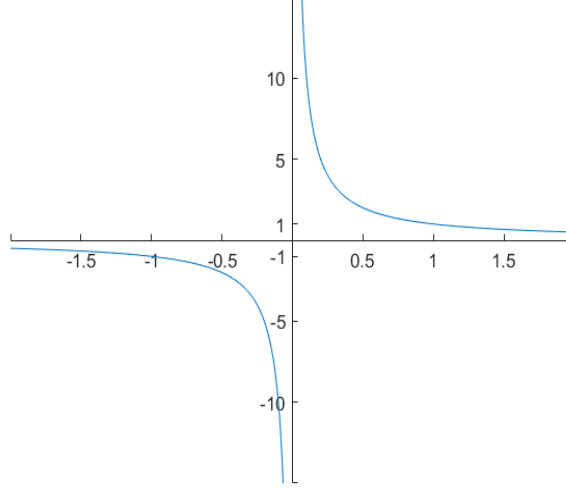


Figure 8: Reciprocal function (hyperbola) $y = \frac{1}{x}$.

known as the root function. A few examples are shown in Figure 7. We note that the square root function is defined only on the positive axis, i.e. the domain of $f(x) = \sqrt{x}$ is $x \in [0, +\infty)$. Note though that $y = \sqrt{x^2} = \pm x = |x|$, $+x$ being the so-called principal value. The cubic root can instead take both negative and positive values.

3. $a = -1$, then for $x \neq 0$

$$y = x^{-1} = \frac{1}{x}.$$

This is the reciprocal function, and it is such that $xy = 1$; it is represented in Figure 8, and it is a hyperbola with the coordinate axes as its asymptotes. Indeed, for $x = 0$ the function is not defined as it tends to infinity. To be precise, as x approaches 0 from the left, the function drops towards $-\infty$, if instead x approaches 0 from the right, the function tends to $+\infty$. This is a first encounter with the concept of limit, which we will explore in more details in Section 1.3.

Remark 2 A fundamental relation in Microeconomics is the so-called utility function and the associated indifference curve. For a utility function $u(x)$, the indifference curve is defined as the set of bundles which yield a constant level of utility

$$\text{Indifference Curve} = \{x : u(x) = k\}.$$

Said differently, we are considering a collection of indifference curves, each one corresponding to a level of utility k . A popular choice is the so-called Cobb-Douglas utility function, which in its symmetric version has form

$$u(x, y) = xy,$$

with x and y denoting the two goods from which the economic agent extracts welfare. The corresponding indifference curve is then

$$y = \frac{k}{x}$$

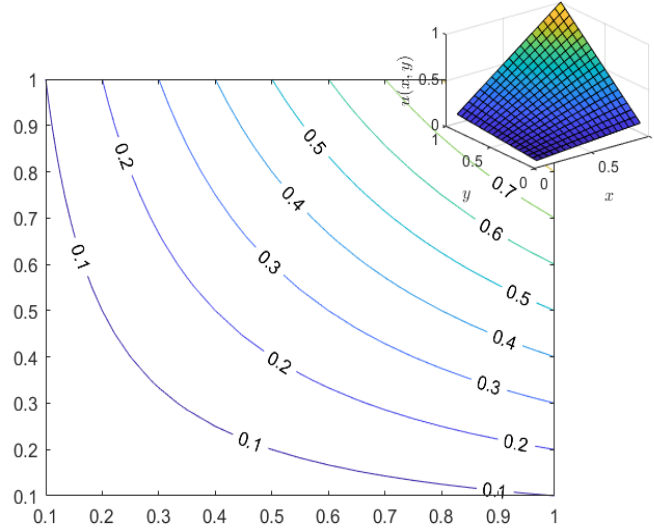


Figure 9: Indifference curves (main plot) and utility function (secondary plot). The levels of utility k are indicated in the main plot.

for a given utility level k . This is a multiple of order k of the reciprocal function (or hyperbola). Figure 9 shows the utility function $u(x, y) = xy$ in a 3-D space and the corresponding indifference curves, which are the related isoquants.

1.2.4 Inverse functions

Example 6 Consider the cubic function $y = x^3$. If $x = 2$ then $y = 8$. Suppose now you have $y = 8$ and you want to know the value of x which has generated y based on the cubic mapping. Then you need to express x as a function of y . This function is the cubic root, i.e. $x = \sqrt[3]{y}$.

Definition 2 Let f be a function whose domain is the set X , and whose range is the set Y . Then f is invertible if there exists a function g with domain Y and image X , with the property

$$f(x) = y \Leftrightarrow g(y) = x.$$

If f is invertible, the function g is unique, which means that there is exactly one function g satisfying this property. That function g is then called the inverse of f , and is usually denoted as f^{-1} .

A condition for f to be invertible is to be a one-to-one function, i.e. a function that never takes on the same value twice. In symbols

$$f(x_1) \neq f(x_2), \text{ whenever } x_1 \neq x_2.$$

If a horizontal line intersects the graph of f in more than one point, then there are numbers x_1 and x_2 such that $f(x_1) = f(x_2)$. This means that f is not one-to-one.

- Example 7** 1. The cubic function $f(x) = x^3$ is a one-to-one function, as we can observe from Figure 7 - right hand panel. Therefore it is invertible and for $y = x^3$, $f^{-1}(y) = \sqrt[3]{y}$, i.e. $x = \sqrt[3]{y}$.
2. The quadratic function $f(x) = x^2$ is not a one-to-one function: we can observe from Figure 7 - left hand panel that, if we draw an horizontal line, this will intersect the graph in two points. For example, if we draw the line at the point $y = 4$, we see that this line crosses the graph at both $x = -2$ and $x = 2$, i.e. for $y = f(x)$

$$f(-2) = 4 = f(2).$$

In order to be able to invert this function, we need to consider separately the cases in which $x \geq 0$ and $x \leq 0$. Then the quadratic function is one-to-one and we can invert. In this case we have

$$x = f^{-1}(y) = \sqrt{y} \text{ if } x \geq 0,$$

and

$$x = f^{-1}(y) = -\sqrt{y} \text{ if } x \leq 0.$$

This can be written in a compact way using the so called indicator function

$$1_A(\omega) = \begin{cases} 1 & \text{if } \omega \in A \\ 0 & \text{otherwise} \end{cases}$$

so that

$$x = f^{-1}(y) = \sqrt{y}1_{x \geq 0} - \sqrt{y}1_{x \leq 0}.$$

Remark 3 (A matter of notation: notation matters!) Do not mistake the $^{-1}$ in $f^{-1}(x)$ for an exponent. Thus

$$f^{-1}(x) \text{ does not mean } \frac{1}{f(x)}$$

The reciprocal $1/f(x)$ could, however, be written as $(f(x))^{-1}$. Mathematics is a very precise form of language, consequently we need to be as precise and respect the rules given by mathematics.

For the inverse function, the following ‘cancellation equations’ hold.

$$f^{-1}(f(x)) = x, \quad f(f^{-1}(y)) = y$$

for $y = f(x)$ and $x = f^{-1}(y)$.

- Example 8** 1. Suppose the function f is such that $f(1) = 5$, $f(3) = 7$ and $f(8) = -10$. Also assume that f is invertible. Then

$$\begin{aligned} f^{-1}(5) &= f^{-1}(f(1)) = 1 \\ f^{-1}(7) &= f^{-1}(f(3)) = 3 \\ f^{-1}(-10) &= f^{-1}(f(8)) = 8. \end{aligned}$$

2. Let $y = f(x) = x^n$. Then $x = f^{-1}(y) = y^{\frac{1}{n}}$. Indeed $(x^n)^{\frac{1}{n}} = x^{\frac{n}{n}} = x$. It is understood that for n even we consider $x \geq 0$.

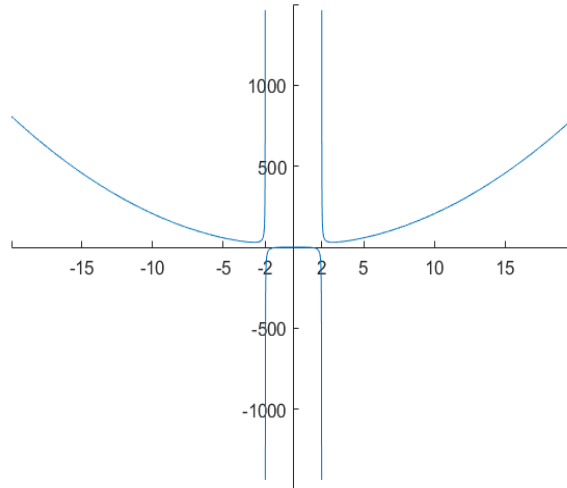


Figure 10: Rational function $f(x) = \frac{2x^4 - x^2 + 1}{x^2 - 4}$.

1.2.5 Rational functions

Rational functions have form

$$f(x) = \frac{P(x)}{Q(x)}$$

for polynomials P and Q . The domain consists of all values of x such that $Q(x) \neq 0$. A simple example is the reciprocal function $f(x) = 1/x$ which is defined for $x \neq 0$ (see Figure 8). Another example is the function

$$f(x) = \frac{2x^4 - x^2 + 1}{x^2 - 4}$$

with domain $X = \{x : x \neq -2, x \neq 2\}$. Its graph is shown in Figure 10.

1.2.6 Exponential and logarithmic functions

The exponential function has form

$$f(x) = a^x$$

for any positive constant a , known as the base. The base a can be any positive constant. Special cases of particular interest are $a = 2$, $a = 10$ (both are found in biology, computer science and engineering), and $a = e$, with e the Euler constant. This latter case is very common in finance.

Some examples of exponential functions are shown in the left hand side panel of Figure 11.

Properties of the exponential function (from the laws of exponents)

1. $a^{x+b} = a^x a^b$
2. $a^{x-b} = a^x a^{-b} = \frac{a^x}{a^b}$
3. $a^{bx} = (a^b)^x = c^x$, for $c = a^b$.

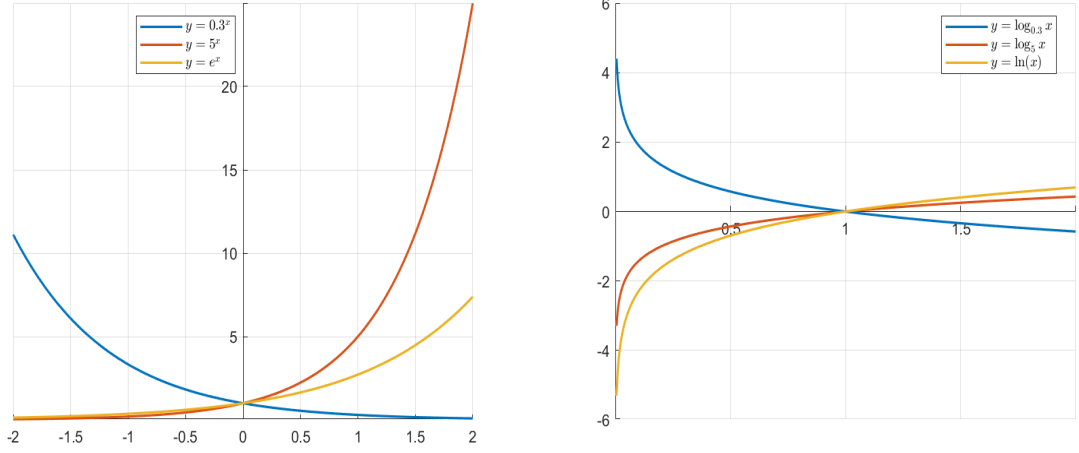


Figure 11: Exponential and logarithm function.

It is advisable to know these relationships by heart.

The logarithm function is defined as the inverse function of the exponential function. In other words.

Definition 3 For a any number such that $a > 0$ and $a \neq 1$, and $x > 0$, then

$$y = \log_a x \text{ is equivalent to } a^y = x.$$

We usually read this as 'log base a of x '.

It follows from the definition that the logarithm function is such that

$$y = \log_a(a^y).$$

A special case is when $a = e$, then we have the so-called natural logarithm $y = \ln x$, which is the inverse function of $x = e^y$.

Some log-functions are graphed in Figure 11 (right hand side panel).

Properties of the logarithm - from the properties of the exponential function

1. $\log_a(xb) = \log_a(x) + \log_a(b)$.

Indeed, let $y = \log_a x$ and $z = \log_a b$; then $x = a^y$ and $b = a^z$. From property (1) of the exponential function, we know that $xb = a^{y+z}$. Take logarithm both sides of the equality to obtain $\log_a(xb) = \log_a(a^{y+z})$; the definition of logarithm implies the result.

Note

$$\log_a(x) + \log_a(b) \neq \log_a(x + b).$$

2. $\log_a\left(\frac{x}{b}\right) = \log_a(x) - \log_a(b)$.

Indeed, let $y = \log_a x$ and $z = \log_a b$; then $x = a^y$ and $b = a^z$. From property (2) of the exponential function, we know that $x/b = a^{y-z}$. Take logarithm both sides of the equality to obtain $\log_a(x/b) = \log_a(a^{y-z})$; the definition of logarithm implies the result.

Note

$$\log_a(x) - \log_a(b) \neq \log_a(x - b).$$

3. $\log_a(x^b) = b \log_a(x)$.

Indeed, let $y = \log_a x$; then $x = a^y$. From property (3) of the exponential function, we know that $x^b = a^{by}$. Take logarithm both sides of the equality to obtain $\log_a(x^b) = \log_a(a^{by})$; the definition of logarithm implies the result.

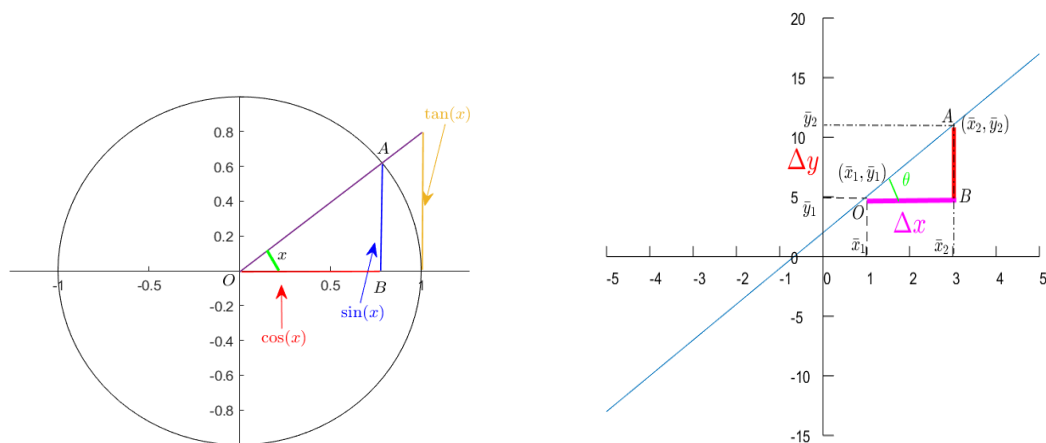
1.2.7 Trigonometric functions

A note about unit of measure with angles The radian (symbol: rad) is the SI unit for measuring angles, and is the standard unit of angular measure used in many areas of mathematics. The length of an arc of a unit circle is numerically equal to the measurement in radians of the angle that it subtends. Thus, in practice. Assume a unit circle. The angle it subtends is 360° . The arc in question is the full circumference which measures 2π . Therefore the radian is

$$\text{rad} = \frac{180}{\pi^\circ}.$$

Then, the 90° angle measures $\pi/2$ in radians. In calculus the convention is that radian measure is always used (except when otherwise indicated). For example, when we use the function $\sin(x)$, it is understood as the sine of the angle whose radian measure is x .

Figure 12: Left hand side panel: The unit circle and the trigonometric functions. Right hand side panel: lines and slopes.



Trigonometric functions (also called circular functions, angle functions or goniometric functions) are real functions which relate an angle of a right-angled triangle to ratios of two side lengths. The most commonly used functions are sine, cosine and tangent.

We define these function with respect to the unit circle, with centre in the origin, O , as illustrated in Figure 12. Inside this circle, consider the right-angled triangle AOB , and its acute angle x ; by construction the hypotenuse of this triangle is the segment AO which has length 1. Then by definition we have

$$\sin(x) = \frac{AB}{AO} = AB; \quad \cos(x) = \frac{OB}{AO} = OB.$$

These quantities are shown in the left hand side panel of Figure 12, which highlights that $\cos(x), \sin(x)$ represent the coordinates of the point A .

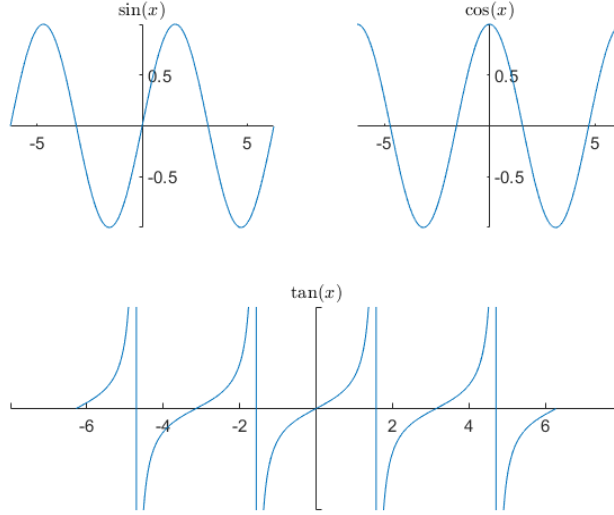


Figure 13: The fundamental trigonometric functions: $\sin(x)$, $\cos(x)$ and $\tan(x)$.

A fundamental relationship linking these functions follows from Pythagoras' theorem (reminder: the unit circle has radius 1)

$$\sin^2(x) + \cos^2(x) = 1.$$

Now, we can re-interpret the meaning of the slope coefficient of a line (right-hand side panel of Figure 12). Consider the line passing through the points O and A . Let θ denote the angle subtended by this line. We know from Section 1.2.1 that the slope is

$$b = \frac{AB}{OB},$$

i.e.

$$b = \frac{\sin(\theta)}{\cos(\theta)}.$$

For simplicity, but without loss of generality, suppose we translate the Cartesian axes to the point O , so that this becomes our new origin. Then, the line passing through the new origin O and A has equation

$$y = bx = \frac{\sin(\theta)}{\cos(\theta)}x.$$

Now. Assume $x = 1$. The corresponding y value is precisely

$$\frac{\sin(\theta)}{\cos(\theta)}.$$

This is what we define to be the tangent of θ , i.e.

$$\tan(\theta) = \frac{AB}{OB} = \frac{\sin(\theta)}{\cos(\theta)}.$$

All the trigonometric functions are represented in Figure 13. We note that function $y = \tan(x)$ has asymptotes corresponding to multiples of $\pi/2$ (can you explain why?).

We can also define the corresponding inverse functions, i.e.

1. $x = \arcsin(y)$ (arcsine) i.e. the function such that $y = \sin(x)$. The arcsine is defined on $-1 \leq y \leq 1$.
2. $x = \arccos(y)$ (arccosine) i.e. the function such that $y = \cos(x)$. The arccosine is defined on $-1 \leq y \leq 1$.
3. $x = \arctan(y)$ (arctangent) i.e. the function such that $y = \tan(x)$. The arctangent is defined on the extended real line to include $\pm\infty$.

Note the restrictions we imposed on the domain of the inverse functions. Can you explain why?

1.2.8 Composite functions

Two functions f and g can be combined to form new functions by applying to them any of the basic operations, exponentiation, and so on. We obtain what is called a composite function.

Definition 4 Given two functions f and g , the composite function $f \circ g$ (also called the composition of f and g) is defined by

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

The domain of $f \circ g$ is the set of all x in the domain of g such that $g(x)$ is in the domain of f . In other words, $(f \circ g)(x)$ is defined whenever both f and g are defined.

Remark 4 (A matter of notation: notation matters!) Remember, the $f \circ g$ notation means that the function g is applied first and then f is applied second. Consequently

$$f \circ g \neq g \circ f,$$

the only exception being the case of the identity function.

Example 9 1. Let $f(x) = \sqrt{x}$ and $g(x) = x^2 + 1$. Then

$$F(x) = f(g(x)) = f(x^2 + 1) = \sqrt{x^2 + 1}.$$

The domain of this function is the entire real line as $x^2 + 1$ is always positive, i.e. $X = \mathbb{R}$. However

$$G(x) = g(f(x)) = g(\sqrt{x}) = x + 1$$

with domain $X = \{x : x \geq 0\} = [0, +\infty)$.

2. Let $f(x) = \sqrt{x}$ and $g(x) = \sqrt{2 - x}$. Then

$$F(x) = f(g(x)) = f(\sqrt{2 - x}) = \sqrt{\sqrt{2 - x}} = \left((2 - x)^{\frac{1}{2}}\right)^{\frac{1}{2}} = \sqrt[4]{2 - x}.$$

The domain of this function is $X = \{x : x \leq 2\} = (-\infty, 2]$.

However

$$G(x) = g(f(x)) = g(\sqrt{x}) = \sqrt{2 - \sqrt{x}};$$

note that $\sqrt{x}$ is defined for $x \geq 0$ and $\sqrt{2 - \sqrt{x}}$ is defined for $2 - \sqrt{x} \geq 0$ i.e. $x \leq 4$. Consequently the domain is $X = [0, 4]$.

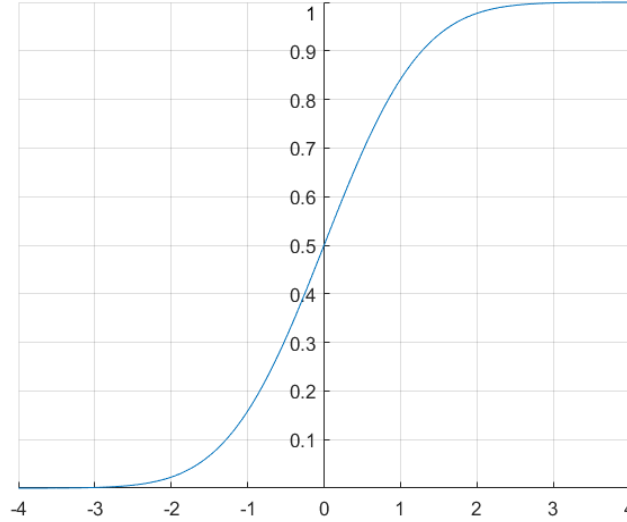


Figure 14: The cumulative distribution function of the standard normal (or Gaussian) distribution.

1.2.9 Other functions

There are a number of other functions that have applications in finance, and a quite large class of these is the one of transcendental functions.

Transcendental functions are usually defined as an integral or a solution to a specific differential equation. The most important example is represented by the cumulative distribution function of the standard normal (or Gaussian) random variable. This is defined as

$$\Phi(x) = \int_{-\infty}^x \frac{1}{\sqrt{2\pi}} e^{-\frac{y^2}{2}} dy.$$

The function is plotted in Figure 14, and we will look at it again in Section 1.6.5.

1.3 Limits

Example 10 Consider the function

$$f(x) = \frac{1}{x(x-1)}$$

which we saw in Example 2.2. In that example we have shown that the domain of this function is $X = \{x : x \neq 0 \text{ and } x \neq 1\}$, as for those values the denominator is zero and therefore the fraction is not defined. So, how can we determine what happens to this function as x gets closer and closer to either 0 or 1? In order to answer to this question, we need to be able to calculate the limit of the function $f(x)$ as x tends to either 0 or 1.

Definition 5 Suppose $f(x)$ is defined when x is near point a . Then

$$\lim_{x \rightarrow a} f(x) = L$$

which reads as ‘the limit of $f(x)$ as x approaches a equals L ’.

The above is equivalent to say: we can make the values of $f(x)$ arbitrarily close to L , by considering X sufficiently close to a (on either side of a) but not equal to a .

In some cases the limit can be simple to compute.

Example 11 1. Let

$$f(x) = \frac{x-1}{x^2-1}.$$

Assume we are interested in understanding the behaviour of the function in 1. By noticing that $x^2 - 1 = (x-1)(x+1)$, we realize that

$$f(x) = \frac{1}{x+1}$$

and

$$\lim_{x \rightarrow 1} f(x) = \frac{1}{2}.$$

2. Consider now

$$f(x) = \frac{x^2-1}{x-1}.$$

Assume we are still interested in understanding the behaviour of the function in 1. By noticing that $x^2 - 1 = (x-1)(x+1)$, we realize that

$$f(x) = x+1$$

and

$$\lim_{x \rightarrow 1} f(x) = 2.$$

3. Consider the reciprocal function

$$f(x) = \frac{1}{x}$$

defined for $x \neq 0$. As x approaches 0 from the right, then x takes values 1, 0.1, 0.01, 0.001 and so on. The corresponding values of the function f are 1, 10, 100, 1000, i.e. they keep increasing. However, we will never be able to actually reach zero on the x -axis, and the function asymptotically reaches $+\infty$. We write this situation in mathematical terms as

$$\lim_{x \rightarrow 0+} f(x) = \lim_{x \rightarrow 0+} \frac{1}{x} = +\infty.$$

(The ‘+’ next to the 0 means that we are approaching from the right.) Suppose now that we approach 0 from the left, so that x takes values -1, -0.1, -0.01, -0.001 and so on. The corresponding values of the function $f(x)$ are -1, -10, -100, -1000, and the function asymptotically reaches $-\infty$, i.e.

$$\lim_{x \rightarrow 0-} f(x) = \lim_{x \rightarrow 0-} \frac{1}{x} = -\infty.$$

(The ‘-’ next to the 0 means that we are approaching from the left.) You can refer to Figure 8 for the graphical illustration.

4. Back to the function

$$f(x) = \frac{1}{x(x-1)},$$

we have

$$\lim_{x \rightarrow 0+} f(x) = \lim_{x \rightarrow 0+} \frac{1}{x(x-1)} \rightarrow -\infty, \quad \lim_{x \rightarrow 0-} f(x) = \lim_{x \rightarrow 0-} \frac{1}{x(x-1)} \rightarrow +\infty,$$

and

$$\lim_{x \rightarrow 1+} f(x) = \lim_{x \rightarrow 1+} \frac{1}{x(x-1)} \rightarrow +\infty, \quad \lim_{x \rightarrow 1-} f(x) = \lim_{x \rightarrow 1-} \frac{1}{x(x-1)} \rightarrow -\infty.$$

The function is graphed in Figure 15.

The above examples concern points in which the function is not defined, i.e. points which are excluded from the domain X . Also $\pm\infty$ are technically excluded from the domain as they are not in $\mathbb{R}$. If we want to understand the behaviour of the function as x becomes larger and larger, we need to look at the limit for $x \rightarrow \pm\infty$.

Example 12 1. The reciprocal function $1/x$ tends to 0 as x becomes larger and larger: for $x = 10, 100, 1000$, in fact, we have $f(x) = 0.1, 0.01, 0.001$ and for $x = -10, -100, -1000$, $f(x) = -0.1, -0.01, -0.001$. As shown in Figure 8 the Cartesian axes are the asymptotes, i.e. the lines such that the distance between the curve and the line approaches zero as one or both of the x or y coordinates tends to infinity. In proper mathematical terms

$$\lim_{x \rightarrow -\infty} \frac{1}{x} = 0; \quad \lim_{x \rightarrow +\infty} \frac{1}{x} = 0.$$

2. For the function

$$f(x) = \frac{1}{x(x-1)},$$

we have that

$$\lim_{x \rightarrow \infty} f(x) = 0.$$

The function is graphed in Figure 15.

Laws of limits Let c be a constant.

1. $\lim_{x \rightarrow a} c = c$; and $\lim_{x \rightarrow a} x = a$.
2. $\lim_{x \rightarrow a} (f(x) + g(x)) = \lim_{x \rightarrow a} f(x) + \lim_{x \rightarrow a} g(x)$.
3. $\lim_{x \rightarrow a} (f(x) - g(x)) = \lim_{x \rightarrow a} f(x) - \lim_{x \rightarrow a} g(x)$.
4. $\lim_{x \rightarrow a} cf(x) = c \lim_{x \rightarrow a} f(x)$.
5. $\lim_{x \rightarrow a} (f(x)g(x)) = (\lim_{x \rightarrow a} f(x))(\lim_{x \rightarrow a} g(x))$.
6. $\lim_{x \rightarrow a} \frac{f(x)}{g(x)} = \frac{\lim_{x \rightarrow a} f(x)}{\lim_{x \rightarrow a} g(x)}$.
7. $\lim_{x \rightarrow a} x^n = a^n$.

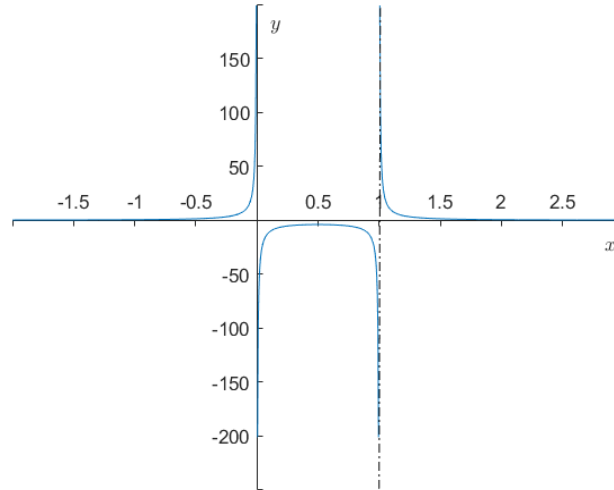


Figure 15: $f(x) = \frac{1}{x(x-1)}$. The asymptote $x = 1$ is shown in dash-dot line.

8. $\lim_{x \rightarrow a} f(x)^n = (\lim_{x \rightarrow a} f(x))^n$.
9. $\lim_{x \rightarrow a} \sqrt[n]{x} = \sqrt[n]{a}$ (if n is even, then we assume $a \geq 0$).
10. $\lim_{x \rightarrow a} \sqrt[n]{f(x)} = \sqrt[n]{\lim_{x \rightarrow a} f(x)}$ (if n is even, then we assume $\lim_{x \rightarrow a} f(x) \geq 0$).
11. DIRECT SUBSTITUTION PROPERTY: If f is a polynomial or a rational function and a is in the domain of f , then $\lim_{x \rightarrow a} f(x) = f(a)$.
12. SQUEEZE THEOREM: Suppose that for every x not equal to a , we have $g(x) \leq f(x) \leq h(x)$, and also suppose that

$$\lim_{x \rightarrow a} g(x) = \lim_{x \rightarrow a} h(x) = L.$$

Then

$$\lim_{x \rightarrow a} f(x) = L.$$

Example 13 1. Consider the the rational function

$$f(x) = \frac{5x^2 + 1}{3x^2 - x},$$

we want to study its behaviour for large values of x , i.e. for $x \rightarrow \infty$. We write then the corresponding limit

$$\lim_{x \rightarrow \infty} \frac{5x^2 + 1}{3x^2 - x}.$$

Both numerator and denominator are polynomial of the same order 2. This is a very useful hint: it is like to say that in the limit we are looking at infinity of the same order. Then,

if this is true, the limit might be finite. Indeed this is the case: you need first to factorize the higher order power of x and then simplify as follows

$$\begin{aligned}
 \lim_{x \rightarrow \infty} \frac{5x^2 + 1}{3x^2 - x} &= \lim_{x \rightarrow \infty} \frac{x^2 \left(5 + \frac{1}{x^2}\right)}{x^2 \left(3 - \frac{x}{x^2}\right)} \\
 &= \lim_{x \rightarrow \infty} \frac{x^2 \left(5 + \frac{1}{x^2}\right)}{x^2 \left(3 - \frac{1}{x}\right)} \\
 &= \lim_{x \rightarrow \infty} \frac{\left(5 + \frac{1}{x^2}\right)}{\left(3 - \frac{1}{x}\right)} \\
 &= \frac{5 + \lim_{x \rightarrow \infty} \frac{1}{x^2}}{3 - \lim_{x \rightarrow \infty} \frac{1}{x}} \\
 &= \frac{5}{3}.
 \end{aligned}$$

We applied Rule 6 for the calculation of the limit of the fraction.

2. Consider the function

$$f(x) = \begin{cases} x + 3 & \text{if } x \leq 2 \\ -x + 7 & \text{if } x > 2. \end{cases}$$

Then

$$\lim_{x \rightarrow 2} f(x) = 5$$

by substitution of the value 2 in $x + 3$ which is the branch of the function for which the value 2 belongs to the domain.

Finally, there is a number of ‘notable limits’ which are quite convenient and it is advisable to remember them at all times. The most important are listed in Table 3.

Table 3: Notable limits

$\lim_{x \rightarrow \infty} \left(1 + \frac{1}{x}\right)^x = e$	$\lim_{x \rightarrow 0} (1 + x)^{\frac{1}{x}} = e$
$\lim_{x \rightarrow 0} \frac{\log_a(1+x)}{x} = \frac{1}{\ln a}$	$\lim_{x \rightarrow 0} \frac{\ln(1+x)}{x} = 1$
$\lim_{x \rightarrow 0} \frac{a^x - 1}{x} = \ln a$	$\lim_{x \rightarrow 0} \frac{e^x - 1}{x} = 1$
$\lim_{x \rightarrow 0} \frac{\sin(x)}{x} = 1$	$\lim_{x \rightarrow 0} \frac{1 - \cos(x)}{x} = 0$

Example 14 The compound interest law is given by

$$V = P \left(1 + \frac{r}{m}\right)^{mt}.$$

In Example 1 we have seen that as the compounding frequency m increases, V increases - as also illustrated in Figure 1. What happens in the limit? Using the table of notable limits and the laws of the limits we have

$$\begin{aligned}
 \lim_{m \rightarrow \infty} P \left(1 + \frac{r}{m}\right)^{mt} &= P \lim_{m \rightarrow \infty} \left(1 + \frac{r}{m}\right)^{mt} \quad (\text{rule 4}) \\
 &= P \left(\lim_{m \rightarrow \infty} \left(1 + \frac{r}{m}\right)^m \right)^t \quad (\text{rule 8}) \\
 &= P \left(\lim_{x \rightarrow 0} (1+x)^{\frac{r}{x}} \right)^t \quad (\text{set } x = \frac{r}{m}) \\
 &= P \left(\lim_{x \rightarrow 0} (1+x)^{\frac{1}{x}} \right)^{rt} \quad (\text{rules of exponents}) \\
 &= P e^{rt} \quad (\text{notable limit}).
 \end{aligned}$$

This is the continuous compounding rule.

Continuous functions The function f is continuous at some point c of its domain if

$$\lim_{x \rightarrow c} f(x) = f(c).$$

The above statement is meaningful if the following holds: f is defined at c , the limit on the left hand side of the above equation exists, and finally, the value of this limit equals $f(c)$.

Example 15 Let us investigate the continuity of the following functions at the origin, i.e. for $c = 0$.

1. $f(x) = x^2$.

We note by direct substitution that $f(0) = 0$ and indeed $\lim_{x \rightarrow 0} x^2 = 0$. Thus the function is continuous.

2. $f(x) = |x|$.

We start by noticing that

$$f(x) = |x| = \begin{cases} x & \text{if } x \geq 0 \\ -x & \text{if } x < 0 \end{cases}$$

Therefore, $\lim_{x \rightarrow 0+} |x| = \lim_{x \rightarrow 0+} x = 0$, and $\lim_{x \rightarrow 0-} |x| = \lim_{x \rightarrow 0-} (-x) = 0$. Consequently, $\lim_{x \rightarrow 0} |x| = 0$ which is, of course, equal to $f(0)$. Hence the function is continuous.

3. $f(x) = \frac{1}{x}$

The function is not continuous in $c = 0$ as this point does not belong to the domain of f , see also Example 11.3.

1.4 Derivatives

The derivative of a function of a real variable measures the sensitivity to changes of the function value (output) with respect to a change in its argument (input). The process of finding a derivative is called differentiation.

Definition 6 The derivative of a function f at a point c , denoted by $f'(c)$, is

$$f'(c) = \lim_{h \rightarrow 0} \frac{f(c+h) - f(c)}{h}$$

if this limit exists.

Example 16 1. Consider the function $f(x) = x^2$. Let us apply the definition in order to find its derivative. Then

$$\begin{aligned} f'(x) &= \lim_{h \rightarrow 0} \frac{(x+h)^2 - x^2}{h} \\ &= \lim_{h \rightarrow 0} \frac{x^2 + 2xh + h^2 - x^2}{h} \\ &= 2x + \lim_{h \rightarrow 0} h = 2x. \end{aligned}$$

Thus, the limit exists and $f'(x) = 2x$ for $f(x) = x^2$.

2. Consider the function $f(x) = |x|$. We know

$$f(x) = |x| = \begin{cases} x & \text{if } x \geq 0 \\ -x & \text{if } x < 0 \end{cases}$$

From the definition, we have

$$f'(x) = \lim_{h \rightarrow 0} \frac{|x+h| - |x|}{h} = \frac{|h|}{h}.$$

However, due to the definition of modulus function, it follows that

$$f'(x) = \begin{cases} 1 & \text{if } h > 0 \\ -1 & \text{if } h < 0 \end{cases}$$

In other words, the limit from the right is 1, while the limit from the left is -1. Thus the function $f(x) = |x|$ is not differentiable.

The latter example shows that in general differentiability is a more restrictive condition than continuity, because it requires something more than continuity. Continuity at a point rules out the presence of a gap, whereas differentiability rules out “sharpness” as well. Therefore, differentiability calls for “smoothness” of a function as well as its continuity. If a unique limit in Definition 6 does not exist, the derivative does not exist.

1.4.1 Meaning of the derivative

From the definition, set $x = c + h$, then

$$f'(c) = \lim_{x \rightarrow c} \frac{f(x) - f(c)}{x - c},$$

in other words, if the input variable changes by an amount h from c to x , the functions changes from $f(c)$ to $f(x)$. Therefore

$$\frac{f(x) - f(c)}{x - c}$$

is the incremental ratio for the function f , and $f'(c)$ measures this ratio in the case in which the increment in the input, i.e. h , is very small; so small actually that is infinitesimal (meaning of $h \rightarrow 0$). Thus, the derivative returns the change in the function due to an infinitesimal change in the input.

To gain more insight in this very important concept, let us consider the following examples.

Example 17 1. Let us refer back to Example 4 and the corresponding Figure 4, in which we consider the linear function $y = f(x)$ with $f(x) = 2 + 3x$. The slope b is defined as the incremental ratio

$$\frac{\Delta y}{\Delta x}.$$

For the points given in Example 4 then

$$b = \frac{11 - 5}{3 - 1} = 3.$$

This is saying that if x changes by 2 units (3-1), then y changes by 6 units (11-5) or, equivalently, if x changes by 1 unit, then y changes by 3 units.

Note that this ratio is constant, i.e. it does not change when x changes. This is typical of a linear function.

2. Let us now consider the function

$$f(x) = \frac{1}{2}x^2.$$

Suppose we are interested in obtaining the equation of the line passing through the point $(2, 2)$ which is tangent to $f(x)$, i.e. the line that only ‘touches’ $f(x)$ once, without ‘cutting’ it. However, from Section 1.2.1 we see that we have too little information as we only know one point, and this tangent line might not necessarily pass through the origin. What we could do to solve this problem is to choose arbitrarily other points on $f(x)$ and draw the line joining these points and $(2, 2)$. These lines are called secants (these lines cut through the given function $f(x)$ at a minimum of two points). Suppose we consider the point $(-3, 9/2)$. Then from Section 1.2.1, we obtain $y = 3 - x/2$, thus the rate of change per 1 unit is -0.5 . Now, let us consider the origin $(0, 0)$. The secant line through $(0, 0)$ and $(2, 2)$ is $y = x$, i.e. the identity function, so that the rate of change is 1. The current value $x = 0$ is ‘only’ 2 units away from the point of interest. Therefore, in order to be able to obtain the tangent line, we will need to repeat this procedure for x approaching 2 as close as possible. This is illustrated in Figure 16. This example tells us that the slope of the tangent line to $f(x)$ in $(2, 2)$ is obtained as

$$f'(2) = \lim_{x \rightarrow 2} \frac{f(x) - f(2)}{x - 2},$$

which is the derivative of $f(x)$ in 2.

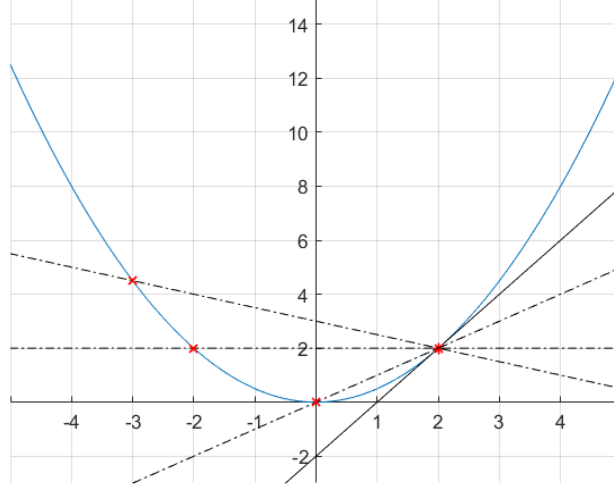


Figure 16: Secants (dash-dot lines) and tangent (solid line) in $(2, 2)$ for $f(x) = \frac{1}{2}x^2$.

The previous example shows that to differentiate a function f in a point c is equivalent to determining the slope of the tangent line to f in c .

The definition of derivative implies that the derivative itself is a function of the original input, and therefore varies according to the values of x on the domain of f . Indeed, the definition implies

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}.$$

The idea is shown in Figure 17, which shows different tangent lines of $f(x) = \frac{1}{2}x^2$ at different points x . We note the different slopes: negative at first, in correspondence of the values of x such that $f(x)$ is decreasing, then flat, finally positive, in correspondence of those values of x such that $f(x)$ is increasing.

Indeed, this is one of the many pieces of information that a derivative can provide us regarding the behaviour of a function: for which values of x $f(x)$ is either increasing, decreasing or not changing at all (i.e. flat). This specific information can be read from the sign of the derivative. Let us consider for example the case in which $f'(x) < 0$. Then, from the definition of derivative, it follows that for small $h > 0$ then

$$\frac{f(x+h) - f(x)}{h} < 0,$$

i.e. $f(x+h) < f(x)$ with $x+h > x$. This implies that the function $f(x)$ is decreasing. Viceversa, if $f'(x) > 0$, then

$$\frac{f(x+h) - f(x)}{h} > 0,$$

i.e. $f(x+h) > f(x)$ with $x+h > x$. This implies that the function $f(x)$ is increasing.

Example 18 Consider the function $f(x) = x^2$. We have seen in Example 16.1 that $f'(x) = 2x$. Thus, $f'(x) < 0$, i.e. $f(x)$ is decreasing, if $x < 0$, $f'(x) = 0$, i.e. the function does not change, if $x = 0$, and $f'(x) > 0$, i.e. $f(x)$ is increasing, if $x > 0$.

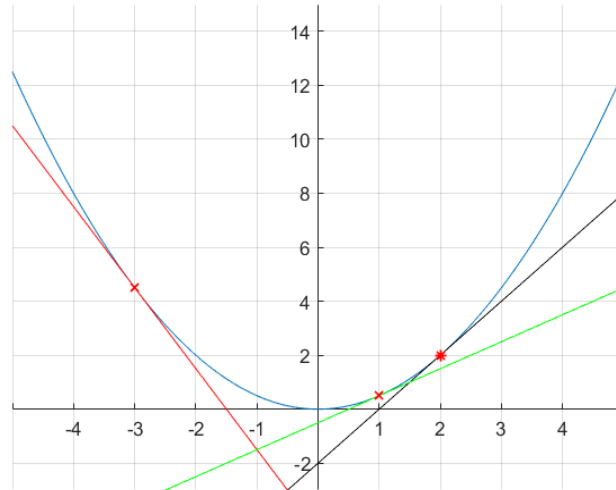


Figure 17: $f(x) = \frac{1}{2}x^2$ and $f'(x)$. Points of interest $(-3, 9/2)$, $(1, 1/2)$, $(2, 2)$.

Note. The point at which the derivative is zero is called stationary point. In the previous example, $x = 0$ is the stationary point for the parabola.

The above proves that we can learn quite a lot about the behaviour of $f(x)$ by looking at its derivative. For this reason, we need to learn how to differentiate. This will be covered next. Before this, though, a matter of notation: the derivative f' of the function f can also be denoted as

$$\frac{dy}{dx} = f'(x) \quad (2)$$

for $y = f(x)$; alternatively, you can find it denoted as

$$\frac{d}{dx}f(x).$$

1.4.2 Rules of differentiation

The rules are accompanied by a short intuitive proof, which aims firstly at giving insight into the result, and secondly gaining practice with the concepts introduced in Section 1.2. Note that each rule can be used also in combination with others. All rules are summarized in Table 4 and need to be known.

1. **CONSTANT FUNCTION:** If $f(x) = c$, for any constant c , then $f'(x) = 0$.
Indeed, by using the definition

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{c - c}{h} = 0.$$

Example 19 Consider the function $f(x) = 5$. Then, $f'(x) = 0$.

2. POWER RULE: If $f(x) = x^n$, any real number n , then $f'(x) = nx^{n-1}$.

Indeed, from the definition

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{(x+h)^n - x^n}{h}. \quad (3)$$

From the binomial theorem, which we met in Section 1.2.2 - equation (1), we know that

$$(x+h)^n = \sum_{k=0}^n \frac{n!}{k!(n-k)!} h^k x^{n-k}$$

i.e.

$$(x+h)^n = x^n + nx^{n-1}h + \frac{1}{2}n(n-1)x^{n-2}h^2 + \cdots + nxh^{n-1} + h^n,$$

therefore, by substituting in (3) and simplifying, we obtain

$$f'(x) = \lim_{h \rightarrow 0} \frac{(x+h)^n - x^n}{h} = nx^{n-1} + \lim_{h \rightarrow 0} \left(\frac{1}{2}n(n-1)x^{n-2}h + \cdots + nxh^{n-2} + h^{n-1} \right)$$

from which the result follows as all terms in h are zero.

Example 20 (a) Consider $f(x) = x^5$. Then $f'(x) = 5x^{5-1} = 5x^4$.

(b) Consider $f(x) = \sqrt[5]{x}$. Note that $f(x) = x^{1/5}$. Then $f'(x) = \frac{1}{5}x^{1/5-1} = \frac{1}{5}\sqrt[5]{x^{-4}}$.

3. CONSTANT TIMES A FUNCTION: If $f(x) = cg(x)$, where c is a constant and g is a differentiable function, then $f'(x) = cg'(x)$.

This is a straightforward application of the rules of limits and the definition of derivative.

Example 21 (a) Consider $f(x) = 10x^5$. Then $f'(x) = 10(5x^{5-1}) = 50x^4$.

(b) Consider $f(x) = 3x$. Then $f'(x) = 3x^{1-1} = 3$.

(c) More in general, let $f(x) = a + bx$, then $f'(x) = b$.

4. SUM OR DIFFERENCE OF FUNCTIONS: If $f(x) = u(x) \pm g(x)$, with u and g differentiable functions, then $f'(x) = u'(x) \pm g'(x)$.

This result follows from the additivity law of limits.

Example 22 Consider $f(x) = 8x^3 - 4x^2 + 5x - 10$. Then $f'(x) = 24x^2 - 8x + 5$.

5. PRODUCT RULE: If $f(x) = u(x)v(x)$, with u and v differentiable functions, then $f'(x) = u'(x)v(x) + u(x)v'(x)$.

To show this, write first the incremental ratio

$$\frac{f(x+h) - f(x)}{h} = \frac{u(x+h)v(x+h) - u(x)v(x)}{h},$$

then, at the numerator add and subtract the amount $u(x+h)v(x)$ (this is a very useful trick which does not alter the expression but allows to simplify considerably the calculations) so that

$$\begin{aligned}\frac{f(x+h) - f(x)}{h} &= \frac{u(x+h)v(x+h) - u(x+h)v(x) + u(x+h)v(x) - u(x)v(x)}{h} \\ &= u(x+h) \frac{v(x+h) - v(x)}{h} + \frac{u(x+h) - u(x)}{h} v(x).\end{aligned}$$

Then, take the limit for $h \rightarrow 0$ and apply rule 4 first, and the definition of derivative to conclude.

Example 23 Consider $f(x) = (x^2 - 5)(x - x^3)$. Then, $u(x) = x^2 - 5$ and $v(x) = x - x^3$. Consequently, rule 2 implies that $u'(x) = 2x$, rule 2 and rule 4 imply that $v'(x) = 1 - 3x^2$. Now put the pieces together

$$f'(x) = 2x(x - x^3) + (x^2 - 5)(1 - 3x^2) = -5x^4 + 18x^2 - 5.$$

6. **CHAIN RULE:** the chain rule applies to composite functions. Thus, consider $F(x) = f(g(x))$ with both f, g differentiable functions. Then $F'(x) = f'(g(x))g'(x)$. The heuristic argument works as follows. Let $u = g(x)$ and $y = f(u)$. Let $\Delta u = g(x+h) - g(x)$ and $\Delta y = f(u+h) - f(u)$. Note that for $h \rightarrow 0$, then $\Delta u \rightarrow 0$. Now, let us consider our derivative

$$\begin{aligned}F'(x) &= \lim_{h \rightarrow 0} \frac{\Delta y}{h} \\ &= \lim_{h \rightarrow 0} \frac{\Delta y}{h} \frac{\Delta u}{\Delta u} \quad \left(\text{multiply by } \frac{\Delta u}{\Delta u} \right) \\ &= \lim_{h \rightarrow 0} \frac{\Delta y}{\Delta u} \frac{\Delta u}{h} \\ &= \left(\lim_{\Delta u \rightarrow 0} \frac{\Delta y}{\Delta u} \right) \left(\lim_{h \rightarrow 0} \frac{\Delta u}{h} \right) \\ &= f'(g(x))g'(x),\end{aligned}$$

where the last but one equality follows from the laws of limits, and the last equality follows from the definition of derivative.

Example 24 Consider $F(x) = \sqrt{x^2 + 1}$ so that $g(x) = x^2 + 1$ and $f(u) = \sqrt{u} = u^{1/2}$. Then, $g'(x) = 2x$ from rule 2, and $f'(u) = \frac{1}{2}u^{-1/2}$ again from rule 2. Then

$$F'(x) = \frac{1}{2}(x^2 + 1)^{-1/2} 2x = \frac{x}{\sqrt{x^2 + 1}}.$$

The chain rule allows us to generalize all the results we established above.

Example 25 POWER OF A FUNCTION: $f(x) = u(x)^n$, for a differentiable function u and n a real number has derivative $f'(x) = n u(x)^{n-1} u'(x)$

7. QUOTIENT RULE: If

$$f(x) = \frac{u(x)}{v(x)},$$

with u and v differentiable functions, then

$$f'(x) = \frac{u'(x)v(x) - u(x)v'(x)}{v(x)^2}.$$

The quotient rule can be proven using the product rule in conjunction with the chain rule by writing the function f as $f(x) = u(x)v(x)^{-1}$.

Example 26 Consider

$$f(x) = \frac{10 - x}{x^2 + 1}.$$

Then

$$f'(x) = \frac{-(x^2 + 1) - (10 - x)2x}{(x^2 + 1)^2} = \frac{x^2 - 20x - 1}{(x^2 + 1)^2}.$$

8. RULE OF EXPONENTIAL FUNCTIONS: If $f(x) = a^x$, for $a > 0$, then $f'(x) = a^x \ln a$.
Indeed, from the definition

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{a^{x+h} - a^x}{h}.$$

You can now factorize a^x and use the laws of limits to write

$$f'(x) = a^x \lim_{h \rightarrow 0} \frac{a^h - 1}{h} = a^x \ln a$$

in virtue of the notable limit for the exponential function (see Table 3).

The chain rule allows us to generalize the above to the case

$$f(x) = a^{u(x)}$$

for which

$$f'(x) = a^{u(x)} u'(x) \ln a$$

holds. The case of the base e exponential functions $f(x) = e^x$ and $f(x) = e^{u(x)}$ follows by setting $a = e$ and noticing that $\ln e = 1$.

Example 27 Consider $f(x) = 10e^{5x^2-4x}$. Then $u(x) = 5x^2 - 4x$ for which $u'(x) = 10x - 4$ and

$$f'(x) = 10(10x - 4)e^{5x^2-4x} = (100x - 40)e^{5x^2-4x}.$$

9. RULE OF THE LOG-FUNCTIONS: If $f(x) = \log_a x$, for $a > 0$, $a \neq 1$, and $x > 0$, then

$$f'(x) = \frac{1}{x \ln a}.$$

Indeed, from the definition

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{\log_a(x+h) - \log_a x}{h}.$$

You can now use the properties of the log-function to write

$$f'(x) = \lim_{h \rightarrow 0} \frac{\log_a \left(1 + \frac{h}{x}\right)}{h}.$$

At this point you can multiply and divide the above expression by $1/x$ so that

$$f'(x) = \frac{1}{x} \lim_{h \rightarrow 0} \frac{\log_a \left(1 + \frac{h}{x}\right)}{\frac{h}{x}};$$

the notable limit for the log-functions (see Table 3) implies the result.

The chain rule allows us to generalize the above to the case

$$f(x) = \log_a u(x)$$

for which

$$f'(x) = \frac{1}{u(x) \ln a} u'(x)$$

holds. The case of the base e exponential functions $f(x) = \ln x$ and $f(x) = \ln u(x)$ follows by setting $a = e$ and noticing that $\ln e = 1$.

Example 28 Consider $f(x) = \ln(4x^2 - 16x)$, then $u(x) = 4x^2 - 16x$, for which $u'(x) = 8x - 16$. Therefore

$$f'(x) = \frac{8x - 16}{4x^2 - 16x} = \frac{2x - 4}{x^2 - 4}.$$

10. **INVERSE FUNCTION RULE:** If f is a one-to-one differentiable function with inverse function $g(x) = f^{-1}(x)$, then the inverse function is differentiable and

$$g'(x) = \frac{d}{dx} f^{-1}(x) = \frac{1}{f'(f^{-1}(x))},$$

for any x such that $f'(f^{-1}(x)) \neq 0$.

Indeed, by definition of inverse function if $y = f^{-1}(x)$, then $x = f(y)$. Take this last relationship and differentiate both sides with respect to the independent variable x , bearing in mind that y is a function of x . So, we will need the chain rule. We obtain

$$\frac{dx}{dx} = f'(y) \frac{dy}{dx};$$

as $dx/dx = 1$, we can rearrange to obtain

$$\frac{dy}{dx} = \frac{d}{dx} f^{-1}(x) = \frac{1}{f'(y)} = \frac{1}{f'(f^{-1}(x))},$$

as required.

Example 29 Let $f(x) = x^3$. We know from the power rule that $f'(x) = 3x^2$. Consider now the inverse function i.e. $y = f^{-1}(x) = x^{1/3}$. The inverse function rule then says that

$$y' = \frac{d}{dx} f^{-1}(x) = \frac{1}{f'(f^{-1}(x))} = \frac{1}{3(x^{1/3})^2},$$

i.e.

$$y' = \frac{1}{3\sqrt[3]{x^2}},$$

which is the same result we would obtain via the power rule.

This example shows that we can obtain the same result using different rules/combination of rules. This gives us a powerful way to check if our calculations are correct.

All these results are summarized in Table 4. These rules need to be known.

Table 4: Summary: the rules of differentiation

$\mathbf{f(x)}$	$\mathbf{f'(x)}$
$f(x) = c$	$f'(x) = 0$
$f(x) = cg(x)$	$f'(x) = cg'(x)$
$f(x) = h(g(x))$	$f'(x) = h'(g(x))g'(x)$ [Chain Rule]
$f(x) = x^n$	$f'(x) = nx^{n-1}$
$f(x) = u(x)^n$	$f'(x) = nu(x)^{n-1}u'(x)$
$f(x) = u(x) \pm g(x)$	$f'(x) = u'(x) \pm g'(x)$
$f(x) = u(x)v(x)$	$f'(x) = u'(x)v(x) + u(x)v'(x)$ [Product Rule]
$f(x) = \frac{u(x)}{v(x)}$	$f'(x) = \frac{u'(x)v(x) - u(x)v'(x)}{v(x)^2}$
$f(x) = a^x$	$f'(x) = a^x \ln a$
$f(x) = a^{u(x)}$	$f'(x) = a^{u(x)}u'(x) \ln a$
$f(x) = e^x$	$f'(x) = e^x$
$f(x) = e^{u(x)}$	$f'(x) = e^{u(x)}u'(x)$
$f(x) = \log_a x$	$f'(x) = \frac{1}{x \ln a}$
$f(x) = \log_a u(x)$	$f'(x) = \frac{u'(x)}{u(x) \ln a}$
$f(x) = \ln x$	$f'(x) = \frac{1}{x}$
$f(x) = \ln u(x)$	$f'(x) = \frac{u'(x)}{u(x)}$
$g(x) = f^{-1}(x)$	$g'(x) = (f^{-1})'(x) = \frac{1}{f'(f^{-1}(x))}$

1.4.3 Higher order derivatives

The derivative f' of the function f is referred to as the first-order derivative of the function.

Higher order derivatives of the function f can be found by applying the rules of differentiation to lower-order derivatives.

Example 30 If $f(x) = 2x^4 + 5x^3 + 3x^2$, then

$$\begin{aligned} f'(x) &= 8x^3 + 15x^2 + 6x, \\ f''(x) &= 24x^2 + 30x + 6, \\ f'''(x) &= 48x + 30, \\ f^{(4)}(x) &= 48, \\ f^{(5)}(x) &= 0. \end{aligned}$$

All additional higher-order derivatives are also equal zero.

Of particular interest for applications in economics and finance is the second-order derivative, f'' , also denoted as

$$\frac{d^2y}{dx^2}, \quad \text{or} \quad \frac{d^2}{dx^2}f(x).$$

The second-order derivative measures the slope and the instantaneous rate of change of the first-order derivative with respect to a change in x , just as the first derivative measures the slope and the rate of change of the original (called primitive) function $f(x)$.

The second derivative corresponds to the curvature or concavity of the graph. The graph of a function with a positive second derivative is upwardly concave (i.e. convex), while the graph of a function with a negative second derivative curves in the opposite way (i.e. concave).

Example 31 1. Let us consider the parabola $f(x) = x^2$ from Example 16.1. We know by now that $f'(x) = 2x$, and therefore the function $f(x)$ is decreasing in the negative part of the real axis ($x < 0$), and increasing in the positive part of the real axis (i.e. $x > 0$). How about concavity/convexity? We can obtain the second-order derivative $f''(x) = 2 > 0$. Thus, the parabola is always convex.

2. Let us now consider the reciprocal function $f(x) = x^{-1}$. Then, $f'(x) = -x^{-2}$ and $f''(x) = 2x^{-3}$. Consequently, the function is always decreasing, it is convex for $x > 0$ and concave for $x < 0$. You can refer to the graph shown in Figure 8.

A note. The points at which the second order derivative of a function is zero are called inflection points.

A more detailed definition of concave/convex function A function whose second derivative is positive, is convex, meaning that the tangent lines will lie below the graph of the function. Similarly, a function whose second derivative is negative, is concave, and its tangent lines will lie above the graph of the function.

The formal definition of concave/convex function is the following. A real-valued function f on an interval is said to be concave if, for any x and y in the interval and for any $\alpha \in [0, 1]$

$$f((1 - \alpha)x + \alpha y) \geq (1 - \alpha)f(x) + \alpha f(y).$$

A function is called strictly concave if

$$f((1 - \alpha)x + \alpha y) > (1 - \alpha)f(x) + \alpha f(y),$$

for any $\alpha \in (0, 1)$ and $x \neq y$.

For a function f , this second definition states that for every z strictly between x and y , the point $(z, f(z))$ on the graph of f is above the straight line joining the points $(x, f(x))$ and $(y, f(y))$.

Differential If $y = f(x)$, for a differentiable function f , the differential dx is an independent variable; that is, dx can be given the value of any real number. The differential dy is then defined in terms of dx by the equation

$$dy = f'(x)dx,$$

which follows from equation (2). Therefore, dy is a dependent variable as it depends on the values of x and dx . If dx is given a specific value and x is taken to be some specific number in the domain of f , then the numerical value of dy is determined.

A note on notation: you can find differentials written in different ways, such as, for example

$$df(x) = f'(x)dx,$$

or

$$df(x) = f(dx).$$

1.5 Applications of derivatives

1.5.1 Linear approximation and Taylor expansion

In this section we explore a useful application of the derivatives of a function $f(x)$ aimed at obtaining convenient approximations.

The argument starts with the assumption that any continuous function $f(x)$, with derivatives of any order, can be approximated by a polynomial of the form

$$P(x) = a_0 + a_1(x - c) + a_2(x - c)^2 + a_3(x - c)^3 + \cdots + a_n(x - c)^n + \cdots,$$

with c a point in the domain of f suitably chosen.

We would like to fix the coefficients $a_0, a_1, a_2, a_3, \dots, a_n, \dots$ such that

$$P(x) = f(x).$$

If the above holds, also the derivatives of any order should be identical.

We proceed in steps as follows.

1. Set $x = c$ then $P(c) = a_0$ and for the identity to work, it follows that

$$a_0 = f(c).$$

2. Now, let us differentiate once so that

$$P'(x) = a_1 + 2a_2(x - c) + 3a_3(x - c)^2 + \cdots + na_n(x - c)^{n-1} + \cdots$$

Set $x = c$; for the identity to hold, i.e. $P'(c) = f'(c)$, it follows that

$$a_1 = f'(c).$$

3. Let us differentiate once more so that

$$P''(x) = 2a_2 + 3 \times 2a_3(x - c) + \cdots + n(n - 1)a_n(x - c)^{n-2} + \cdots$$

Set $x = c$; for the identity to hold, i.e. $P''(c) = f''(c)$, it follows that

$$a_2 = \frac{1}{2}f''(c).$$

4. We can repeat the procedure once more to obtain

$$P'''(x) = 3 \times 2a_3 + \dots + n(n-1)(n-2)a_n(x-c)^{n-3} + \dots$$

Set $x = c$; for the identity to hold, i.e. $P'''(c) = f'''(c)$, it follows that

$$a_3 = \frac{1}{3 \times 2} f'''(c) = \frac{1}{3!} f'''(c).$$

Thus, there is a pattern in the setting of the coefficients; indeed, in general we have

$$a_n = \frac{1}{n!} f^{(n)}(c).$$

Consequently, approximating the function $f(x)$ via the polynomial $P(x)$ leads to writing the function $f(x)$ as

$$\begin{aligned} f(x) &= f(c) + f'(c)(x-c) + \frac{1}{2!} f''(c)(x-c)^2 + \frac{1}{3!} f'''(c)(x-c)^3 + \dots + \frac{1}{n!} f^{(n)}(c)(x-c)^n + \dots \\ &= \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(c)(x-c)^n. \end{aligned}$$

This is called the Taylor expansion of the function f at c (or about c). If we decide to use only the first N terms, we will incur into an approximation error due to the truncation of the infinite summation at an arbitrarily chosen value of N . This error is called ‘remainder’ and is of higher order than N . In other words

$$f(x) = \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(c)(x-c)^n = \sum_{n=0}^N \frac{1}{n!} f^{(n)}(c)(x-c)^n + \underbrace{\mathcal{O}(N)}_{\text{Remainder}}.$$

A convenient point of choice is $c = 0$, so that the Taylor expansion reads

$$f(x) = \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(0)x^n.$$

Example 32 Consider the function $f(x) = e^x$. We know that all derivatives of the exponential function equal e^x . Consider the point $c = 0$; as $e^c = e^0 = 1$, it follows that the Taylor expansion at 0 of e^x is

$$e^x = \sum_{n=0}^{\infty} \frac{x^n}{n!}.$$

In Figure 18 we show in the top left panel the quality of the approximation when we use only the first two terms, i.e. $n = 0, 1$ so that

$$e^x \approx 1 + x.$$

In the remaining panels we use higher order terms, i.e. n up to 2, 3 and 4. We can observe the improvement in the quality of the approximation. The error (remainder) becomes smaller and smaller.

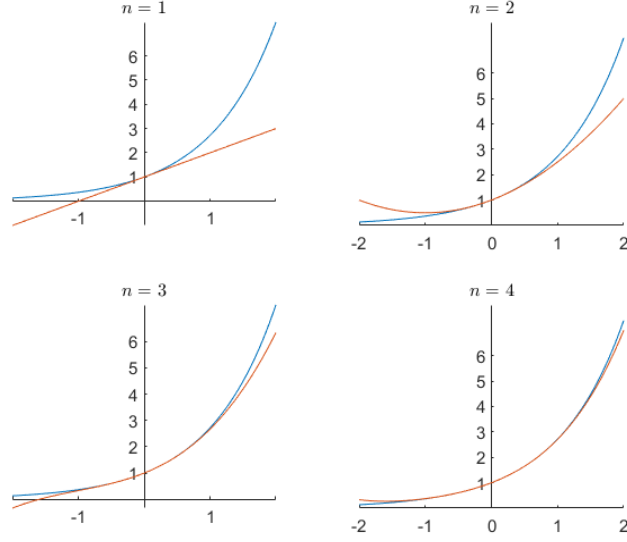


Figure 18: Taylor expansion of the exponential function $f(x) = e^x$ at zero. Several orders of approximation.

Remark 5 Using the Taylor expansion of the exponential function, we can show that

$$\sum_{n=0}^{\infty} e^{-\lambda} \frac{\lambda^n}{n!} = 1, \quad \lambda > 0.$$

The expression

$$e^{-\lambda} \frac{\lambda^n}{n!}$$

for $\lambda > 0$ is therefore a legitimate probability function. Indeed it defines the probability function of the Poisson distribution.

Example 33 (Derivatives, Bonds, Duration, Modified Duration and Convexity.) The price of a coupon bond, $P(r)$, with face value F , paying a fixed and known coupon d at the coupon dates $t = 1, 2, \dots, T$ and maturity T , with yield r is (based on the compound interest law)

$$P(r) = \sum_{t=1}^T \frac{d}{\left(1 + \frac{r}{m}\right)^{mt}} + \frac{F}{\left(1 + \frac{r}{m}\right)^{mT}},$$

with m the compounding frequency per year.

Important concepts when working with bonds are the one of duration and modified duration. We build these quantities as follows. Let us consider the first-order derivative of the bond price with respect to the yield r , that is

$$\frac{dP}{dr} = -\frac{1}{1 + \frac{r}{m}} \left(\sum_{t=1}^T t \frac{d}{\left(1 + \frac{r}{m}\right)^{mt}} + T \frac{F}{\left(1 + \frac{r}{m}\right)^{mT}} \right).$$

Thus, if D denotes the Macaulay's duration and MD the modified duration, then

$$\begin{aligned} D &= \frac{\sum_{t=1}^T t \frac{d}{(1+\frac{r}{m})^{mt}} + T \frac{F}{(1+\frac{r}{m})^{mT}}}{P(r)} = -\frac{(1+\frac{r}{m})}{P(r)} \frac{dP}{dr} \\ MD &= -\frac{1}{P(r)} \frac{dP}{dr}. \end{aligned}$$

In other words, the Macaulay duration D is the weighted average maturity of cash flows; the modified duration MD , instead, measures the bond sensitivity to changes in the interest rates. A further concept is the one of bond convexity. Bond convexity is a measure of the non-linear relationship of bond prices to changes in interest rates, i.e. the second derivative of the price of the bond with respect to interest rates.

Now consider the second order Taylor expansion around r of the bond price

$$P(r+h) = P(r) + P'(r)h + \frac{1}{2}P''(r)h^2 + \mathcal{O}(h^2).$$

Then

$$\begin{aligned} \frac{P(r+h) - P(r)}{P(r)} &= \frac{P'(r)}{P(r)}h + \frac{1}{2} \frac{P''(r)}{P(r)}h^2 + \mathcal{O}(h^2) \\ &= -MD h + \frac{1}{2}C h^2 + \mathcal{O}(h^2) \end{aligned}$$

with C denoting the convexity of the bond.

The Taylor expansion is very useful in Statistics also to prove the connection between Moment Generating Functions (and/or Characteristic Functions) and raw moments of a random variable, and between the log of the Moment Generating Function and the central moments of a random variable. These are fundamental quantities in Mathematical Finance and Computational Finance in order to study stochastic models for asset pricing and develop useful related numerical routines. We explore this in the following.

A digression in Statistics: Moments, Cumulants and the Generating Functions The Moment Generating Function (MGF) of a random variable X with cumulative distribution function $F_X(\cdot)$ is defined as follows.

$$M_X(k) = \mathbb{E} \left(e^{kX} \right) = \int_{-\infty}^{\infty} e^{kx} dF_X(x).$$

As the name of this function suggests, it can be used to study the distribution features of random variables as it helps to quickly recover the relevant moments. In more details, we distinguish between.

- The raw moments of the random variable X , which are defined as

$$\mu_n = \mathbb{E} (X^n),$$

and which can be recovered by direct differentiation of the MGF as

$$\mu_n = \left. \frac{\partial^n M_X(k)}{\partial k^n} \right|_{k=0}.$$

Proof. Let us use the Taylor expansion of the exponential function to rewrite the MGF as

$$\begin{aligned} M_X(k) &= \mathbb{E} \left(\sum_{n=0}^{\infty} \frac{(kX)^n}{n!} \right) \\ &= \sum_{n=0}^{\infty} \frac{k^n}{n!} \mathbb{E}(X^n) \\ &= 1 + k\mathbb{E}(X) + \frac{1}{2}k^2\mathbb{E}(X^2) + \dots \end{aligned}$$

First order differentiation of the above with respect to k returns

$$\frac{\partial M_X(k)}{\partial k} = \mathbb{E}(X) + k\mathbb{E}(X^2) + \dots;$$

second order differentiation of the above with respect to k returns

$$\frac{\partial^2 M_X(k)}{\partial k^2} = \mathbb{E}(X^2) + k\mathbb{E}(X^3) + \dots,$$

and so on. The required results follows setting $k = 0$ in the above. ■

- The central moments of the random variable X , defined as

$$m_n = \mathbb{E}[(X - \mathbb{E}(X))^n] \quad n \geq 2.$$

The central moments can be easily obtained from the cumulants which are defined as

$$c_n = \left. \frac{\partial^n \ln M_X(k)}{\partial k^n} \right|_{k=0}.$$

Proof. First order differentiation of $\ln M_X(k)$ returns

$$\frac{\partial \ln M_X(k)}{\partial k} = \frac{1}{M_X(k)} \frac{\partial M_X(k)}{\partial k};$$

setting $k = 0$ returns the first order cumulant which coincides with the mean. Second order differentiation returns

$$\frac{\partial^2 \ln M_X(k)}{\partial k^2} = -M_X(k)^{-2} \left(\frac{\partial M_X(k)}{\partial k} \right)^2 + \frac{1}{M_X(k)} \frac{\partial^2 M_X(k)}{\partial k^2};$$

setting $k = 0$ returns $-\mathbb{E}(X)^2 + \mathbb{E}(X^2)$, i.e. the variance. The higher order cumulants can be obtained following the same procedure. ■

In particular, the following identities hold.

$$\begin{aligned} c_1 &= \mathbb{E}(X) \\ c_2 &= \mathbb{E}[(X - \mathbb{E}(X))^2] = \text{Var}(X) = m_2 \\ c_3 &= \mathbb{E}[(X - \mathbb{E}(X))^3] = m_3 \\ c_4 &= \mathbb{E}[(X - \mathbb{E}(X))^4] - 3\text{Var}^2(X) = m_4 - 3m_2^2. \end{aligned}$$

The skewness and the kurtosis of the distribution of a random variable are measured respectively by the index of skewness

$$\gamma_1 = \frac{c_3}{c_2^{3/2}}$$

and by the index of excess kurtosis

$$\gamma_2 = \frac{c_4}{c_2^2}.$$

The distribution of a random variable is said to be asymmetric to the left if $\gamma_1 < 0$, symmetric if $\gamma_1 = 0$, and asymmetric to the right otherwise. It is platikurtic if $\gamma_2 < 0$, mesokurtic if $\gamma_2 = 0$ and leptokurtic otherwise.

As the following example shows, the Gaussian distribution is both symmetric and mesokurtic. If any other random variable is, for example, leptokurtic, it means that the tails of its distribution are thicker than the ones of the Gaussian random variable. Therefore, ‘extreme’ events are more likely than in the Gaussian case.

Example 34 Few (and very important, as we will use them throughout the entire year) examples of random variables.

1. As already seen before, the normal (or Gaussian) random variable $X \sim N(\mu, \sigma^2)$ is a continuous random variable defined by the density function

$$f_X(x) = \frac{e^{-\frac{(x-\mu)^2}{2\sigma^2}}}{\sigma\sqrt{2\pi}}.$$

You can easily show that

$$\mathbb{E}(X) = \mu; \mathbb{V}ar(X) = \sigma^2; c_3 = 0 = c_4; \quad M_X(k) = e^{k\mu + \frac{k^2\sigma^2}{2}}$$

2. The Poisson random variable is an example of discrete random variable. More precisely, a Poisson random variable $N \sim Poi(\lambda)$, with rate λ has probability mass

$$p_N(n) = \frac{e^{-\lambda}\lambda^n}{n!}$$

from which it follows that

$$\mathbb{E}(N) = \lambda = \mathbb{V}ar(N) = c_3 = c_4; \quad M_N(k) = e^{\lambda(e^k - 1)}.$$

Hence the Poisson distribution is asymmetric to the right and leptokurtic.

3. Assume $X \sim \Gamma(\alpha, \lambda)$, $\alpha > 0$. Then X is a non-negative random variable which follows a gamma distribution; its density function is given by

$$f(x) = \frac{1}{\Gamma(\alpha)} \lambda^\alpha x^{\alpha-1} e^{-\lambda x},$$

where $\Gamma(\alpha)$ is the gamma function, which is defined as

$$\Gamma(\alpha) = \int_0^\infty x^{\alpha-1} e^{-x} dx,$$

and has the property that

$$\Gamma(\alpha) = \Gamma(\alpha - 1)(\alpha - 1).$$

This means that

$$\Gamma(\alpha) = (\alpha - 1)!$$

where α is a positive integer. The MGF of X is

$$M_X(k) = \frac{1}{\Gamma(\alpha)} \lambda^\alpha \int_0^\infty x^{\alpha-1} e^{-x(\lambda-k)} dx.$$

Set $y = x(\lambda - k)$, then

$$M_X(k) = \frac{1}{\Gamma(\alpha)} \lambda^\alpha \int_0^\infty \left(\frac{y}{\lambda - k} \right)^{\alpha-1} \frac{e^{-y}}{\lambda - k} dy = \left(\frac{\lambda}{\lambda - k} \right)^\alpha.$$

Note that if $\alpha = 1$, then X follows an exponential distribution with rate λ . Using the MGF you can show that the gamma random variable has mean $\mu = \alpha/\lambda$ and variance $\nu = \alpha/\lambda^2$. The parameter α is the shape parameter, whilst λ is the scale parameter. You can also show that the gamma random variable is asymmetric to the right and leptokurtic.

1.5.2 Computing limits: L'Hôpital

Suppose $f(x)$ and $g(x)$ are differentiable and $g'(x) \neq 0$ near a (except possibly at a). Suppose that

$$\lim_{x \rightarrow a} f(x) = 0 \text{ and } \lim_{x \rightarrow a} g(x) = 0$$

or that

$$\lim_{x \rightarrow a} f(x) = \pm\infty \text{ and } \lim_{x \rightarrow a} g(x) = \pm\infty.$$

Consider the limit

$$\lim_{x \rightarrow a} \frac{f(x)}{g(x)};$$

under the given assumption, this originates an indeterminate form of the type $0/0$ or ∞/∞ . Then

$$\lim_{x \rightarrow a} \frac{f(x)}{g(x)} = \lim_{x \rightarrow a} \frac{f'(x)}{g'(x)}$$

if the limit on the right side exists (or is $\pm\infty$).

Example 35 Consider the notable limit

$$\lim_{x \rightarrow 0} \frac{\ln(1+x)}{x} = 1$$

(see Table 3). We can verify this is indeed the case. Let us differentiate

$$\lim_{x \rightarrow 0} \frac{\ln(1+x)}{x} = \lim_{x \rightarrow 0} \frac{\frac{1}{1+x}}{1} = 1$$

where the final step follows by substitution.

1.5.3 Implicit differentiation

So far we have dealt with explicit functions, i.e. functions for which the dependent variable is expressed in terms of the independent variable explicitly, for example

$$y = \sqrt{x^3 + 1}.$$

Occasionally, though, such explicit expressions are not possible and implicit functions are encountered. These are functions which are defined implicitly by a relation between x and y such as

$$x^3 + y^3 = 6xy.$$

Some implicit functions can be easily converted to explicit functions by solving for the dependent variable in terms of the independent variable; for others this is not possible, or at least not convenient. However, in order to find the derivative, we do not need to solve an equation for y in terms of x . Instead we can use the method of implicit differentiation.

This consists of differentiating both sides of the equation with respect to x and then solving the resulting equation for y' . Note that differentiating with respect to x amounts to applying the chain rule as y is a function of x , albeit implicit. We show the procedure in a few examples.

Example 36 1. Consider $3y^5 - 6y^4 + 5x^6 = 243$. By considering y as $y(x)$, differentiate both sides, so that

$$15y^4y' - 24y^3y' + 30x^5 = 0,$$

from which it follows that

$$y' = \frac{-10x^5}{5y^4 - 8y^3}.$$

2. Consider our previous example $x^3 + y^3 = 6xy$. Then, by differentiating both sides we obtain

$$x^2 + y^2y' = 2y + 2xy'$$

and therefore

$$y' = \frac{2y - x^2}{y^2 - 2x}.$$

With this knowledge we can now for example obtain the tangent line in $(3, 3)$. The slope is the derivative in the given point i.e. set $x = y = 3$, therefore

$$y' = -1.$$

The tangent line then has equation

$$y - 3 = -(x - 3)$$

i.e.

$$y = 6 - x.$$

The implicit function and the tangent line are represented in Figure 19.

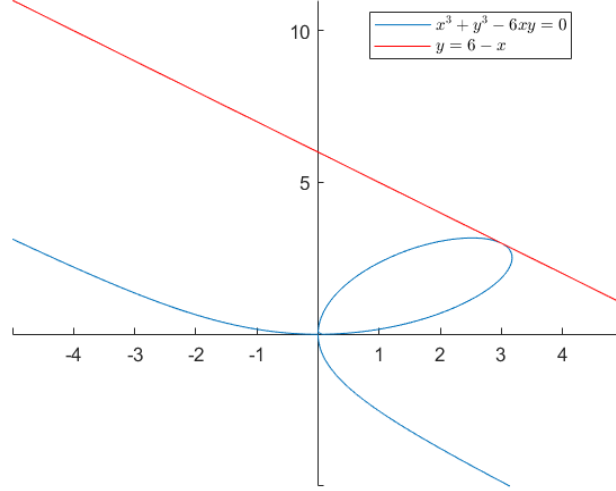


Figure 19: Descartes folium $x^3 + y^3 - 6xy = 0$ and its tangent line in $(3, 3)$.

1.5.4 Minima & Maxima of a function

In Section 1.4 we have seen that the sign of the first and the second order derivatives are providing useful information concerning the shape of the graph of a function, i.e. in which part of the domain it increases/decreases or it is convex/concave.

This information is important in order to be able to detect minima and maxima (if they exist) of the function, a concept strictly linked to the idea of optimization which lies at the core of many applications in Economics and Finance.

We start our analysis by providing the formal definitions of minima and maxima and the different type of minima and maxima which we can encounter.

Definition 7 Consider a continuous and differentiable function f with domain X . Let c be a number in the domain X of f . Then.

1. c is a point of global maximum if $f(c) \geq f(x)$ for all $x \in X$ (c also known as point of absolute maximum).
2. c is a point of global minimum if $f(c) \leq f(x)$ for all $x \in X$ (c also known as point of absolute minimum).
3. c is a point of local maximum if $f(c) \geq f(x)$ for all x near c .
4. c is a point of local minimum if $f(c) \leq f(x)$ for all x near c .
5. If f is continuous on a closed interval $[a, b]$, then f attains an absolute maximum value $f(c)$ and an absolute minimum value $f(d)$ at some points c, d in $[a, b]$.

Example 37 Let us revisit some functions which we have already studied in Section 1.4.

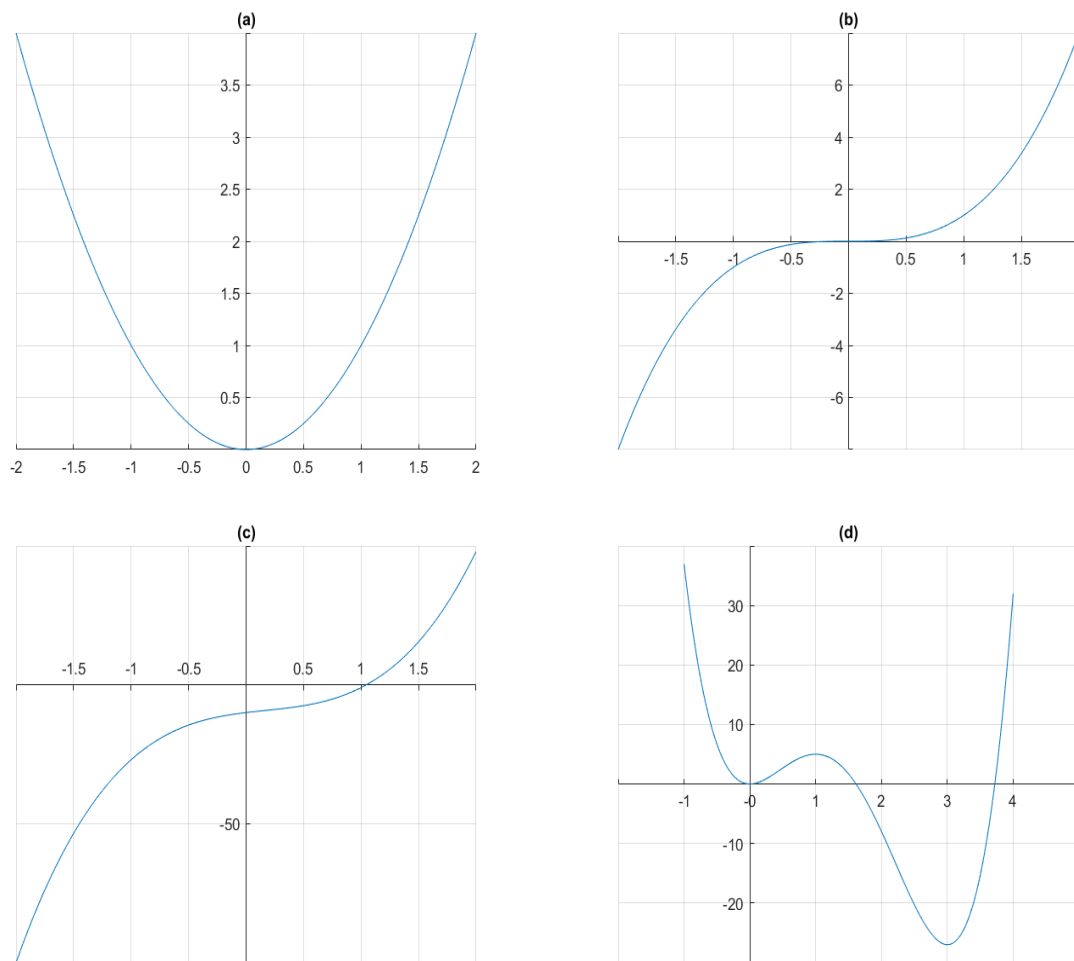


Figure 20: Global and local minima and maxima. Panel (a): x^2 . Panel (b): x^3 , Panel (c): $8x^3 - 4x^2 + 5x - 10$. Panel (d): $3x^4 - 16x^3 + 18x^2$ on $[-1, 4]$.

1. Consider $f(x) = x^2$. We have previously established that this function is decreasing for $x < 0$, increasing for $x > 0$, it has a stationary point in $x = 0$ and it is convex. Its graph - reported for convenience in Figure 20.(a) - shows that this function has a point of global minimum for $x = 0$ and no maxima. Indeed, $f(0) = 0 \leq f(x)$ for all $x \in \mathbb{R}$, which is the domain of the parabola.
2. Consider $f(x) = x^3$ with domain $\mathbb{R}$. Then, $f'(x) = 3x^2$ and $f''(x) = 6x$. Therefore, f' is always positive and the function is always increasing on $\mathbb{R}$. $f'(x) = 0$ for $x = 0$ which is then a stationary point. The function is also concave for $x < 0$ and convex for $x > 0$, and $x = 0$ is the inflection point. Its graph in Figure 20.(b) shows that the function has no minima or maxima. The stationary point $x = 0$ is therefore called a saddle point.
3. Consider the function $f(x) = 8x^3 - 4x^2 + 5x - 10$ from Section ??, Exercise 6. We have seen that the function has no stationary points, and it is always increasing. The graph of the function shown in Figure 20.(c) shows that the function has no minima or maxima.

4. Consider the function $f(x) = 3x^4 - 16x^3 + 18x^2$ on the interval $[-1, 4]$. We start with the first order derivative

$$f'(x) = 12x^3 - 48x^2 + 36x = 12x(x^2 - 4x + 3).$$

To find the stationary points, let us find the zeros of $f'(x)$, i.e. let us solve

$$12x(x^2 - 4x + 3) = 0,$$

which gives as roots $x_1 = 0$, and $x_2 = 1$, $x_3 = 3$. These are the stationary points. Now. We want to know in which regions of the domain $[-1, 4]$ $f'(x) > 0$. Thus, we need either $x > 0$ and at the same time $x^2 - 4x + 3 > 0$, or $x < 0$ and at the same time $x^2 - 4x + 3 < 0$. We note that the inequality

$$x^2 - 4x + 3 > 0$$

is satisfied for $x < 1$ or $x > 3$. We can put all the information together in a scheme as follows.

	$x < 0$	$0 < x < 1$	$1 < x < 3$	$x > 3$
x	—	+	+	+
$x^2 - 4x + 3$	+	+	—	+
$f'(x)$	—	+	—	+
$f(x)$	$\searrow$	$\nearrow$	$\searrow$	$\nearrow$

Then

$$f'(x) < 0 \text{ for } \{-1 < x < 0\} \cup \{1 < x < 3\}$$

and

$$f'(x) > 0 \text{ for } \{0 < x < 1\} \cup \{3 < x < 4\}.$$

Consequently, $f(x)$ is decreasing on $\{-1 < x < 0\} \cup \{1 < x < 3\}$ and increasing on $\{0 < x < 1\} \cup \{3 < x < 4\}$. This suggests that $f(x)$ has minima for $x = 0$ and $x = 3$ and a maximum for $x = 1$.

Are these local or global extreme points? We need to check the value of the primitive function: $f(0) = 0$, $f(1) = 5$ and $f(3) = -27$. We also note that at the limits of the given domain we have $f(-1) = 37$ and $f(4) = 32$. Therefore, the function $f(x)$ has a local minimum in $(0, 0)$, a global minimum in $(3, -27)$, a local maximum in $(1, 5)$ and a global maximum in $(-1, 37)$. In $(4, 32)$, the function has neither a local nor an absolute maximum.

Now, let us turn our attention to the concavity/convexity of f . The second order derivative is

$$f''(x) = 36x^2 - 96x + 36.$$

The zeros are $x_1 = (4 - \sqrt{7})/3 \cong 0.45$ and $x_2 = (4 + \sqrt{7})/3 \cong 2.21$. Consequently, as $f''(x)$ is a parabola with positive leading coefficient, then

$$f''(x) < 0 \text{ for } \left\{ \frac{4 - \sqrt{7}}{3} < x < \frac{4 + \sqrt{7}}{3} \right\}$$

and

$$f'(x) > 0 \text{ for } \left\{ -1 < x < \frac{4 - \sqrt{7}}{3} \right\} \cup \left\{ \frac{4 + \sqrt{7}}{3} < x < 4 \right\}.$$

We note in particular the following.

$$\begin{aligned} f''(-1) &= 168 > 0 \\ f''(0) &= 36 > 0 \\ f''(1) &= -24 < 0 \\ f''(3) &= 72 > 0. \end{aligned}$$

Thus, the function f is concave at the maxima and convex at the minima. The graph is shown in Figure 20.(d).

The latter example illustrates how we can use the derivatives of first and second order to identify the minima and maxima of any function. We can therefore formulate two tests, one based exclusively on the first-order derivative, and one which uses both the first and second-order derivatives. For this purpose we require the following result (stated without proof).

Theorem 8 (Mean Value Theorem) *Let f be a function that satisfies the following hypotheses.*

- (a) f is continuous on the closed interval $[x_1, x_2]$.
- (b) f is differentiable on the open interval (x_1, x_2) .

Then there is a number c in $[x_1, x_2]$ such that

$$f'(c) = \frac{f(x_2) - f(x_1)}{x_2 - x_1}.$$

First-order derivative test. Let c be a stationary point of the function f , i.e. a point in the domain of f such that $f'(c) = 0$. Then the following situations might occur.

- (a) c is such that $f'(x) < 0$ for any $x < c$ and $f'(x) > 0$ for any $x > c$. Then c is a minimum of the function $f(x)$.

Proof. Let us consider first the case in which $x < c$. Then $f'(x) < 0$ by assumption. The Mean Value Theorem 8 implies that on the interval $[x, c]$

$$\frac{f(c) - f(x)}{c - x} < 0,$$

which implies that $f(c) < f(x)$. Let us now consider the case in which $x > c$ so that $f'(x) > 0$. The Mean Value Theorem 8 implies that on the interval $[c, x]$

$$\frac{f(x) - f(c)}{x - c} > 0,$$

i.e. $f(x) > f(c)$. Thus, $f(c) < f(x)$ for all x near c (c excluded), and $f(c) \leq f(x)$ for all x near c with c included. Then by Definition 7 c is a (local) minimum. ■

- (b) c is such that $f'(x) > 0$ for any $x < c$ and $f'(x) < 0$ for any $x > c$. Then c is a maximum of the function $f(x)$.

Proof. It follows from the same argument as above, mutatis mutandis. ■

Example 38 Consider the functions which we studied in Example 37. The function $f(x) = x^2$ has a minimum in 0. Indeed $f'(x) = 2x$ which is negative for $x < 0$ and positive for $x > 0$, and therefore it satisfies the first order test (condition (a)).

Let us consider now $f(x) = 3x^4 - 16x^3 + 18x^2$. We found a minimum in 0: indeed for $x < 0$ $f'(x) < 0$ whilst $f'(x) > 0$ for $x > 0$ up to 1. After 1, $f'(x) < 0$ and by the first order test (condition (b)) f has a maximum in 1. Note that after 3, $f'(x) > 0$ which gives us a minimum in 3. Which is what we also obtained in Example 37.

Example 39 A firm short-run production function is given by

$$Q = 6L^2 - 0.2L^3$$

where L denotes the number of workers.

- (a) Find the size of the workforce that maximises output.
 (b) Find the size of the workforce that maximises the average product of labour AP_L . Calculate the marginal product of labor MP_L and AP_L at this value of L .

We start from part (a). It means we are looking for the value of L such that the production function hits its (global) maximum. This means we need to find the first order derivative, and study its sign around any possible stationary point. So let us start.

First order derivative

$$Q' = 12L - 0.6L^2 = 0.6L(20 - L).$$

Therefore the zeros are $L_1 = 0$ and $L_2 = 20$. As the leading coefficient of the polynomial of second order is negative, then $Q' > 0$ if $0 < L < 20$. Our magic scheme looks like

	$L < 0$	$0 < L < 20$	$L > 20$
L	−	+	+
$20 - L$	+	+	−
$Q'(L)$	−	+	−
$Q(x)$	↘	↗	↘

Thus, the output is maximized for $L = 20$ and $Q(20) = 800$.

Let us continue with part (b). We need to obtain the function AP_L and find its global maximum.

Thus,

$$AP_L = \frac{Q}{L} = 6L - 0.2L^2;$$

the first order derivative is

$$AP'_L = 6 - 0.4L,$$

which has stationary point in $L = 6/0.4 = 15$ and AP'_L is positive for $L < 15$ and negative otherwise. Thus the global maximum is $L = 15$. The corresponding value of AP_L is $AP_L(15) = 45$. The function MP_L is actually $Q' = 0.6L(20 - L)$ and for $L = 15$ we obtain $MP_L = Q'(15) = 45$.

Second-order derivative test Let c be a stationary point of the function f , i.e. a point in the domain of f such that $f'(c) = 0$. Then the following situations might occur.

- (a) $f'(c) = 0$ and $f''(c) > 0$, then c is a (local) minimum.

Proof. As $f''(c) > 0$, then

$$0 < f''(c) = \lim_{h \rightarrow 0} \frac{f'(c+h) - f'(c)}{h} = \lim_{h \rightarrow 0} \frac{f'(c+h)}{h}$$

which follows as $f'(c) = 0$. Therefore for sufficiently small h

$$\frac{f'(c+h)}{h} > 0.$$

The above implies that $f'(c+h) < 0$ if $h < 0$, i.e. f is decreasing as it approaches c from the left; similarly $f'(c+h) > 0$ if $h > 0$, i.e. f is increasing on the right side of c . By the first-order derivative test, f has a minimum in c . ■

- (b) $f'(c) = 0$ and $f''(c) < 0$, then c is a (local) maximum.

Proof. It follows from the same argument as above, *mutatis mutandis*. ■

Example 40 We can revisit Example 37 with the second order test. For $f(x) = x^2$ we have that $f''(x) = 2 > 0$. Thus, the function is convex and the stationary point $x = 0$ is indeed a minimum.

For the case of the function $f(x) = 3x^4 - 16x^3 + 18x^2$, we have already computed the values of the second order derivatives in the stationary points - see Exercise 37. These results are in line with the second order test.

Example 41 Consider the maximization problem in Example 39. We can verify the sign of the second order derivative in the stationary points we identified as maxima. For part (a), we have

$$Q'' = 0.6(20 - 2L),$$

therefore $Q''(20) = -12 < 0$. So $L = 20$ is indeed the level of workforce which allows the maximization of the output. For part (b), we have

$$AP''_L = -0.4 < 0,$$

in other words the average product of labour is a parabola with downwards concavity (the leading coefficient is negative) and can therefore only admit one global maximum.

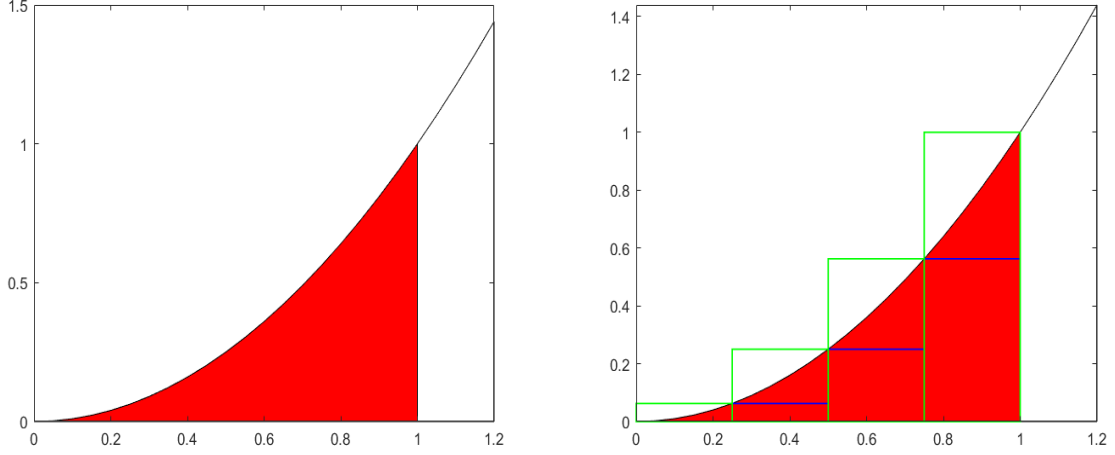


Figure 21: Area between 0 and 1 underneath x^2 : an integration problem.

In the previous part, we considered local minima and maxima. As the examples show, any given function can present more than one local minima/maxima. In order to determine which one is the actual global minimum/maximum, we need to check the values of f at such candidate points in order to select the actual global minimum and maximum.

A note. When working with functions of 1 variable, i.e. functions of the form $y = f(x)$ such as the ones considered so far, the first-order derivative test is usually sufficient in order to determine minima and maxima of a function f . However, in more general cases when we consider functions of more than one variable, i.e. of the form $z = f(x, y)$ for example, the first-order test becomes quite difficult to apply, whilst the second-order test is of more immediate application.

1.6 Integrals

Consider the graph of the function $f(x) = x^2$ in Figure 21, left hand panel. We are focussing specifically on the shaded portion of the figure, i.e. all those points which lie below the curve $f(x)$ and the interval $[0, 1]$ on the horizontal axis. We want to measure this area. We can proceed as follows.

- Divide the range of values of x in the interval of interest $[a, b]$ into n subinterval, each of size $\Delta x = (b - a)/n$. In the case illustrated in Figure 21, right hand panel, we have fixed $n = 4$, so that $\Delta x = 0.25$ as $[a, b] = [0, 1]$.
- The start-points of each interval are $\{0, 0.25, 0.5, 0.75\}$ and the corresponding values of the function $f(x) = x^2$ are $\{0, 0.0625, 0.25, 0.5625\}$, whilst the end-points of each interval are $\{0.25, 0.5, 0.75, 1\}$ and the corresponding values of the function $f(x) = x^2$ are $\{0.0625, 0.25, 0.5625, 1\}$.
- Calculate the area of each rectangle formed by using the start-point of each interval and its corresponding value of the function, and sum them. As all these rectangles have the same base Δx we compute

$$S_4^l = 0.25 (0 + 0.0625 + 0.25 + 0.5625) = 0.21875,$$

which corresponds to the sum of the rectangles with blue perimeter. We can write the above in general terms as

$$S_n^l = \sum_{i=1}^n f(x_{i-1})\Delta x,$$

and we call this a lower sum.

Alternatively, calculate the area of each rectangle formed by using the end-point of each interval and its corresponding value of the function, and sum them, i.e.

$$S_4^u = 0.25 (0.0625 + 0.25 + 0.5625 + 1) = 0.46875,$$

which corresponds to the sum of the rectangles with green perimeter. We can write the above in general terms as

$$S_n^u = \sum_{i=1}^n f(x_i)\Delta x,$$

and we call this an upper sum. In the case illustrated in Figure 21, right hand panel, this sum would be

As Figure 21 itself shows, both summations are quite a ‘crude’ approximation, and our scope is to find a way of improving it. The idea is to increase the number of subinterval n (from 4 to 8, say, then to 16, 32, 64 and so on) so that the sum of the resulting rectangles becomes more and more similar to the shaded area which we want to measure. If this is the case, then both lower sum and upper sum will return an identical value.

In other words, what we are actually computing is a limit for $n \rightarrow \infty$, or equivalently $\Delta x \rightarrow 0$, of S_n^l and S_n^u ; thus, the area A we want to measure is given by

$$A = \lim_{n \rightarrow \infty} \sum_{i=1}^n f(x_{i-1})\Delta x = \lim_{n \rightarrow \infty} \sum_{i=1}^n f(x_i)\Delta x.$$

The summations

$$\sum_{i=1}^n f(x_{i-1})\Delta x, \quad \sum_{i=1}^n f(x_i)\Delta x$$

are usually referred to as the Riemann sums. The limit of the Riemann sums is known as the integral of $f(x)$ and it is symbolically written as

$$A = \int_a^b f(x)dx.$$

The scope of this section is to formalize the intuition given above in a way that is practical for actual applications: the ‘slicing’ procedure shown above requires many slices, many rectangles, many sums,... too many calculations!

1.6.1 The Definite Integral

Consider a function f on an interval $[a, b]$. Consider a partition, i.e. a subdivision, of this interval in n subintervals identified by the points $a = x_0 < x_1 < \dots < x_i < \dots < x_n = b$.

Example 42 Referring to the previous illustration with $f(x) = x^2$ on $[0, 1]$, we consider $n = 4$ subintervals $[0, 0.25]$, $[0.25, 0.5]$, $[0.5, 0.75]$ and $[0.75, 1]$, thus the points are $0, 0.25, 0.5, 0.75, 1$.

Let us denote by

$$m_k = \min_{x_{i-1} \leq x_k \leq x_i} f(x_k)$$

i.e. the minimum value taken by f over each subinterval, and

$$M_k = \max_{x_{i-1} \leq x_k \leq x_i} f(x_k)$$

i.e. the maximum value taken by f over each subinterval. Then, the following holds.

Definition 9 (Definite Integral) *If f is a function defined on $[a, b]$, the definite integral of f from a to b is the number*

$$\int_a^b f(x)dx = \lim_{n \rightarrow \infty} \sum_{k=1}^n m_k \Delta x = \lim_{n \rightarrow \infty} \sum_{k=1}^n M_k \Delta x$$

provided that the limits exist and coincide. If they do exist and coincide, we say that f is (Riemann) integrable on $[a, b]$.

$\int$ is the integral sign, $f(x)$ is called the integrand function, $[a, b]$ are the limits of integration, with a being the lower limit and b the upper limit; for now the quantity dx indicates that x is the independent variable and that we ‘slice’ the x range in order to build our rectangles.

Example 43 Referring to the previous illustration with $f(x) = x^2$ on $[0, 1]$, m_k coincides with the values of the functions at the start-point of each subinterval, whilst M_k coincides with the values of the functions at the end-point of each subinterval. Consequently

$$\sum_{k=1}^n m_k \Delta x = \sum_{i=1}^n f(x_{i-1}) \Delta x = S_n^l,$$

i.e. it is the lower (Riemann) sum, and

$$\sum_{k=1}^n M_k \Delta x = \sum_{i=1}^n f(x_i) \Delta x = S_n^u,$$

i.e. it is the upper (Riemann) sum.

The following properties are very useful when it comes to calculate integrals.

Properties of the definite integral

1. For any constant c

$$\int_a^b c dx = c(b - a);$$

2. For any constant c

$$\int_a^b c f(x) dx = c \int_a^b f(x) dx;$$

- 3.

$$\int_a^b (f(x) + g(x)) dx = \int_a^b f(x) dx + \int_a^b g(x) dx;$$

4.

$$\int_a^b (f(x) - g(x))dx = \int_a^b f(x)dx - \int_a^b g(x)dx;$$

5.

$$\int_a^c f(x)dx + \int_c^b f(x)dx = \int_a^b f(x)dx; \quad a \leq c \leq b$$

6.

$$\int_a^b (-f(x))dx = \int_b^a f(x)dx.$$

7. If $f(x) \geq 0$ for $a \leq x \leq b$, then

$$\int_a^b f(x)dx \geq 0.$$

8. If $f(x) \geq g(x)$ for $a \leq x \leq b$, then

$$\int_a^b f(x)dx \geq \int_a^b g(x)dx.$$

9. If m and M are constants such that $m \leq f(x) \leq M$ for $a \leq x \leq b$, then

$$m(b-a) \leq \int_a^b f(x)dx \leq M(b-a).$$

We have not yet answered the question how do we actually calculate integrals. This is covered in the next section.

1.6.2 How to evaluate integrals: the Indefinite Integral

Evaluation of integrals is given in the Fundamental Theorem of Calculus. As the name suggest, this is a very important result. So, first we look at it and its proof. Then we try to understand what it is really saying, i.e. its meaning and its implication, in order to understand what we are really looking for.

Before all of this, though, we need to introduce the concept of indefinite integral. It will help us to simplify a little the notation.

Definition 10 (Indefinite Integral) *If f is a continuous function, the indefinite integral of f is the family of functions*

$$\int f(x)dx = F(x) + C$$

for any constant of integration C , and F is any antiderivative (or primitive function) of f such that $F' = f$. Further, the definite integral is the number

$$\int_a^b f(x)dx = \left. \int f(x)dx \right|_a^b.$$

Now we are ready.

Theorem 11 (Fundamental Theorem of Calculus) *Suppose f is continuous on $[a, b]$.*

(a) If

$$F(x) = \int_a^x f(y)dy,$$

then

$$F'(x) = f(x).$$

(b)

$$\int_a^b f(x)dx = F(x)|_a^b = F(b) - F(a).$$

Proof. Let us start with part (a). Consider a quantity $h > 0$ (you can repeat the whole argument for $h < 0$ the result will be the same). We have assumed that $F(x) = \int_a^x f(y)dy$, consequently

$$F(x+h) = \int_a^{x+h} f(y)dy;$$

by the properties of the integral we saw above, we have

$$F(x+h) - F(x) = \int_x^{x+h} f(y)dy.$$

Divide by h both sides to obtain

$$\frac{F(x+h) - F(x)}{h} = \frac{\int_x^{x+h} f(y)dy}{h}. \quad (4)$$

Since f is continuous on $[x, x+h]$, there are numbers u and v in $[x, x+h]$ such that $f(u) = m$ and $f(v) = M$, where m and M are the absolute minimum and maximum values of f on $[x, x+h]$. By property 9 of the definite integral

$$mh \leq \int_x^{x+h} f(y)dy \leq Mh,$$

i.e.

$$f(u)h \leq \int_x^{x+h} f(y)dy \leq f(v)h.$$

As by assumption $h > 0$, we can divide all sides by h and use equation (4) to write

$$f(u) \leq \frac{F(x+h) - F(x)}{h} \leq f(v).$$

Now, take the limit for $h \rightarrow 0$. Notice that $\lim_{h \rightarrow 0} x+h = x$ and therefore $\lim_{h \rightarrow 0} u = x$ and $\lim_{h \rightarrow 0} v = x$, so that $\lim_{h \rightarrow 0} f(u) = \lim_{h \rightarrow 0} f(v) = f(x)$. Consequently, the squeeze theorem for limits implies together with the definition of derivative that

$$F'(x) = \lim_{h \rightarrow 0} \frac{F(x+h) - F(x)}{h} = f(x).$$

We now move to part (b) of the Theorem. We follow the ‘partition’ procedure which we saw at the beginning of this Section. Thus, subdivide $[a, b]$ in n equidistant points $a = x_0, x_1, \dots, x_n = b$ such that $x_i - x_{i-1} = \Delta x$ for $i = 1, \dots, n$, with

$$\Delta x = \frac{b-a}{n}.$$

It follows that

$$F(b) - F(a) = F(x_n) - F(x_0).$$

Then, add and subtract the values of the function F at each points in the interval, i.e.

$$F(b) - F(a) = F(x_n) - F(x_0) = F(x_n) - F(x_{n-1}) + F(x_{n-1}) - \cdots - F(x_1) + F(x_1) - F(x_0).$$

This can be written in a compact manner as

$$F(b) - F(a) = \sum_{i=1}^n (F(x_i) - F(x_{i-1})).$$

We now remind ourselves that from the Mean Value Theorem 8 there exists a point $x_i^* \in [x_{i-1}, x_i]$ such that

$$F'(x_i^*) = \frac{F(x_i) - F(x_{i-1})}{\Delta x}.$$

This can be rewritten as

$$F(x_i) - F(x_{i-1}) = f(x_i^*)\Delta x$$

as by definition $F' = f$. Therefore

$$F(b) - F(a) = \sum_{i=1}^n (F(x_i) - F(x_{i-1})) = \sum_{i=1}^n f(x_i^*)\Delta x.$$

Now we take the limit of each side of this equation for $n \rightarrow \infty$. The left side is a constant and the right side is a Riemann sum for the function f , so

$$F(b) - F(a) = \lim_{n \rightarrow \infty} \sum_{i=1}^n f(x_i^*)\Delta x = \int_a^b f(x)dx,$$

where the last step follows from the definition of definite integral (see Definition 9). ■

These results are telling us what an integral is and what the operation of integration is: the integral is the inverse function of a derivative, and integration is the inverse operation of differentiation.

This means that when we calculate or solve an integral, we are really looking for the function which, once it is differentiated, gives us the integrand function.

This is very helpful information: we know how to differentiate a large family of functions. Therefore, we should be able to trace back the primitive function which originated a given derivative. This will help us to solve the indefinite integral.

Once we have the primitive function, i.e. we know the solution to the indefinite integral, the definite integral can be easily calculated according to Theorem 11.(b).

Example 44 Let us consider the example we started with in Figure 21. So: $f(x) = x^2$ on $[0, 1]$. The area is then

$$\int_0^1 x^2 dx.$$

In order to solve this integral, we first look at the indefinite integral

$$\int x^2 dx.$$

The Fundamental Theorem of Calculus 11.(a) is telling us to first find the primitive function of $f(x) = x^2$. From our classes on derivatives we know that for $F(x) = x^{n+1}$

$$F'(x) = (n+1)x^n,$$

therefore, for $n = 2$, we write

$$\int x^2 dx = \int \frac{3}{3} x^2 dx = \frac{1}{3} \int F'(x) dx.$$

This implies that the primitive of x^n is

$$F(x) = \frac{x^{n+1}}{n+1}.$$

Consequently the solution to the indefinite integral is

$$\int x^2 dx = \frac{x^3}{3} + C.$$

Note the constant of integration C . By rules of differentiation, we know that the derivative of C is zero. Thus, we only know the solution up to the value of this constant. How do we find it? We use the information that we are looking at the area only between $x = 0$ and $x = 1$. Indeed, as $F(x) = \frac{x^3}{3} + C$

$$F(0) = C$$

and

$$F(1) = \frac{1}{3} + C.$$

The Fundamental Theorem of Calculus 11.(b) tells us that the value of our definite integral is

$$\int_0^1 x^2 dx = F(1) - F(0) = \frac{1}{3} + C - C = \frac{1}{3}.$$

Remark 6 (A note for the future) There is a generalized version of part (a) of the Fundamental Theorem of Calculus due to Leibniz which might turn out fairly useful in future applications. It is known as Leibniz rule and states that

$$\frac{d}{dx} \int_{a(x)}^{b(x)} f(y; x) dy = f(x, b(x)) \frac{d}{dx} b(x) - f(x, a(x)) \frac{d}{dx} a(x) + \int_{a(x)}^{b(x)} \frac{d}{dx} f(y; x) dy.$$

1.6.3 Rules of integration for the Indefinite Integral

As we know how to differentiate, we also know how to integrate. We need our Table summarising the rules of derivation and we are ready to go.

For ease of illustration, the Table is repeated here, now written in terms of integrand function and its primitive. We need to pay a little care when it comes to the chain rule and the rule of the product. The corresponding rules in integration are called ‘integration by substitution’ and ‘integration by parts’ respectively. So these two cases are treated in more details in the following.

Table 5: Summary: the rules of integration

$\int \mathbf{f}(\mathbf{x})d\mathbf{x}$	$\mathbf{F}(\mathbf{x})$
$\int k dx$	$kx + C$
$\int k g(x) dx$	$k \int g(x) dx + C$
$\int h(g(x))g'(x)dx = \int h(u)\frac{du}{dx}dx$	$\int h(u)du$ [Integration by Substitution]
$\int x^n dx$	$\frac{x^{n+1}}{n+1} + C$
$\int u(x)^n u'(x) dx$	$\frac{u(x)^{n+1}}{n+1} + C$
$\int (u(x) \pm g(x)) dx$	$\int u(x) dx \pm \int g(x) dx$
$\int e^x dx$	$e^x + C$
$\int e^{u(x)} u'(x) dx$	$e^{u(x)} + C$
$\int \frac{1}{x} dx$	$\ln(x) + C$
$\int \frac{u'(x)}{u(x)} dx$	$\ln(u(x)) + C$
$\int u(x)v'(x) dx$	$u(x)v(x) - \int u'(x)v(x) dx$ [Integration by Parts]

Example 45 1.

$$\int (-2) dx = -2 \int dx = -2x + C.$$

Verify by differentiation: $F(x) = -2x + C$ implies $F'(x) = -2$

2.

$$\int x^{-2} dx = \frac{x^{-2+1}}{-2+1} + C = -\frac{1}{x} + C.$$

Verify by differentiation: $F(x) = -\frac{1}{x} + C$ implies $F'(x) = -(-1)x^{-1-1} = x^{-2}$.

3.

$$\int 10x^5 dx = 10 \int x^5 dx = 10 \frac{x^6}{6} + C = \frac{5}{3}x^6 + C.$$

Verify by differentiation: $F(x) = \frac{5}{3}x^6 + C$ implies $F'(x) = \frac{5}{3}6x^{6-1} = 10x^5$.

4.

$$\int \sin(x) dx = -\cos(x) + C.$$

We have seen in fact that $\cos'(x) = -\sin(x)$.

5.

$$\int \cos(x) dx = \sin(x) + C.$$

We have seen in fact that $\sin'(x) = \cos(x)$.

Integration by Substitution This is the reverse of the chain rule and it states that

$$\int h(g(x))g'(x)dx = \int h(u)du.$$

Indeed, consider first

$$\int h(g(x))g'(x)dx.$$

Set $u = g(x)$. Now differentiate to obtain

$$\frac{du}{dx} = g'(x)$$

or, equivalently using the differential

$$du = g'(x)dx.$$

Then we note

$$\int \underbrace{h(g(x))}_u \underbrace{g'(x)dx}_{du}$$

and the result follows.

Example 46 1. Solve

$$\int (x+1)^3 dx.$$

Set $u = x + 1$. Then $du = dx$ and the integral now reads

$$\int (x+1)^3 dx = \int u^3 du.$$

We can solve this using the power rule so that

$$\int u^3 du = \frac{u^4}{4} + C.$$

Final step: remember that $u = x + 1$, thus

$$\int (x+1)^3 dx = \frac{(x+1)^4}{4} + C.$$

2. Consider the integral

$$\int (ax+b)^n dx.$$

How do we solve it? Let's go by substitution by setting $u = ax + b$. Then $du = adx$ i.e. $dx = du/a$. Consequently

$$\int (ax + b)^n dx = \int u^n \frac{du}{a} = \frac{1}{a} \int u^n du.$$

Using the power rule and then switching back to the variable x we obtain

$$\int (ax + b)^n dx = \frac{1}{a} \int u^n du = \frac{1}{a(n+1)} u^{n+1} + C = \frac{1}{a(n+1)} (ax + b)^{n+1} + C.$$

3. Solve

$$\int \frac{3x^2 + 2x + 4}{x^3 + x^2 + 4x} dx.$$

Note that if $u = x^3 + x^2 + 4x$, then $du = (3x^2 + 2x + 4)dx$. Therefore by substitution

$$\int \frac{3x^2 + 2x + 4}{x^3 + x^2 + 4x} dx = \int \frac{du}{u} = \ln(u) + C = \ln(x^3 + x^2 + 4x) + C.$$

4. Solve

$$\int \frac{x}{x^2 + 1} dx.$$

We note that if we set $u = x^2 + 1$, then $du = 2xdx$. Consequently

$$\int \frac{x}{x^2 + 1} dx = \int \frac{du}{2u} = \frac{1}{2} \ln(u) + C = \frac{1}{2} \ln(x^2 + 1) + C$$

5. Solve

$$\int 2xe^{x^2} dx.$$

Set $u = x^2$ so that $du = 2xdx$. Then

$$\int 2xe^{x^2} dx = \int e^u du = e^u + C = e^{x^2} + C.$$

6. Solve

$$\int \cos(x^2) 2x dx.$$

We note that $2xdx$ is the differential of x^2 . Then set $u = x^2$, so that $du = 2xdx$ and substitute to obtain

$$\int \cos(x^2) 2x dx = \int \cos(u) du = \sin(u) + C = \sin(x^2) + C.$$

This holds as when we differentiate $\sin(u)$ we obtain $\cos(u)$.

Remark 7 We can now give a proper meaning to the dx inside the integral: it denotes that we have differentiated with respect to x . It is indeed your differential. Consequently

$$F(x) = \int f(x) dx = \int dF(x).$$

Integration by parts This rule is also known as the reverse product rule (of derivatives). It states that

$$\int f(x)g'(x)dx = f(x)g(x) - \int f'(x)g(x)dx.$$

Indeed, let us define the function $q(x) = f(x)g(x)$. Let us now use the product rule to find the derivative of the function $q(x)$:

$$q'(x) = \frac{dq}{dx} = f'(x)g(x) + f(x)g'(x).$$

Hence, $dq = (f'(x)g(x) + f(x)g'(x))dx$, or equivalently

$$d(f(x)g(x)) = f'(x)g(x)dx + f(x)g'(x)dx.$$

Rearrange the terms to obtain

$$f(x)g'(x)dx = dq - f'(x)g(x)dx.$$

Now, integrate both sides and note that $\int dq = q = f(x)g(x)$ (up to the constant of integration).

Example 47 1. Solve

$$\int x^2 \ln(x)dx.$$

Let us set $f(x) = \ln(x)$ (we choose the one which is not quite immediate to solve) and $g'(x) = x^2$, for which we know that $g(x) = x^3/3$. Remember that $f'(x) = 1/x$. Then

$$\int x^2 \ln(x)dx = \ln(x)\frac{x^3}{3} - \int \frac{1}{x} \frac{x^3}{3}dx = \ln(x)\frac{x^3}{3} - \frac{1}{3} \int x^2 dx.$$

We can now solve the final integral by means of the power rule, so that

$$\int x^2 \ln(x)dx = \frac{x^3}{3} \ln(x) - \frac{x^3}{9} + C$$

2. Solve

$$\int xe^x dx.$$

Let us set $f(x) = x$ - we choose in this way as in the resolution step we need to differentiate f and for the identity function the derivative is simply 1 - very convenient! Also, we take $g'(x) = e^x$ - as it is quite straightforward to find its primitive function, i.e. $g(x) = e^x$. Then by the rule by parts we obtain

$$\int xe^x dx = xe^x - \int e^x dx = xe^x - e^x + C = e^x(x - 1) + C.$$

3. Solve

$$\int e^x \sin(x)dx.$$

We can try by choosing $f(x) = \sin(x)$ and $g'(x) = e^x$. Then we obtain

$$\int e^x \sin(x)dx = e^x \sin(x) - \int e^x \cos(x)dx,$$

as $\sin'(x) = \cos(x)$. At this stage, it does not look like we made much progress. We might try to apply the rule by parts once more, setting $g'(x) = e^x$ again and $f(x) = \cos(x)$. If we do this we obtain that

$$\int e^x \cos(x) dx = e^x \cos(x) + \int e^x \sin(x) dx.$$

Thus, if we plug in into the original integral we were trying to solve, we obtain

$$\begin{aligned} \int e^x \sin(x) dx &= e^x \sin(x) - \int e^x \cos(x) dx \\ &= e^x \sin(x) - \left(e^x \cos(x) + \int e^x \sin(x) dx \right) \\ &= e^x \sin(x) - e^x \cos(x) - \int e^x \sin(x) dx. \end{aligned}$$

Have we managed to solve the integral? Well, yes: we just need to rearrange the terms to write

$$\int e^x \sin(x) dx = e^x \frac{\sin(x) - \cos(x)}{2} + C.$$

The above examples show that you might have to use integration by parts more than once in order to be able to arrive at the correct solution. The choice of which function is f and which one is g' is very important and can dictate whether or not we are going to be successful in solving the integral. It takes quite a bit of practice.

1.6.4 Applications of integrals - some

Integration in Finance appears in many instances, from the calculation of the present value of a continuous investment flow under continuous compounding, to the price of sophisticated securities called derivatives such as European vanilla options. This section, aims at providing a quick overview.

Present Value of a continuous investment flow Suppose the investor has a continuous investment flow S per year over the next T years, an important information regarding this investment flow is its present value today. Assume we use continuous compounding, then the present value is

$$PV_0 = \int_0^T S e^{-rt} dt.$$

Example 48 Assume we want the present value of £1,000 to be paid for 3 years when $r = 5\%$ p.a. Then

$$PV_0 = \int_0^3 1,000 e^{-0.05t} dt = 1,000 \frac{1 - e^{-0.15}}{0.05} = 2,785.8.$$

Capital formation By definition the net investment rate I is the rate of change in the capital formation $K(t)$ overtime. If the process of capital formation is continuous in time, then

$$I(t) = \frac{d}{dt}K(t)$$

and therefore

$$K(t) = \int I(t)dt$$

up to a constant of integration C .

Example 49 Assume $I(t) = 140t^{3/4}$ and $K(0) = 150$. Then

$$K(t) = 80t^{7/4} + 150.$$

Total vs Marginal We have seen in Section 1.5, that total costs and revenues are linked to marginal costs and revenues. Seen in the light of integration, this relationship reads

$$TC(q) = \int MC(q)dq$$

and

$$TR(q) = \int MR(q)dq.$$

Example 50 Assume $MC(q) = 100 + q$ and that $TC(100) = 40,000$. Then

$$TC(q) = 25,000 + 100q + \frac{q^2}{2}.$$

1.6.5 Improper Integrals

In our study of definite integrals so far, we have considered an integrand function f defined on a finite interval $[a, b]$. We now extend the concept of a definite integral to the case in which the interval is infinite and also to the case in which f has an infinite discontinuity in $[a, b]$. In either case the integral is called an improper integral - type I in the first case, and type II in the second case.

Example 51 1. Consider the integral

$$\int_0^{\infty} \lambda e^{-\lambda x} dx, \quad \lambda > 0.$$

This is an improper integral of type I as the upper limit of integration tends to infinity.

2. Similarly the integral

$$\int_0^{\infty} e^x dx$$

is improper of type I.

3. Consider the integral

$$\int_0^1 \frac{1}{x-1} dx.$$

This is an improper integral of type II as the upper limit of integration is a point of discontinuity (singularity) of the integrand function.

4. Consider the integral

$$\int_0^\infty x^{-p} dx.$$

This is both an integral of type I and type II, as the upper limit tends to infinity whilst the lower limit is a point of discontinuity.

There are quite a few occurrences of improper integrals in financial applications, mainly when tools from Statistics are used such as expectation, variance, and moment generating function for example. Thus, the scope of this section is to understand how we can determine a priori if the integral is convergent or not. We want to have this piece of knowledge before attempting solving the integral because, as we have seen in the previous sections, it might not always be that straightforward solving integrals, so we want to attempt this only when we know it is worth doing.

Improper Integral of Type I

Definition 12 (a) Let $f(x)$ be a continuous function on the interval $[a, \infty)$. We define the improper integral as

$$\int_a^\infty f(x) dx = \lim_{b \rightarrow \infty} \int_a^b f(x) dx$$

provided this limit exists (as a finite number).

(b) Similarly, if a continuous function $f(x)$ is given on the interval $(-\infty, b]$, the improper integral of $f(x)$ is defined as

$$\int_{-\infty}^b f(x) dx = \lim_{a \rightarrow -\infty} \int_a^b f(x) dx$$

provided this limit exists (as a finite number).

The improper integrals are called convergent if the corresponding limit exists and divergent if the limit does not exist.

(c) If both

$$\int_a^\infty f(x) dx$$

and

$$\int_{-\infty}^a f(x) dx$$

are convergent, then we define

$$\int_{-\infty}^\infty f(x) dx = \int_a^\infty f(x) dx + \int_{-\infty}^a f(x) dx$$

for any $a \in \mathbb{R}$.

Example 52 Let us consider some of the integrals introduced in Example 51.

1. Let us first consider

$$\int_0^{\infty} \lambda e^{-\lambda x} dx, \quad \lambda > 0.$$

We solve the indefinite form by substitution, so that

$$\int \lambda e^{-\lambda x} dx = -e^{-\lambda x} + C.$$

The improper integral is then

$$\int_0^{\infty} \lambda e^{-\lambda x} dx = -\lim_{b \rightarrow \infty} e^{-\lambda b} + 1 = 1$$

as

$$\lim_{b \rightarrow \infty} e^{-\lambda b} = 0.$$

Thus, the given integral is always convergent.

2. Consider the integral

$$\int_0^{\infty} e^x dx.$$

This integral is divergent as

$$\int_0^{\infty} e^x dx = \lim_{b \rightarrow \infty} e^b - 1$$

but the exponential function grows to infinity, and therefore

$$\int_0^{\infty} e^x dx = \infty.$$

3. Let us now consider

$$\int_1^{\infty} x^{-p} dx.$$

Let us first consider the case $p \neq 1$. Then, the indefinite form has solution

$$\int x^{-p} dx = \frac{1}{1-p} \frac{1}{x^{p-1}} + C.$$

Consequently, the given improper integral has value

$$\int_1^{\infty} x^{-p} dx = \frac{1}{1-p} \lim_{b \rightarrow \infty} \frac{1}{b^{p-1}} - \frac{1}{1-p}.$$

Does the limit always exist? Are there values of p such that the integral is convergent, and other values of p such that the integral diverges? We need to explore this in more details. Let us start from the case $p > 1$ so that $p - 1 > 0$. Then

$$\lim_{b \rightarrow \infty} \frac{1}{b^{p-1}} = 0$$

and

$$\int_1^{\infty} x^{-p} dx = \frac{1}{p-1}.$$

So it is convergent. Now, let us consider the case $p < 1$ so that $p - 1 < 0$. Then

$$\lim_{b \rightarrow \infty} \frac{1}{b^{p-1}} = \lim_{b \rightarrow \infty} b^{1-p} = \infty.$$

Because this integral does not exist, the original integral

$$\int_1^\infty x^{-p} dx$$

is divergent.

The final case to consider is $p = 1$. In this case the given solution to the indefinite form does not work; as a matter of fact we know that

$$\int x^{-1} dx = \ln(x) + C.$$

We also know that the $\ln(x) \rightarrow \infty$ for $x \rightarrow \infty$. Therefore the integral is divergent. Summarising, we have that

$$\int_1^\infty x^{-p} dx$$

exists and converges for $p > 1$, and diverges for $p \leq 1$. This is particularly useful information for the application of the so-called comparison test, which we will see next.

Remark 8 The exponential distribution has density function

$$\lambda e^{-\lambda x}, \quad x > 0$$

The previous example (part 1) shows that the integrated density function over its domain returns 1.

Proposition 13 (Comparison Test: Type I) *Let $f(x)$ and $g(x)$ be continuous functions on the interval $[a, \infty)$. Suppose that $0 \leq g(x) \leq f(x)$ for all x in the interval $[a, \infty)$.*

- (a) *If $\int_a^\infty f(x) dx$ is convergent, then $\int_a^\infty g(x) dx$ is also convergent.*
- (b) *If $\int_a^\infty g(x) dx$ is divergent, then $\int_a^\infty f(x) dx$ is also divergent.*
- (c) *If $\int_a^\infty |f(x)| dx$ is convergent, then $\int_a^\infty f(x) dx$ is also convergent, and in this case, we say that the integral $\int_a^\infty f(x) dx$ is absolutely convergent.*

Similar results hold for $f(x)$ and $g(x)$ continuous functions on the interval $(-\infty, b]$.

The meaning of the above result is essentially the following: an area smaller than something that converges, converges. An area bigger than something that diverges, diverges.

Example 53 1. Consider the integral

$$\int_1^{\infty} e^{-x^2} dx.$$

We note that $0 \leq e^{-x^2} \leq e^{-x}$ for any $x \geq 1$. We know from Example 52 that integrals of the form

$$\int_1^{\infty} e^{-x} dx$$

converges, and so does

$$\int_1^{\infty} e^{-x^2} dx.$$

We do not solve this integral here. It is actually not easy to solve as the primitive is a transcendental function. However, we are going to have a look at it again later on in this section, albeit on a different interval of integration.

2. Consider the integral

$$\int_1^{\infty} \frac{\sin^2(x)}{x^2} dx.$$

As $-1 \leq \sin(x) \leq 1$, then

$$0 \leq \frac{\sin^2(x)}{x^2} \leq \frac{1}{x^2}.$$

From Example 52 we know that

$$\int_1^{\infty} \frac{1}{x^2} dx$$

converges ($p > 1$) and therefore also the original integral converges.

The above result is kind of implying some sort of proportionality, in the limit, between the integrand function, say f , and the test function g . The point then becomes to choose a good test function which tells us everything we need to know about the improper integral we want to quantify. Such a useful choice is the function x^{-p} . From the above examples, in fact, we know how the improper integral of x^{-p} behaves depending on the value of p . Thus, we reformulate our Comparison Test in a useful and operational manner as follows.

Limit Comparison Test - Type I

Consider the improper integrals of type I

$$\int_a^{\infty} f(x) dx, \quad \int_{-\infty}^b f(x) dx,$$

and the corresponding limit comparison test (LCT)

$$\lim_{x \rightarrow \pm\infty} \frac{f(x)}{\left(\frac{1}{x}\right)^p} = L.$$

Then

$p > 1, L \in \mathbb{R}$	the integral exists and is finite, i.e. it converges
$p \leq 1, L \in \mathbb{R}^* \setminus \{0\}$	the integral does not exist, i.e. it diverges
$p \leq 1, L = 0$	no answer from the test.

Note $\mathbb{R}^*$ in the above means the set of the real numbers extended to accommodate $\pm\infty$. Further, if $p > 1$ but $L \rightarrow \pm\infty$, the integral diverges: f , when compared to g , is infinitely bigger and therefore the improper integral diverges.

Finally, the test does not provide an answer if $p \leq 1, L = 0$ because it means that g is infinitely bigger than f , so it could be that the improper integral of f converges anyhow even if the one of g does not.

Example 54 1. Consider the integral

$$\int_a^\infty \frac{\gamma}{\pi(\gamma^2 + x^2)} dx.$$

We use the LCT and look at the limit

$$\lim_{x \rightarrow \infty} \frac{\frac{\gamma}{\pi(\gamma^2 + x^2)}}{x^{-p}} = \lim_{x \rightarrow \infty} \frac{\gamma x^p}{\pi(\gamma^2 + x^2)} = \frac{\gamma}{\pi}$$

if $p = 2 > 1$. We know that the type I improper integral of x^{-p} converges for $p > 1$, therefore as in this case $p = 2$, the integral

$$\int_a^\infty \frac{\gamma}{\pi(\gamma^2 + x^2)} dx$$

converges as well.

2. Consider the integral

$$\int_a^\infty \frac{\gamma x}{\pi(\gamma^2 + x^2)} dx.$$

We follow the same reasoning as in the previous example, and check the limit

$$\lim_{x \rightarrow \infty} \frac{\frac{\gamma x}{\pi(\gamma^2 + x^2)}}{x^{-p}} = \lim_{x \rightarrow \infty} \frac{\gamma x^{p+1}}{\pi(\gamma^2 + x^2)} = \frac{\gamma}{\pi}$$

if $p + 1 = 2$, i.e. $p = 1$. However, we know that the improper integral of type I of x^{-1} diverges. And therefore

$$\int_a^\infty \frac{\gamma x}{\pi(\gamma^2 + x^2)} dx$$

diverges as well.

3. Consider the integral

$$\int_a^\infty \frac{\gamma e^x}{\pi(\gamma^2 + x^2)} dx.$$

Via the LCT we obtain for $p = 2$

$$\lim_{x \rightarrow \infty} \frac{\frac{\gamma e^x}{\pi(\gamma^2 + x^2)}}{x^{-p}} = \lim_{x \rightarrow \infty} \frac{\gamma x^p e^x}{\pi(\gamma^2 + x^2)} = \infty$$

therefore the improper integral diverges.

4. Consider the integral

$$\int_0^{\infty} x^{a-1} e^{-x} dx,$$

for $a > 0$. We focus only on the upper limit - we will come back to the lower limit once we cover the type II improper integral. Then by LCT with $p = 2$

$$\lim_{x \rightarrow \infty} \frac{x^{a-1} e^{-x}}{x^{-p}} = \lim_{x \rightarrow \infty} \frac{x^{a-1} x^p}{e^x} = 0.$$

Thus, the integral converges.

Remark 9 The Cauchy distribution has density function

$$\frac{\gamma}{\pi(\gamma^2 + x^2)}, \quad x \in \mathbb{R}.$$

The previous example (part 1) shows that the integrated density function over its domain returns 1. However, it can be deduced from part 2 and 3 of the previous example that for this distribution we cannot compute the mean and the moment generating function. As a matter of fact we cannot obtain any moment.

Remark 10 The improper integral that we saw in the previous example - part 4 - is known as the Gamma function with parameter a , denoted as $\Gamma(a)$. It is the generalization of the factorial to non integer numbers. Indeed you can prove by integrating by parts that for

$$\Gamma(a) = \int_0^{\infty} x^{a-1} e^{-x} dx$$

the following recursion holds

$$\Gamma(a+1) = a\Gamma(a)$$

Example 55 We want to prove that

$$I_1 = \int_0^{\infty} e^{-x^2} dx = \frac{\sqrt{\pi}}{2}.$$

To begin with, consider the following integral defined for any $s > 0$

$$\int_0^{\infty} e^{-s y^2} dy.$$

We tackle this integral by substitution. Thus, set $x = y\sqrt{s}$. Consequently $dx = \sqrt{s}dy$. Then

$$\int_0^{\infty} e^{-s y^2} dy = \frac{1}{\sqrt{s}} \int_0^{\infty} e^{-x^2} dx = \frac{I_1}{\sqrt{s}}. \quad (5)$$

Now, consider the quantity I_1^2 . It follows that

$$\begin{aligned} I_1^2 &= I_1 \int_0^{\infty} e^{-x^2} dx \\ &= \int_0^{\infty} I_1 e^{-x^2} dx. \end{aligned}$$

By substitution, setting $u = x^2$ so that $du = 2x dx$, i.e. $dx = du/(2x) = du/(2\sqrt{u})$, we obtain

$$\begin{aligned} I_1^2 &= \int_0^\infty I_1 e^{-x^2} dx \\ &= \int_0^\infty \frac{I_1}{\sqrt{u}} \frac{e^{-u}}{2} du \\ &= \int_0^\infty \int_0^\infty \frac{e^{-u}}{2} e^{-u y^2} dy du \end{aligned}$$

where the last step follows from equation (5). Thus

$$I_1^2 = \frac{1}{2} \int_0^\infty \left(\int_0^\infty e^{-u(1+y^2)} du \right) dy;$$

use substitution by setting $x = u(1+y^2)$ so that $dx = du(1+y^2)$. Then

$$\begin{aligned} I_1^2 &= \frac{1}{2} \int_0^\infty \left(\int_0^\infty e^{-u(1+y^2)} du \right) dy \\ &= \frac{1}{2} \int_0^\infty \left(\frac{-e^{-u(1+y^2)}}{1+y^2} \Big|_0^{b \rightarrow \infty} \right) dy \\ &= \frac{1}{2} \int_0^\infty \frac{1}{1+y^2} dy \\ &= \frac{1}{2} \arctan(y) \Big|_0^{b \rightarrow \infty} \\ &= \frac{\pi}{4}. \end{aligned}$$

Therefore

$$I_1 = \frac{\sqrt{\pi}}{2}$$

as required.

Now consider the integral

$$\int_{-\infty}^\infty e^{-\frac{x^2}{2}} dx.$$

We can use the result above to calculate its value. Indeed, operate by substitution setting $y = x/\sqrt{2}$ to obtain

$$\int_{-\infty}^\infty e^{-\frac{x^2}{2}} dx = \sqrt{2} \left(\int_{-\infty}^\infty e^{-y^2} dy \right),$$

and note that

$$\int_{-\infty}^\infty e^{-y^2} dy = \int_0^\infty e^{-y^2} dy + \int_0^\infty e^{-z^2} dz = 2I_1,$$

where in the second integral we operated the substitution $y = -z$ with corresponding readjustment of the interval of integration. Then,

$$\int_{-\infty}^\infty e^{-\frac{x^2}{2}} dx = \sqrt{2} \times 2 \times \frac{\sqrt{\pi}}{2}$$

which returns

$$\int_{-\infty}^\infty e^{-\frac{x^2}{2}} dy = \sqrt{2\pi}.$$

An alternative proof uses polar coordinates.

Remark 11 A fundamental statistical distribution used in Finance is the standard Normal distribution which has density function

$$\frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}.$$

The previous example shows that the integrated density function over the entire real line returns 1, i.e.

$$\int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} = 1.$$

Another digression in Statistics: Characteristic Functions Moment generating functions suffer the disadvantage that the integrals which define them may not always be finite.

Example 56 A Cauchy random variable X has density function

$$f(x) = \frac{1}{\pi(1+x^2)} \quad x \in \mathbb{R}.$$

Hence the MGF of X is given by

$$M_X(k) = \int_{-\infty}^{\infty} \frac{e^{kx}}{\pi(1+x^2)} dx.$$

This is an improper integral of the 1st kind which does not converge unless $k = 0$ (which of course is a ‘nonsense’...) In fact, if you perform the convergence test, you obtain that

$$\lim_{x \rightarrow \infty} \frac{\frac{e^{kx}}{\pi(1+x^2)}}{\left(\frac{1}{x}\right)^\alpha} \stackrel{\alpha=2}{=} \frac{1}{\pi} \lim_{x \rightarrow \infty} e^{kx} = \begin{cases} 0 & \text{if } k < 0 \\ \infty & \text{if } k > 0, \end{cases}$$

$$\lim_{x \rightarrow -\infty} \frac{\frac{e^{kx}}{\pi(1+x^2)}}{\left(\frac{1}{x}\right)^\alpha} \stackrel{\alpha=2}{=} \frac{1}{\pi} \lim_{x \rightarrow -\infty} e^{kx} = \begin{cases} 0 & \text{if } k > 0 \\ \infty & \text{if } k < 0. \end{cases}$$

Hence, the MGF of a Cauchy random variable does not exist.

Characteristic functions are another class of functions equally useful and whose finiteness is guaranteed (so you might actually want to say that they are more useful than MGFs...).

Definition 14 The characteristic function of X is the function $\phi_X : \mathbb{R} \rightarrow \mathbb{C}$ defined by

$$\phi_X(u) = \mathbb{E}(e^{iuX})$$

for $i = \sqrt{-1}$.

This is a common transformation and is often called the Fourier transform of the density f of X if this quantity exists. In this case

$$\phi_X(u) = \int e^{iux} dF_X(x) = \int e^{iux} f_X(x) dx.$$

The characteristic function of a random variable has several nice properties. Firstly it always exists and it is finite (in $\mathbb{L}^1$): note that

$$\phi_X(u) = \mathbb{E}(e^{iuX}) = \mathbb{E}(\cos(uX) + i \sin(uX)),$$

hence¹

$$|\cos(uX) + i \sin(uX)| := \sqrt{\cos(uX)^2 + \sin(uX)^2} = 1.$$

Then

$$|\mathbb{E}(e^{iuX})| \leq \mathbb{E}(|e^{iuX}|) = 1.$$

Moreover.

- There is bijection between cumulative distribution functions/probability density function and characteristic functions: two distinct probability distributions never share the same characteristic function.
- if X and Y are independent random variables, $\phi_{X+Y}(u) = \phi_X(u) \phi_Y(u)$;
- if $a, b \in \mathbb{R}$ and $Y = aX + b$, then $\phi_Y(u) = e^{iub} \phi_X(au)$.

Examples of characteristic functions * Calculations of integrals involving complex numbers are not always pleasant; usually you should know about contour integration... but for our purposes you can get away with only knowing about analytic continuation.

Analytic continuation provides a way of extending the domain over which a complex function is defined. Let us start from a *complex function* f (like the characteristic function); this function is *complex differentiable* at z_0 and has derivative A if and only if

$$f(z) = f(z_0) + A(z - z_0) + o(z - z_0), \quad \forall z \in \mathbb{C}.$$

A complex function is said to be *analytic* on a region D if it is complex differentiable at every point in D (i.e. has no singularities, i.e. points at which the function ‘blows up’ or becomes degenerate). Now, let f_1 and f_2 be analytic functions on domains D_1 and D_2 respectively, with $D_1 \subset D_2$, such that $f_1 = f_2$ on $D_1 \cap D_2$. Then f_2 is called the *analytic continuation* of f_1 to D_2 . Moreover, if it exists, the analytic continuation of f_1 to D_2 is unique.

Consider now the MGF M_X of some random variable X ; we can say that the function

$$\mathcal{M}_X(z) = \int_{-\infty}^{\infty} f(x) e^{zx} dx \quad z \in \mathbb{C}$$

is the analytic continuation of M_X to the complex plane, if it respects the condition above. Then, the characteristic function of X , ϕ_X , is the restriction of $\mathcal{M}_X$ to the imaginary axis, i.e.

$$\phi_X(u) = \mathcal{M}_X(iu)$$

¹Note that this is the complex square of the complex number $z = \cos(uX) + i \sin(uX)$, and you can interpret the notation as a norm.

Example 57 1. Let $X \sim N(0, 1)$. The characteristic function is

$$\phi_X(u) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{iux - \frac{x^2}{2}} dx.$$

Now consider the real valued function

$$M_X(k) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{kx - \frac{x^2}{2}} dx = e^{\frac{k^2}{2}},$$

i.e. the MGF of X . Since $\mathbb{R} \cap \mathbb{C} \neq \emptyset$, then M_X has analytic continuation on the complex plane given by

$$\mathcal{M}_X(z) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{zx - \frac{x^2}{2}} dx = e^{\frac{z^2}{2}} \quad z \in \mathbb{C}.$$

Therefore, by analytic continuation

$$\phi_X(u) = \mathcal{M}_X(iu) = e^{-\frac{u^2}{2}}.$$

2. Let X be a Poisson random variable with rate u . You can apply the same argument as above (i.e. analytic continuation) to show that

$$\phi_X(u) = \mathbb{M}_X(iu) = e^{\lambda(e^{iu}-1)}.$$

3. Consider now the gamma distribution. Analytic continuation implies that

$$\phi_X(u) = \left(\frac{\lambda}{\lambda - iu} \right)^\alpha.$$

4. Assume X is a Cauchy random variable, i.e.

$$f(x) = \frac{1}{\pi(1+x^2)}.$$

We cannot use the analytic continuation argument because the function is not analytic (can you spot why?). The CF exists anyhow and is equal to

$$\phi_X(u) = e^{-|u|}.$$

However, to prove the result you need to use contour integration and the residue theorem - which are beyond the scope of this module.

Similarly to the MGF, the characteristic function can be used to obtain the raw moments and the central moments of a random variable. In particular, the following hold.

$$\begin{aligned} i^n \mu_n &= \left. \frac{\partial \phi_X(u)}{\partial u} \right|_{u=0}, \\ i^n c_n &= \left. \frac{\partial \ln \phi_X(u)}{\partial u} \right|_{u=0} \end{aligned}$$

(the proof follows the same steps as in the case in which the MGF is used).

Improper Integral of Type II

Definition 15 (a) If f is continuous on $[a, b)$ and is discontinuous at b , then

$$\int_a^b f(x)dx = \lim_{y \rightarrow b^-} \int_a^y f(x)dx$$

if this limit exists (as a finite number).

(b) If f is continuous on $(a, b]$ and is discontinuous at a , then

$$\int_a^b f(x)dx = \lim_{y \rightarrow a^+} \int_y^b f(x)dx$$

if this limit exists (as a finite number).

The improper integral $\int_a^b f(x)dx$ is called convergent if the corresponding limit exists and divergent if the limit does not exist.

(c) If f has a discontinuity at c , for $a < c < b$, and both $\int_a^c f(x)dx$ and $\int_c^b f(x)dx$ are convergent, then we define

$$\int_a^b f(x)dx = \int_a^c f(x)dx + \int_c^b f(x)dx.$$

Example 58 1. Consider the integral

$$\int_0^4 \frac{1}{\sqrt{4-x}} dx.$$

We note that the square root at the denominator of the integrand function is defined for $x \leq 4$. However, for the fraction to be well defined $x \neq 4$. Therefore we have a type 2 improper integral. This is a simple integral to solve. So, let's first consider the indefinite form

$$\int \frac{1}{\sqrt{4-x}} dx = \int (4-x)^{-1/2} dx$$

and solve by substitution setting $u = 4 - x$ so that $du = -dx$. Then

$$\int \frac{1}{\sqrt{4-x}} dx = - \int u^{-1/2} du = -2\sqrt{4-x} + C.$$

Now we can compute the improper integral as definite integral

$$\int_0^4 \frac{1}{\sqrt{4-x}} dx = -2\sqrt{4-x} \Big|_0^{\lim_{x \rightarrow 4^-}} = 4.$$

So, this integral converges.

2. Consider the integral

$$\int_0^1 x^{-p} dx.$$

We have already seen the primitive of the indefinite form. This is

$$\frac{1}{1-p} \frac{1}{x^{p-1}}.$$

Then

$$\int_0^1 x^{-p} dx = \frac{1}{1-p} \frac{1}{x^{p-1}} \Big|_{\lim_{x \rightarrow 0+}}^1.$$

Similarly to the case of type I improper integral, we need to discuss the values of p as this will change the solution. We start with $p > 1$. Then $p - 1 > 0$ and the integral returns

$$\int_0^1 x^{-p} dx = \frac{1}{1-p} - \infty = -\infty.$$

Then the integral diverges.

Set now $p < 1$. Then

$$\int_0^1 x^{-p} dx = \frac{1}{1-p} x^{1-p} \Big|_{\lim_{x \rightarrow 0+}}^1 = \frac{1}{1-p}.$$

Thus the integral converges.

Finally, if $p = 1$, then

$$\int_0^1 x^{-1} dx = \ln(x) \Big|_{\lim_{x \rightarrow 0+}}^1 = \infty$$

and the integral diverges.

Summarising, the type II improper integral

$$\int_0^1 x^{-p} dx$$

converges for $p < 1$ and diverges for $p \geq 1$.

Also in the case of type II improper integrals, we would like to have a test which allows us to check if the integral converges or not before solving it. We have the LCT for this scope; in the case of the type II integral it reads as follows.

Limit Comparison Test - type II

Consider the improper integrals of type II

$$\int_a^b f(x) dx$$

with either a , or b or any $x \in (a, b)$, or all at the same time not in the domain of f , and the corresponding limit comparison test (LCT)

$$\lim_{x \rightarrow a+} \frac{f(x)}{\left(\frac{1}{x-a}\right)^p} = L; \quad \lim_{x \rightarrow b-} \frac{f(x)}{\left(\frac{1}{b-x}\right)^p} = L.$$

Then

$p < 1, L \in \mathbb{R}$	the integral exists and is finite
$p \geq 1, L \in \mathbb{R}^* \setminus \{0\}$	the integral does not exist
$p \geq 1, L = 0$	no answer from the test.

Example 59 1. Consider the integral

$$\int_0^3 \frac{dx}{x-1}.$$

Note that the domain of the integrand function does not include $x = 1$, as for $x = 1$ the function is undefined. Therefore, let us split the definite integral in two type II improper integrals

$$\int_0^3 \frac{dx}{x-1} = \int_0^1 \frac{dx}{x-1} + \int_1^3 \frac{dx}{x-1},$$

and let's check the LCT for both integrals. We start with the first one

$$\lim_{x \rightarrow 1^-} \frac{\frac{1}{x-1}}{(1-x)^{-p}} = \lim_{x \rightarrow 1^-} \frac{(1-x)^p}{x-1} = -1$$

for $p = 1$; therefore the integral $\int_0^1 \frac{dx}{x-1}$ diverges. Note that this implies that the original integral diverges as well, as it is enough for one of the parts to be undefined. But you can check the second integral as well.

We observe that in this case the integral could have been solved as well: the indefinite integral is

$$\int \frac{dx}{x-1} = \ln|x-1| + C.$$

For $x = 1$ the argument is 0 and the logarithm of 0 is $-\infty$. The function is portrayed in Figure 22.

2. Consider the Gamma function

$$\Gamma(a) = \int_0^\infty x^{a-1} e^{-x} dx.$$

We have already discussed the convergence for the upper limit. Now we want to look at the convergence for the lower limit $x \rightarrow 0$. The LCT leads to

$$\lim_{x \rightarrow 0_+} \frac{x^{a-1} e^{-x}}{x^{-p}} = \lim_{x \rightarrow 0_+} \frac{x^{a-1} x^p}{e^x}.$$

Let us consider first the case in which $0 < a < 1$ so that $a - 1 < 0$. Then

$$\lim_{x \rightarrow 0_+} \frac{x^{a-1} x^p}{e^x} = \lim_{x \rightarrow 0_+} \frac{x^p}{x^{1-a} e^x} = 1$$

for $p = 1 - a$; given the assumptions, this means $p < 1$. Therefore in this case we have convergence. Consider now the case in which $a > 1$. Then

$$\lim_{x \rightarrow 0_+} \frac{x^{a-1} x^p}{e^x} = 1$$

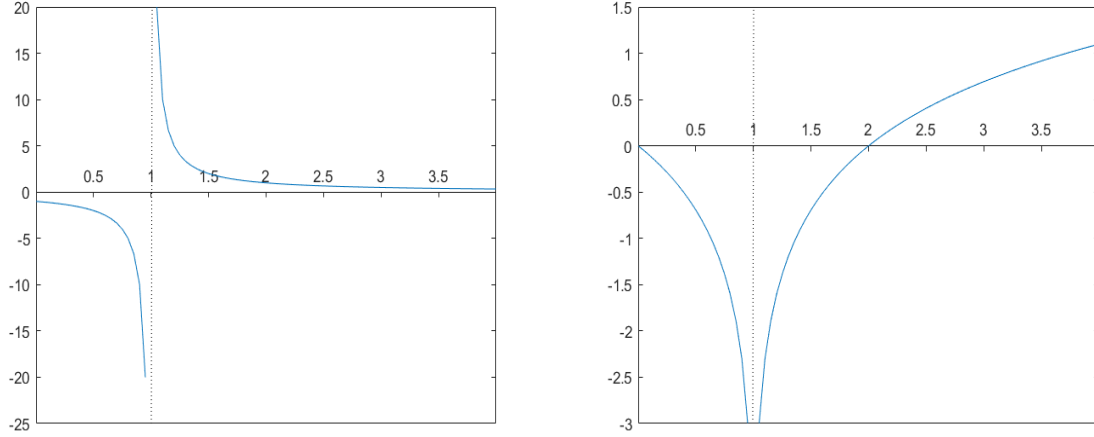


Figure 22: $f(x) = \frac{1}{x-1}$ and $F(x) = \ln|x-1|$. The asymptote $x = 1$ is indicated.

for $p = 1 - a$ which means $p < 0$. Therefore we have convergence also in this case. Summarising, the integral of the Gamma function converges.

2 Linear Algebra

Linear algebra is so named because it studies linear operators. Roughly speaking, this means that inputs of these operators are proportional to their outputs and that the operators are additive. Applications in Finance are almost ubiquitous. An example is as follows.

Example 60 You collected the daily returns of three stocks (Apple, Goldman Sachs, AT&T) for a (trading) week. These are

$$r_{AAPL} = \{0.005, -0.001, 0.006, -0.023, 0.002\}$$

$$r_{GS} = \{0.006, -0.012, -0.011, -0.004, 0.007\}$$

$$r_T = \{-0.015, 0.009, -0.007, 0.008, -0.004\}.$$

They can also be represented in the form of a table.

Stock (Ticker)	Trading Day				
	Day 1	Day 2	Day 3	Day 4	Day 5
Apple (AAPL)	0.005	-0.001	0.006	-0.023	0.002
Goldman Sachs (GS)	0.006	-0.012	-0.011	-0.004	0.007
AT&T (T)	-0.015	0.009	-0.007	0.008	-0.004

Now, suppose you have a budget of £100 invested in the these three assets, specifically you have £30 in AAPL, £35 in GS and the remaining in T. Then, you can calculate your total return each trading day as

$$r_{tot} = 0.30 \times r_{AAPL} + 0.35 \times r_{GS} + 0.35r_T.$$

One way of doing this, is to manually compute the return taking the data provided in the table. A more efficient alternative is to realize that the portfolio return is a linear combination of the returns of the components, and then use appropriate linear algebra operators. We will come back to this example in the following section.

The building blocks of linear algebra are matrices and vectors. We start from here.

2.1 Matrices

We all are familiar with numbers (integers, fractions, reals, ...). In linear algebra we refer to these numbers as scalars.

Matrices are an “ordered” collection of numbers which are stored in an object (the matrix). The reason to store those numbers in a matrix is mainly because matrix operations (and linear algebra in general) allows to write and compute linear combinations (i.e. combinations of sums and products) more efficiently and faster. Therefore, it becomes easier and more efficient to perform operations on these data, such as for example find a mathematical rule connecting the data points (a task that you will explore in Financial Econometrics in the following years). Ultimately, you can think of a matrix as a table.

Example 61 Back to Example 60. We have the returns of the three assets represented in the form of a table.

Stock (Ticker)	Trading Day				
	Day 1	Day 2	Day 3	Day 4	Day 5
Apple (AAPL)	0.005	-0.001	0.006	-0.023	0.002
Goldman Saches (G)	0.006	-0.012	-0.011	-0.004	0.007
AT&T (T)	-0.015	0.009	-0.007	0.008	-0.004

These numbers can be also stored in a 3×5 matrix such as

$$\mathbf{R} = \begin{bmatrix} 0.005 & -0.001 & 0.006 & -0.023 & 0.002 \\ 0.006 & -0.012 & -0.011 & -0.004 & 0.007 \\ -0.015 & 0.009 & -0.007 & 0.008 & -0.004 \end{bmatrix}$$

where the rows represent the corporations, while the columns represent the trading days. If we set $i = \{AAPL, GS, T\} = \{1, 2, 3\}$ and $j = \{\text{Day 1, Day 2, Day 3, Day 4, Day 5}\} = \{1, 2, 3, 4, 5\}$, each element of the matrix is $r_{i,j}$ that is the return for company i on day j .

More formally, a $r \times c$ matrix $\mathbf{M}$ is a table containing numbers ordered in r rows and c columns, that is

$$\mathbf{M} = \begin{bmatrix} m_{1,1} & m_{1,2} & \dots & m_{1,c} \\ m_{2,1} & m_{2,2} & \dots & m_{2,c} \\ \dots & \dots & \dots & \dots \\ m_{r,1} & m_{r,2} & \dots & m_{r,c} \end{bmatrix}.$$

Remark 12 When writing the element of a matrix, as well as the dimension of a matrix, we indicate first the rows and then the column. That is $m_{2,3}$ is the element of the matrix corresponding to row 2 column 3.

Further, the matrix dimension is indicated as $r \times c$, but other letters can be used such as $n \times m$ for example

A matrix is usually written with capital bold letters, while the single elements are indicated with lower case letters.

2.1.1 Special matrices

Square Matrix A square matrix is such that $r = c$ (same number of rows and column). If $r \neq c$, the matrix is called rectangular matrix.

Given a square matrix $\mathbf{M}$, the elements $m_{i,j}$, with $i = \{1, \dots, r\}$ and $j = \{1, \dots, r\}$, such that $i = j$, constitute the main diagonal of the matrix.

Further, if $m_{i,j} = m_{j,i}$ the matrix is called symmetric.

Example 62 Consider the matrices

$$\mathbf{A} = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \quad \mathbf{B} = \begin{bmatrix} 1 & 2 & \pi \\ 3 & 4 & e \end{bmatrix} \quad \mathbf{C} = \begin{bmatrix} 0.01 & 0.05 & -4 \\ 0.05 & 7 & -1.2 \\ -4 & -1.2 & -6 \end{bmatrix}$$

Using the definitions, $\mathbf{A}$ and $\mathbf{C}$ are square matrices, $\mathbf{B}$ is a rectangular matrix, and only $\mathbf{C}$ is symmetric. The main diagonal of $\mathbf{A}$ is $\{1, 4\}$ and the main diagonal of $\mathbf{C}$ is $\{0.01, 7, -6\}$. Note that the matrix $\mathbf{C}$ is symmetric: the off-diagonal elements are equal to each other in pairs, i.e. $c_{ij} = c_{ji}$. Indeed $c_{12} = c_{21}$, $c_{13} = c_{31}$ and $c_{23} = c_{32}$.

Diagonal Matrix A diagonal matrix is a square matrix with non-zero elements on the principal diagonal (the diagonal running from top left to bottom right of the matrix), and zeros elsewhere.

Example 63

$$\mathbf{A} = [21]; \quad \mathbf{B} = \begin{bmatrix} 1 & 0 \\ 0 & 7 \end{bmatrix}; \quad \mathbf{C} = \begin{bmatrix} \pi & 0 & 0 \\ 0 & 7 & 0 \\ 0 & 0 & \ln(5) \end{bmatrix};$$

Upper/Lower Triangular Matrix A matrix is upper triangular if all of the entries below the diagonal are zero, i.e., if $m_{rc} = 0$ whenever $r > c$.

A matrix is lower triangular if all of the entries above the diagonal are zero, i.e., if $m_{rc} = 0$ whenever $r < c$.

Example 64

$$\mathbf{A} = \begin{bmatrix} 3 & 8 & 20 \\ 0 & 7 & 5 \\ 0 & 0 & 1 \end{bmatrix}; \quad \mathbf{B} = \begin{bmatrix} 1 & 0 & 0 \\ 4 & 7 & 0 \\ 3 & 10 & 2 \end{bmatrix}$$

$\mathbf{A}$ is upper triangular; $\mathbf{B}$ is lower triangular.

Identity Matrix The identity matrix is defined as the $n \times n$ square matrix that has on the main diagonal 1s; all the other elements are 0s. The matrix is usually indicated as $\mathbf{I}_n$ (or just $\mathbf{I}$ if the dimension is clear).

$$\mathbf{I}_2 = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \quad \mathbf{I}_3 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad \dots \quad \mathbf{I}_5 = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \quad \dots$$

It is a special case of the diagonal matrix.

Vectors A r -dimensional (column) vector is a $r \times 1$ matrix, i.e. a matrix with r rows and one column

$$\mathbf{v} = \begin{pmatrix} v_1 \\ v_2 \\ \dots \\ v_r \end{pmatrix}.$$

(You can also define row vectors of dimension c as a $1 \times c$ matrix).

Vectors are usually indicated with bold lower-case letters. The n -dimensional vector of 0s is usually written as $\mathbf{0}_n$ and the n -dimensional vector of 1s as $\mathbf{1}_n$ (or just $\mathbf{0}$ and $\mathbf{1}$ respectively).

Example 65 Rethinking about Example 60, the quantities we hold in our portfolio of each

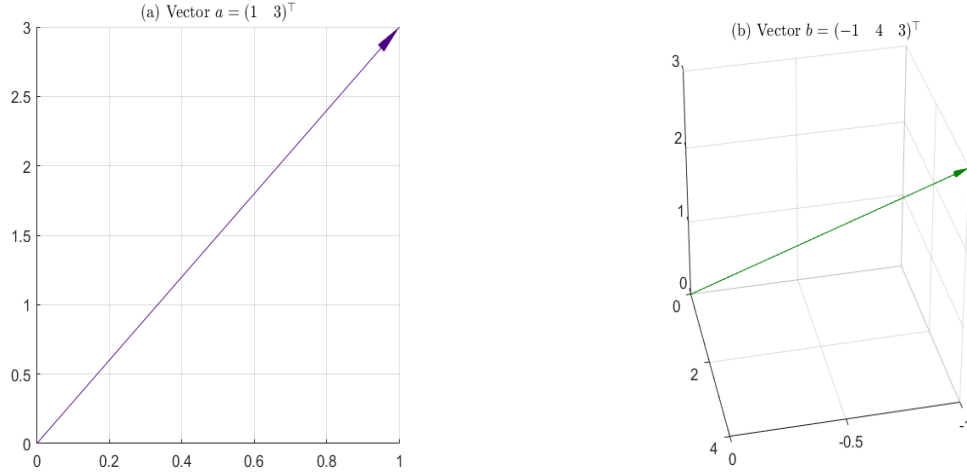


Figure 23: 2D and 3D vectors.

assets can be represented as a column vector

$$\mathbf{w} = \begin{pmatrix} 30/100 \\ 35/100 \\ 35/100 \end{pmatrix}.$$

Notice that vectors can also be interpreted as points on a plane.

Example 66 The vectors

$$\mathbf{a} = \begin{pmatrix} 1 \\ 3 \end{pmatrix} \quad \mathbf{b} = \begin{pmatrix} -1 \\ 4 \\ 3 \end{pmatrix} \quad \mathbf{c} = \begin{pmatrix} 2 \\ 3 \\ 0.5 \\ 1 \\ 3 \end{pmatrix}$$

are all points in the respective dimensional space. That is $\mathbf{a}$ is the point of x -coordinate 1 and y -coordinate 3 in the two-dimensional Cartesian plane; $\mathbf{b}$ is the point of x -coordinate -1 and y -coordinate 4 and z -coordinate 3 in the three-dimensional space; $\mathbf{c}$ is the point of x_1 -coordinate 2, x_2 -coordinate 3, x_3 -coordinate 0.5, x_4 -coordinate 1 and x_5 -coordinate 3 in a five-dimensional space. We show the vectors in Figure 23. Of course, only vectors with at most 3 elements can be drawn.

We can measure the length of a vector, i.e. the distance between the origin and the point with coordinates given by the elements of the vector, using the Euclidean norm

$$\|\mathbf{v}\| = \sqrt{v_1^2 + v_2^2 + \cdots + v_n^2}.$$

Example 67 Let us use the same vectors as in the previous example. Then

$$\|\mathbf{a}\| = \sqrt{1^2 + 3^2} = \sqrt{10}.$$

$$\|\mathbf{b}\| = \sqrt{(-1)^2 + 4^2 + 3^2} = \sqrt{26}.$$

$$\|\mathbf{c}\| = \sqrt{2^2 + 3^2 + 0.5^2 + 1^2 + 3^2} = \sqrt{23.25}.$$

We pay particular attention to linearly dependent vectors, as these will have an important role when we talk about rank of a matrix, determinant of a matrix and solution to linear systems.

Definition 16 (Linearly dependent vectors) A sequence of vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$ is said to be linearly dependent if there exist scalars $w_1, w_2, \dots, w_n$ not all zero such that

$$w_1 \mathbf{v}_1 + w_2 \mathbf{v}_2 + \dots + w_n \mathbf{v}_n = \mathbf{0}.$$

We call the expression

$$w_1 \mathbf{v}_1 + w_2 \mathbf{v}_2 + \dots + w_n \mathbf{v}_n$$

a linear combination of the vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$.

Example 68 Consider the vectors

$$\mathbf{v}_1 = \begin{bmatrix} 1 \\ 3 \end{bmatrix}, \quad \mathbf{v}_2 = \begin{bmatrix} 2 \\ 4 \end{bmatrix}, \quad \mathbf{v}_3 = \begin{bmatrix} 8 \\ 18 \end{bmatrix}.$$

These are linearly dependent as $\mathbf{v}_3 = 2\mathbf{v}_1 + 3\mathbf{v}_2$. Thus, the linear combination of these three vectors with weights $w_1 = 2$, $w_2 = 3$ and $w_3 = -1$ returns zero, i.e.

$$2\mathbf{v}_1 + 3\mathbf{v}_2 - \mathbf{v}_3 = \mathbf{0}$$

Note that in order to carry out the algebraic sum in the previous example, we operated the sum element by element. This is indeed how sums (and subtractions) work with vectors and matrices. We formalize this better in the following section.

Finally, linear combinations are very important in many fields such as Econometrics, for linear regression analysis, and Finance, for portfolio theory and derivative pricing. Indeed, portfolios of assets are linear combinations of these securities (see Example 60); in this case the weights are represented by the percentage of the total fund invested in each security.

2.1.2 Matrix operations

Transposition

Definition 17 Given a $n \times m$ matrix $\mathbf{A}$, its transpose is the $m \times n$ matrix $\mathbf{A}^\top$ obtained by turning all the rows of $\mathbf{A}$ into columns (and vice versa), that is

$$\mathbf{A} = \begin{bmatrix} a_{1,1} & a_{1,2} & \cdots & a_{1,m} \\ a_{2,1} & a_{2,2} & \cdots & a_{2,m} \\ \cdots & \cdots & \cdots & \cdots \\ a_{n,1} & a_{n,2} & \cdots & a_{n,m} \end{bmatrix} \quad \mathbf{A}^\top = \begin{bmatrix} a_{1,1} & a_{2,1} & \cdots & a_{n,1} \\ a_{1,2} & a_{2,2} & \cdots & a_{n,2} \\ \cdots & \cdots & \cdots & \cdots \\ a_{1,m} & a_{2,m} & \cdots & a_{n,m} \end{bmatrix}$$

Example 69 Given the matrices

$$\mathbf{A} = \begin{bmatrix} 1 & 4 & 6.5 \\ 0 & 0.8 & 13 \end{bmatrix} \quad \mathbf{B} = \begin{bmatrix} 0 & -7 & 8 \\ -7 & 11 & 9 \\ 8 & -5 & \sqrt{3} \end{bmatrix}$$

their transpose are

$$\mathbf{A}^\top = \begin{bmatrix} 1 & 0 \\ 4 & 0.8 \\ 6.5 & 13 \end{bmatrix} \quad \mathbf{B}^\top = \begin{bmatrix} 0 & -7 & 8 \\ -7 & 11 & -5 \\ 8 & 9 & \sqrt{3} \end{bmatrix}.$$

If we transpose a n -dimensional column vector, we obtain the correspondent n -dimensional row vector

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ \cdots \\ x_n \end{bmatrix} \quad \rightarrow \quad \mathbf{x}^\top = \begin{pmatrix} x_1 & x_2 & \cdots & x_n \end{pmatrix}.$$

Algebraic sum (sum and subtraction)

Definition 18 Given two $n \times m$ matrices $\mathbf{A}$ and $\mathbf{B}$,

$$\mathbf{A} = \begin{bmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,m} \\ a_{2,1} & a_{2,2} & \dots & a_{2,m} \\ \dots & \dots & \dots & \dots \\ a_{n,1} & a_{n,2} & \dots & a_{n,m} \end{bmatrix} \quad \mathbf{B} = \begin{bmatrix} b_{1,1} & b_{1,2} & \dots & b_{1,m} \\ b_{2,1} & b_{2,2} & \dots & b_{2,m} \\ \dots & \dots & \dots & \dots \\ b_{n,1} & b_{n,2} & \dots & b_{n,m} \end{bmatrix}$$

the algebraic sum of the two matrices is a $n \times m$ matrix $\mathbf{C} := \mathbf{A} \pm \mathbf{B}$ such that

$$\mathbf{C} = \begin{bmatrix} a_{1,1} \pm b_{1,1} & a_{1,2} \pm b_{1,2} & \dots & a_{1,m} \pm b_{1,m} \\ a_{2,1} \pm b_{2,1} & a_{2,2} \pm b_{2,2} & \dots & a_{2,m} \pm b_{2,m} \\ \dots & \dots & \dots & \dots \\ a_{n,1} \pm b_{n,1} & a_{n,2} \pm b_{n,2} & \dots & a_{n,m} \pm b_{n,m} \end{bmatrix} = \begin{bmatrix} c_{1,1} & c_{1,2} & \dots & c_{1,m} \\ c_{2,1} & c_{2,2} & \dots & c_{2,m} \\ \dots & \dots & \dots & \dots \\ c_{n,1} & c_{n,2} & \dots & c_{n,m} \end{bmatrix}.$$

As you can see, the sum of matrices is defined only if and only if the matrices have the same dimension. Furthermore, given a matrix $\mathbf{A}$, the matrix $\mathbf{O} := \mathbf{A} - \mathbf{A}$ is called null matrix and is composed by 0s only.

Example 70 Consider the matrices

$$\mathbf{A} = \begin{bmatrix} 1 & 1 & 1 \\ 2 & 2 & 2 \\ 3 & 3 & 3 \end{bmatrix} \quad \mathbf{B} = \begin{bmatrix} 3 & 2 & 1 \\ 2 & 1 & 0 \\ 1 & 0 & -1 \end{bmatrix} \quad \mathbf{C} = \begin{bmatrix} 0 & 1 & 2 \\ 3 & 4 & 5 \end{bmatrix} \quad \mathbf{D} = \begin{bmatrix} 0 & -1 \\ -1 & -2 \\ -2 & -3 \end{bmatrix}.$$

Then.

1. $\mathbf{E} = \mathbf{A} + \mathbf{B}$ is given by

$$\mathbf{E} = \begin{bmatrix} 1+3 & 1+2 & 1+1 \\ 2+2 & 2+1 & 2+0 \\ 3+1 & 3+0 & 3-1 \end{bmatrix} = \begin{bmatrix} 4 & 3 & 2 \\ 4 & 3 & 2 \\ 4 & 3 & 2 \end{bmatrix}.$$

2. $\mathbf{F} = \mathbf{A}^\top - \mathbf{B}^\top$ is given by

$$\mathbf{F} = \begin{bmatrix} 1-3 & 2-2 & 3-1 \\ 1-2 & 2-1 & 3-0 \\ 1-1 & 2-0 & 3-(-1) \end{bmatrix} = \begin{bmatrix} -2 & 0 & 2 \\ -1 & 1 & 3 \\ 0 & 2 & 4 \end{bmatrix}.$$

3. $\mathbf{G} = \mathbf{C}^\top + \mathbf{D}$ is given by

$$\mathbf{G} = \begin{bmatrix} 0+0 & -1 \\ 1-1 & 4-2 \\ 2-2 & 5-3 \end{bmatrix} = \begin{bmatrix} 0 & 2 \\ 0 & 2 \\ 0 & 2 \end{bmatrix}.$$

4. $\mathbf{H} = \mathbf{C} + \mathbf{D}^\top$

$$\mathbf{H} = \begin{bmatrix} 0+0 & 1-1 & 2-2 \\ 3-1 & 4-2 & 5-3 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 \\ 2 & 2 & 2 \end{bmatrix}.$$

5. $\mathbf{J} = \mathbf{C} + \mathbf{D}$ cannot be calculated as $\mathbf{C}$ is a 2×3 matrix, whilst $\mathbf{D}$ is a 3×2 matrix (the dimensions are different).

Properties of the (algebraic) sum

1. $\mathbf{A} + \mathbf{B} = \mathbf{B} + \mathbf{A}$ (commutative)
2. $\mathbf{A} + (\mathbf{B} + \mathbf{C}) = (\mathbf{A} + \mathbf{B}) + \mathbf{C}$ (associative)
3. $\mathbf{A} \pm \mathbf{O} = \mathbf{A}$ (identity element with respect to the algebraic sum)

Multiplication by a scalar

Definition 19 Given a $n \times m$ matrix $\mathbf{A}$,

$$\mathbf{A} = \begin{bmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,m} \\ a_{2,1} & a_{2,2} & \dots & a_{2,m} \\ \dots & \dots & \dots & \dots \\ a_{n,1} & a_{n,2} & \dots & a_{n,m} \end{bmatrix}$$

the matrix obtained by multiplying $\mathbf{A}$ and a scalar c , say $\mathbf{C} := c \cdot \mathbf{A}$, is a $n \times m$ matrix whose element are the elements of $\mathbf{A}$ multiplied by c , that is

$$\mathbf{C} = \begin{bmatrix} c \cdot a_{1,1} & c \cdot a_{1,2} & \dots & c \cdot a_{1,m} \\ c \cdot a_{2,1} & c \cdot a_{2,2} & \dots & c \cdot a_{2,m} \\ \dots & \dots & \dots & \dots \\ c \cdot a_{n,1} & c \cdot a_{n,2} & \dots & c \cdot a_{n,m} \end{bmatrix} = \begin{bmatrix} c_{1,1} & c_{1,2} & \dots & c_{1,m} \\ c_{2,1} & c_{2,2} & \dots & c_{2,m} \\ \dots & \dots & \dots & \dots \\ c_{n,1} & c_{n,2} & \dots & c_{n,m} \end{bmatrix}.$$

Example 71 Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 1 & 7 \\ 6 & \pi \\ \ln(64) & 0 \end{bmatrix}$$

and $c = \frac{1}{2}$. Then $\mathbf{C} = c \cdot \mathbf{A}$ is given by

$$\mathbf{C} = \begin{bmatrix} \frac{1}{2} \cdot 1 & \frac{1}{2} \cdot 7 \\ \frac{1}{2} \cdot 6 & \frac{1}{2} \cdot \pi \\ \frac{1}{2} \cdot \ln(64) & \frac{1}{2} \cdot 0 \end{bmatrix} = \begin{bmatrix} 0.5 & 3.5 \\ 3 & \frac{\pi}{2} \\ \ln(8) & 0 \end{bmatrix}$$

Properties of the multiplication by a scalar

1. $b \cdot (c \cdot \mathbf{A}) = (bc) \cdot \mathbf{A}$ (associative)
2. $c \cdot (\mathbf{A} \pm \mathbf{B}) = c \cdot \mathbf{A} \pm c \cdot \mathbf{B}$ (distributive 1)
3. $(b \pm c) \cdot \mathbf{A} = b \cdot \mathbf{A} \pm c \cdot \mathbf{A}$ (distributive 2)
4. $1 \cdot \mathbf{A} = \mathbf{A}$ (identity element with respect to the multiplication by a scalar)
5. $0 \cdot \mathbf{A} = \mathbf{O}$ (property of the zero)
6. $c \cdot \mathbf{O} = \mathbf{O}$ (property of the null matrix)

Hadamard product For two matrices $\mathbf{A}$ and $\mathbf{B}$ of the same dimension $m \times n$, the Hadamard product $\mathbf{A} \circ \mathbf{B}$ is a matrix of the same dimension as the operands, with elements given by $a_{i,j}b_{i,j}$. For matrices of different dimensions ($m \times n$ and $p \times q$, where $m \neq p$ or $n \neq q$) the Hadamard product is undefined.

Example 72 The Hadamard product for a 3×3 matrix $\mathbf{A}$ with a 3×3 matrix $\mathbf{B}$ is

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix} \circ \begin{bmatrix} b_{11} & b_{12} & b_{13} \\ b_{21} & b_{22} & b_{23} \\ b_{31} & b_{32} & b_{33} \end{bmatrix} = \begin{bmatrix} a_{11}b_{11} & a_{12}b_{12} & a_{13}b_{13} \\ a_{21}b_{21} & a_{22}b_{22} & a_{23}b_{23} \\ a_{31}b_{31} & a_{32}b_{32} & a_{33}b_{33} \end{bmatrix}.$$

Example 73 Consider the matrices

$$\mathbf{A} = \begin{bmatrix} 3 & 5 & 7 \\ 4 & 9 & 8 \end{bmatrix} \quad \mathbf{B} = \begin{bmatrix} 1 & 6 & 3 \\ 0 & 2 & 9 \end{bmatrix}.$$

The matrix $\mathbf{C} = \mathbf{A} \circ \mathbf{B}$ is then

$$\mathbf{C} = \begin{bmatrix} 3 & 30 & 21 \\ 0 & 18 & 72 \end{bmatrix}.$$

Properties of the Hadamard product

For $\mathbf{A}$, $\mathbf{B}$, and $\mathbf{C}$ matrices of the same size, and k a scalar.

1. $\mathbf{A} \circ \mathbf{B} = \mathbf{B} \circ \mathbf{A}$ (commutative property)
2. $\mathbf{A} \circ (\mathbf{B} \circ \mathbf{C}) = (\mathbf{A} \circ \mathbf{B}) \circ \mathbf{C}$ (associative property)
3. $\mathbf{A} \circ (\mathbf{B} + \mathbf{C}) = \mathbf{A} \circ \mathbf{B} + \mathbf{A} \circ \mathbf{C}$ (distributive property)
4. $(k\mathbf{A}) \circ \mathbf{B} = \mathbf{A} \circ (k\mathbf{B}) = k(\mathbf{A} \circ \mathbf{B})$
5. $\mathbf{A} \circ \mathbf{0} = \mathbf{0} \circ \mathbf{A} = \mathbf{0}$.

Matrix multiplication The above product should not be confused with matrix multiplication, defined in the following.

Definition 20 Given an $n \times p$ matrix $\mathbf{A}$ and $p \times m$ matrix $\mathbf{B}$, the product matrix

$$\underset{n \times m}{\mathbf{C}} = \underset{n \times p}{\mathbf{A}} \times \underset{p \times m}{\mathbf{B}}$$

is an $n \times m$ matrix defined as

$$\mathbf{C} = \begin{bmatrix} c_{1,1} & c_{1,2} & \cdots & c_{1,m} \\ c_{2,1} & c_{2,2} & \cdots & c_{2,m} \\ \cdots & \cdots & \cdots & \cdots \\ c_{n,1} & c_{n,2} & \cdots & c_{n,m} \end{bmatrix} \quad \text{with} \quad c_{i,j} = \sum_{k=1}^p a_{ik}b_{kj},$$

for $i = \{1, \dots, n\}$, $j = \{1, \dots, m\}$.

Despite the definition seems complicated, what it says is very simple: firstly, to multiply two matrices they have to have dimensions such that the product matrix exists (that is p must be the same).

Secondly, the dimension of the product matrix is $n \times m$.

Finally, the element i, j of the matrix $\mathbf{C}$ is given by the element-by-element product of the i th-row of $\mathbf{A}$ with the j th-column of $\mathbf{B}$ (which have the same number of elements because $\mathbf{A}$ and $\mathbf{B}$ agree in p). This is illustrated schematically as follows.

$$\begin{array}{c|c}
 & \begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{bmatrix} \\
 \hline
 \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \\ a_{31} & a_{32} \end{bmatrix} & \begin{bmatrix} a_{11}b_{11} + a_{12}b_{21} & a_{11}b_{12} + a_{12}b_{22} \\ a_{21}b_{11} + a_{22}b_{21} & a_{21}b_{12} + a_{22}b_{22} \\ a_{31}b_{11} + a_{32}b_{21} & a_{31}b_{12} + a_{32}b_{22} \end{bmatrix}
 \end{array}$$

Example 74 Consider the matrices

$$\mathbf{A} = \begin{bmatrix} 4 & 5 \\ 1 & 10 \\ 3 & 7 \end{bmatrix} ; \quad \mathbf{B} = \begin{bmatrix} 2 & 3 \\ 4 & 1 \end{bmatrix} .$$

Then $C = A \times B$ is given by

$$\begin{array}{c|c}
 & \mathbf{B} = \begin{bmatrix} 2 & 3 \\ 4 & 1 \end{bmatrix} \\
 \hline
 \mathbf{A} = \begin{bmatrix} 4 & 5 \\ 1 & 10 \\ 3 & 7 \end{bmatrix} & \begin{bmatrix} 8+20 & 12+5 \\ 2+40 & 3+10 \\ 6+28 & 9+7 \end{bmatrix} = \begin{bmatrix} 28 & 17 \\ 42 & 13 \\ 34 & 16 \end{bmatrix} = \mathbf{C}
 \end{array}$$

i.e.

$$\mathbf{C} = \begin{bmatrix} 28 & 17 \\ 42 & 13 \\ 34 & 16 \end{bmatrix} .$$

We can use the same logic to multiply a matrix and vector together.

Example 75 We go back once more to Example 60 and finally calculate the portfolio return. Let us remind ourselves that we have three assets, of which we know the returns over the past 5 trading days. We summarised this information in the matrix

$$\mathbf{R} = \begin{bmatrix} 0.005 & -0.001 & 0.006 & -0.023 & 0.002 \\ 0.006 & -0.012 & -0.011 & -0.004 & 0.007 \\ -0.015 & 0.009 & -0.007 & 0.008 & -0.004 \end{bmatrix} ,$$

whilst the weights are summarised in the vector

$$\mathbf{w} = \begin{pmatrix} 0.30 \\ 0.35 \\ 0.35 \end{pmatrix}.$$

Then the portfolio returns in each trading day are obtained via matrix multiplication, i.e.

$$\begin{aligned} r_{tot} &= \mathbf{w}^\top \mathbf{R} \\ &= \begin{pmatrix} 0.30 & 0.35 & 0.35 \end{pmatrix} \begin{bmatrix} 0.005 & -0.001 & 0.006 & -0.023 & 0.002 \\ 0.006 & -0.012 & -0.011 & -0.004 & 0.007 \\ -0.015 & 0.009 & -0.007 & 0.008 & -0.004 \end{bmatrix} \\ &= \begin{pmatrix} -0.00165 & -0.00135 & -0.0045 & -0.0055 & 0.00165 \end{pmatrix}. \end{aligned}$$

Properties of matrix multiplication

1. $\mathbf{A}(\mathbf{BC}) = (\mathbf{AB})\mathbf{C}$ (associative)
2. $\mathbf{A}(\mathbf{B} \pm \mathbf{C}) = \mathbf{AB} \pm \mathbf{AC}$ (distributive 1)
3. $(\mathbf{A} \pm \mathbf{B})\mathbf{C} = \mathbf{AC} \pm \mathbf{BC}$ (distributive 2)
4. $\mathbf{IA} = \mathbf{AI} = \mathbf{A}$ (identity element with respect to the matrix multiplication)
5. $\mathbf{OA} = \mathbf{OA} = \mathbf{O}$ (property of the null matrix)
6. Commutative law does not hold for matrix multiplication, i.e.

$$\mathbf{A} \times \mathbf{B} \neq \mathbf{B} \times \mathbf{A};$$

see the previous Exercise.

A final note. Multiplying matrices with pencil and paper may take long (and even simple mistakes can lead to very different results); luckily most software provide efficient algorithms to calculate the product of two matrices.

Matrix inversion Matrix inversion applies to square matrices only.

Definition 21 Given a n -dimensional square matrix $\mathbf{A}$, its inverse matrix $\mathbf{A}^{-1}$, if it exists, is the n -dimensional square matrix such that

$$\mathbf{AA}^{-1} = \mathbf{A}^{-1}\mathbf{A} = \mathbf{I}.$$

Also, the matrix $\mathbf{A}^{-1}$ is unique.

This is the equivalent of asking what is the reciprocal of a real number a . Of course this number is a^{-1} because

$$a \cdot a^{-1} = a^{-1} \cdot a = 1.$$

Example 76 Given the matrix

$$\mathbf{A} = \begin{bmatrix} 2 & 5 & 7 \\ 6 & 3 & 4 \\ 5 & -2 & -3 \end{bmatrix},$$

verify that

$$\mathbf{A}^{-1} = \begin{bmatrix} 1 & -1 & 1 \\ -38 & 41 & -34 \\ 27 & -29 & 24 \end{bmatrix},$$

is its inverse.

$$\begin{aligned} \mathbf{A}\mathbf{A}^{-1} &= \begin{bmatrix} 2 \cdot 1 + 5 \cdot (-38) + 7 \cdot 27 & 2 \cdot (-1) + 5 \cdot 41 + 7 \cdot (-29) & 2 \cdot 1 + 5 \cdot (-34) + 7 \cdot 24 \\ 6 \cdot 1 + 3 \cdot (-38) + 4 \cdot 27 & 6 \cdot (-1) + 3 \cdot 41 + 4 \cdot (-29) & 6 \cdot 1 + 3 \cdot (-34) + 4 \cdot 24 \\ 5 \cdot 1 + (-2) \cdot (-38) + (-3) \cdot 27 & 5 \cdot (-1) + (-2) \cdot 41 + (-3) \cdot (-29) & 5 \cdot 1 + (-2) \cdot (-34) + (-3) \cdot 24 \end{bmatrix} \\ &= \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \mathbf{I} \end{aligned}$$

$$\begin{aligned} \mathbf{A}^{-1}\mathbf{A} &= \begin{bmatrix} 1 \cdot 2 + (-1) \cdot 6 + 1 \cdot 5 & 1 \cdot 5 + (-1) \cdot 3 + 1 \cdot (-2) & 1 \cdot 7 + (-1) \cdot 4 + 1 \cdot (-3) \\ -38 \cdot 2 + 41 \cdot 6 + (-34) \cdot 5 & -38 \cdot 5 + 41 \cdot 3 + (-34) \cdot (-2) & -38 \cdot 7 + 41 \cdot 4 + (-34) \cdot (-3) \\ 27 \cdot 2 + (-29) \cdot 6 + 24 \cdot 5 & 27 \cdot 5 + (-29) \cdot 3 + 24 \cdot (-2) & 27 \cdot 7 + (-29) \cdot 4 + 24 \cdot (-3) \end{bmatrix} \\ &= \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \mathbf{I} \end{aligned}$$

The calculation of the reciprocal of a number is straightforward: $a^{-1} = 1 : a = \frac{1}{a}$. However, the answer in the matrix case is not as straightforward as for scalars because ‘matrix division’ is not defined. We will learn later on how to obtain the inverse matrix. But most importantly, we are very interested in understanding if the inverse matrix is unique and if exists, i.e. when we are allowed to invert.

The inverse matrix is unique, like the reciprocal of a real number is unique.

About the existence, for scalars the answer is straightforward: if $a \neq 0$, the reciprocal exists as it is defined. It would be tempting to think that if $\mathbf{A} \neq \mathbf{O}$, the inverse always exists. Unfortunately, the answer is a bit more complicated because there are matrices, different from the null matrix, which are not invertible. We require the knowledge of nonsingular matrix, which is explored in Section 2.2.

Other properties A useful list of other properties - which is recommended to keep in mind!

1. $(\mathbf{A}^\top)^\top = \mathbf{A}$
2. $(\mathbf{A} \pm \mathbf{B})^\top = \mathbf{A}^\top \pm \mathbf{B}^\top$
3. $(\mathbf{AB})^\top = \mathbf{B}^\top \mathbf{A}^\top$
4. $(c \cdot \mathbf{A})^\top = c \cdot \mathbf{A}^\top$
5. $(\mathbf{A}^{-1})^{-1} = \mathbf{A}$
6. $(c \cdot \mathbf{A})^{-1} = \frac{1}{c} \cdot \mathbf{A}^{-1}$
7. $(\mathbf{A}^\top)^{-1} = (\mathbf{A}^{-1})^\top$

2.2 Rank & nonsingular matrix

For the scope of assessing whether or not a matrix is invertible, solving linear systems, and studying minima and maxima of functions of more than one variable, we require the knowledge of two fundamental quantities: rank and determinant. The rank can be calculated for any matrix. Determinants can be calculated for square matrices only.

We begin with the rank.

Definition 22 Consider a $r \times c$ matrix $\mathbf{A}$. The rank of $\mathbf{A}$, denoted as $\text{rank}(\mathbf{A})$ is defined as

- (a) the maximum number of linearly independent column vectors in a matrix;
- (b) the maximum number of linearly independent row vectors in a matrix.

Statements (a) and (b) are equivalent. If no linear combination is present, then $\text{rank}(\mathbf{A}) = \min(r, c)$.

Example 77 Consider the matrices

$$\mathbf{A} = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}, \quad \mathbf{B} = \begin{bmatrix} 1 & 2 & 8 \\ 3 & 6 & 24 \end{bmatrix}, \quad \mathbf{C} = \begin{bmatrix} 1 & 2 \\ 3 & 4 \\ 7 & 13 \end{bmatrix}, \quad \mathbf{D} = \begin{bmatrix} 1 & 2 & 5 & \pi \\ 3 & 4 & 13 & \ln(3) \end{bmatrix}.$$

$\mathbf{A}$ is 2×2 , $\mathbf{B}$ is 2×3 , $\mathbf{C}$ is 3×2 and $\mathbf{D}$ is 2×4 . Then $\text{rank}(\mathbf{A}) = 2$ as all column/row vectors are linearly independent. Instead, $\text{rank}(\mathbf{B}) = 1$ as the second row is a linear combination of row 1 ($r_2 = 3r_1$). Matrices $\mathbf{C}$ and $\mathbf{D}$ have both rank equal to 2 as they do not contain any linear combination.

The knowledge of the rank of a matrix allows us to identify whether a matrix is invertible or not, i.e. if the inverse matrix exists. This follows from the concept of nonsingular matrix.

Definition 23 *A $n \times n$ matrix $\mathbf{A}$ such that $\text{rank}(\mathbf{A}) = n$ is called nonsingular. A $n \times n$ matrix $\mathbf{A}$ such that $\text{rank}(\mathbf{A}) < n$ is called singular.*

Note that this definition applies only to square matrices.

Example 78 The matrix

$$\mathbf{A} = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$$

from Example 77 is nonsingular. Instead, the matrix

$$\mathbf{B} = \begin{bmatrix} 1 & 2 \\ 3 & 6 \end{bmatrix}$$

is singular. Indeed we notice that there is a linear combination: column 2 is twice column 1. Consequently $\text{rank}(\mathbf{B}) = 1 < 2$.

Finally, a square matrix $\mathbf{A}$ is invertible, i.e. $\mathbf{A}^{-1}$ exists, if and only if $\mathbf{A}$ is nonsingular.

2.3 Determinant of a matrix

The determinant of a matrix $\mathbf{A}$, denoted as $|\mathbf{A}|$ or $\det(\mathbf{A})$, is a number associated to each square matrix, that provides a convenient criterion for determining whether $\mathbf{A}$ has full rank, or equivalently if it is nonsingular, or equivalently if it is invertible.

We are going to learn to calculate the determinant of a matrix one step at the time, beginning with a few simple cases of small size matrices. We will then take a small break to explore some special determinants, which will help us in a number of occasions in the future, and which ultimately will take us to a general rule for the calculation of the determinant of larger matrices via the co-factors rule.

We begin our journey with the determinant of small size matrices.

2.3.1 Determinant of a 1×1 matrix

This is a trivial case of a matrix of the form

$$\mathbf{A} = [a_{11}].$$

Then $|\mathbf{A}| = a_{11}$.

Example 79 The determinant of the matrix $\mathbf{A} = [-4]$ is

$$|\mathbf{A}| = -4.$$

2.3.2 Determinant of 2×2 matrix: the Leibniz formula

The Leibniz formula for the determinant of a 2×2 matrix

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix}$$

is

$$|\mathbf{A}| = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{12}a_{21}.$$

In other words, first multiply the terms on the main diagonal, then subtract the product of the elements on the secondary diagonal.

Example 80 Let

$$\mathbf{A} = \begin{bmatrix} 1 & 4 \\ 3 & 0 \end{bmatrix}.$$

Then

$$\begin{aligned} |\mathbf{A}| &= \begin{vmatrix} 1 & 4 \\ 3 & 0 \end{vmatrix} \\ &= 0 - 12 = -12. \end{aligned}$$

2.3.3 Determinant of a 3×3 matrix: Sarrus rule

We look at Sarrus rule for the determinant of a 3×3 matrix

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix},$$

as in spirit it is not too different from the Leibniz rule for the 2×2 matrix which we have seen above. We still operate along the diagonals starting from the main one. We can refer to the mnemonic scheme below.

1. The first step is to write out the first two columns of the matrix to the right of the third column, giving five columns in a row (additional elements reported in red).

$$\begin{array}{ccccc} a_{11} & a_{12} & a_{13} & a_{11} & a_{12} \\ a_{21} & a_{22} & a_{23} & a_{21} & a_{22} \\ a_{31} & a_{32} & a_{33} & a_{31} & a_{32} \end{array}$$

2. The second step is to add the products of the elements on the diagonals going from top to bottom (elements in blue).

$$\begin{array}{ccc|cc}
 a_{11} & a_{12} & a_{13} & a_{11} & a_{12} \\
 a_{21} & a_{22} & a_{23} & a_{21} & a_{22} \\
 a_{31} & a_{32} & a_{33} & a_{31} & a_{32} \\
 \hline
 a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} & & & &
 \end{array}$$

3. Final step: subtract the products of the diagonals going from bottom to top (elements in green).

$$\begin{array}{ccc|cc}
 a_{11} & a_{12} & a_{13} & a_{11} & a_{12} \\
 a_{21} & a_{22} & a_{23} & a_{21} & a_{22} \\
 a_{31} & a_{32} & a_{33} & a_{31} & a_{32} \\
 \hline
 a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{31}a_{22}a_{13} - a_{32}a_{23}a_{11} - a_{33}a_{21}a_{12} & & & &
 \end{array}$$

Then

$$\begin{aligned}
 |\mathbf{A}| &= \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} \\
 &= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{31}a_{22}a_{13} - a_{32}a_{23}a_{11} - a_{33}a_{21}a_{12} \quad (6)
 \end{aligned}$$

Example 81 The determinant of

$$\mathbf{A} = \begin{bmatrix} 1 & 4 & 7 \\ 3 & 0 & 5 \\ -1 & 9 & 11 \end{bmatrix}$$

is

$$\begin{aligned}
 |\mathbf{A}| &= \begin{vmatrix} 1 & 4 & 7 \\ 3 & 0 & 5 \\ -1 & 9 & 11 \end{vmatrix} \\
 &= \begin{vmatrix} 1 & 4 & 7 & | & 1 & 4 \\ 3 & 0 & 5 & | & 3 & 0 \\ -1 & 9 & 11 & | & -1 & 9 \end{vmatrix} \\
 &= 0 - 20 + 189 - (0 + 45 + 132) = -8
 \end{aligned}$$

Remark 13 This scheme for calculating the determinant of a 3×3 matrix does not carry over into higher dimensions.

Because of the previous remark, we need a more general rule, which is called the co-factor expansions. This though require the knowledge of some special types of determinants which we introduce next.

2.3.4 Minors and co-factors

There are some special determinants which we need to get to know, as they will be crucial for our calculations.

Definition 24 Let $\mathbf{A}$ be a square matrix of dimension n .

- (a) The minor of the entry in the i -th row and j -th column (also called the (i, j) minor) is the determinant of the sub-matrix formed by deleting the i -th row and j -th column. This number is often denoted $M_{i,j}$.
- (b) If $i = j$ then $M_{i,j} = M_{i,i}$ is called a principal minor.
- (c) A principal minor of order k is the determinant of the $k \times k$ sub-matrices formed by including elements from the principal diagonal. We denote it as $|\mathbf{A}_{k,ij}|$.
- (d) A leading principal minor of order k is the principal minor obtained from the sub-matrix whose row and column indices are $1, 2, \dots, k$. We denote it as $|\mathbf{A}_k|$.

Example 82 Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 1 & 4 & 7 \\ 3 & 0 & 5 \\ -1 & 9 & 11 \end{bmatrix}$$

The minor $M_{2,3}$ is the determinant of the matrix we obtain from $\mathbf{A}$ if we eliminate the 2nd row and the third column, i.e.

$$\mathbf{A} = \begin{bmatrix} 1 & 4 \\ -1 & 9 \end{bmatrix}$$

Thus by Leibniz rule

$$M_{2,3} = \begin{vmatrix} 1 & 4 \\ -1 & 9 \end{vmatrix} = 5.$$

We can obtain the principal minor $M_{3,3}$ in the same manner, by deleting from $\mathbf{A}$ both the third row and column, i.e.

$$\mathbf{A} = \begin{bmatrix} 1 & 4 \\ 3 & 0 \end{bmatrix}$$

so that by Leibniz rule we obtain (see Example 80)

$$M_{3,3} = \begin{vmatrix} 1 & 4 \\ 3 & 0 \end{vmatrix} = -12.$$

This is principal minor of order 2. Other principal minors of order 2 are $M_{1,1}$ and $M_{2,2}$. But we can also look at principal minors of order 1. These are

$$\begin{aligned} |\mathbf{A}_{1,11}| &= |a_{11}| = 1; \\ |\mathbf{A}_{1,22}| &= |a_{22}| = 0; \\ |\mathbf{A}_{1,33}| &= |a_{33}| = 11. \end{aligned}$$

The principal minor of order 3 is of course the determinant of the original matrix, i.e. -8 (verify this via Sarrus).

The leading principal minors are instead the determinants of the sub-matrices obtained from the ‘north-west’ corner of the matrix $\mathbf{A}$ proceeding along the main diagonal. Therefore, using the rules we have learnt so far, we obtain that

$$\begin{aligned} \text{the leading principal minor of order 1 is } |\mathbf{A}_1| &= |1| \\ \text{the leading principal minor of order 2 is } |\mathbf{A}_2| &= \begin{vmatrix} 1 & 4 \\ 3 & 0 \end{vmatrix} = -12 \\ \text{the leading principal minor of order 3 is } |\mathbf{A}_3| &= \begin{vmatrix} 1 & 4 & 7 \\ 3 & 0 & 5 \\ -1 & 9 & 11 \end{vmatrix} = -8 \end{aligned}$$

(see Example 81).

Minors will help us with the calculation of the determinant via the co-factors; but we will encounter minors once again when we study quadratic forms in Section 2.5.

We now look at the so-called co-factors.

Definition 25 Consider a $n \times n$ matrix $\mathbf{A} = [a_{ij}]$. Let M_{ij} denote the minor of a_{ij} . We define the co-factor of a_{ij} , denoted as A_{ij} , the signed minor

$$A_{ij} = (-1)^{i+j} M_{ij}.$$

Example 83 Reconsider Example 82. Then, the co-factor of a_{23} is

$$A_{23} = (-1)^{2+3}M_{23} = -1 \times 5 = -5.$$

Remark 14 The sign $(-1)^{i+j}$ of the co-factor form a chessboard pattern with +s on the main diagonal

$$\begin{bmatrix} + & - & + & \dots \\ - & + & - & \dots \\ + & - & + & \dots \\ \dots & & & \end{bmatrix}$$

2.3.5 Determinant of $n \times n$ matrix: co-factor expansion

First of all, let us revisit the 3×3 case for which

$$\begin{aligned} |\mathbf{A}| &= \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} \\ &= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{31}a_{22}a_{13} - a_{32}a_{23}a_{11} - a_{33}a_{21}a_{12} \end{aligned}$$

(see equation (6)). Note that we can rearrange the previous formula as follows

$$\begin{aligned} |\mathbf{A}| &= \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} \\ &= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{31}a_{22}a_{13} - a_{32}a_{23}a_{11} - a_{33}a_{21}a_{12} \\ &= a_{11}a_{22}a_{33} - a_{33}a_{21}a_{12} \\ &\quad + a_{12}a_{23}a_{31} - a_{32}a_{23}a_{11} \\ &\quad + a_{13}a_{21}a_{32} - a_{31}a_{22}a_{13} \\ &= a_{33} \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} - a_{23} \begin{vmatrix} a_{11} & a_{12} \\ a_{31} & a_{32} \end{vmatrix} + a_{13} \begin{vmatrix} a_{21} & a_{22} \\ a_{31} & a_{32} \end{vmatrix} \\ &= (-1)^{3+3}a_{33}M_{33} + (-1)^{2+3}a_{23}M_{23} + (-1)^{1+3}a_{13}M_{13}. \end{aligned}$$

This means that the determinant can be calculated as

$$|\mathbf{A}| = \sum_{i=1}^n (-1)^{i+k} a_{i\mathbf{k}} M_{i,\mathbf{k}} = \sum_{i=1}^n a_{i\mathbf{k}} A_{i,\mathbf{k}};$$

For the specific case of our example we chose $k = 3$, i.e. we fixed the third column. Equivalently, we could rewrite our determinant as

$$\begin{aligned}
 |\mathbf{A}| &= \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} \\
 &= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{31}a_{22}a_{13} - a_{32}a_{23}a_{11} - a_{33}a_{21}a_{12} \\
 &= a_{11} \begin{vmatrix} a_{22} & a_{23} \\ a_{32} & a_{33} \end{vmatrix} - a_{12} \begin{vmatrix} a_{21} & a_{23} \\ a_{31} & a_{33} \end{vmatrix} + a_{13} \begin{vmatrix} a_{21} & a_{22} \\ a_{31} & a_{32} \end{vmatrix} \\
 &= (-1)^{1+1}a_{11}M_{11} + (-1)^{1+2}a_{12}M_{12} + (-1)^{1+3}a_{13}M_{13} \\
 &= \sum_{k=1}^n (-1)^{i+k} a_{ik} M_{ik} = \sum_{k=1}^n a_{ik} A_{ik},
 \end{aligned}$$

where we fix $i = 1$.

So here is our co-factor expansion (or Laplace expansion/minor expansion formula), according to which for a generic $n \times n$ matrix $\mathbf{A}$ we have

$$|\mathbf{A}| = \sum_{k=1}^n a_{ik} A_{ik} \quad (7)$$

$$= \sum_{i=1}^n a_{ik} A_{i,k}. \quad (8)$$

The above says that you can operate in two alternative ways.

You can fix a row i and compute all the co-factors that can be obtained on this row; then you multiply each co-factor by the corresponding element of the matrix along the chosen row, and finally you sum.

Alternatively, you can fix column k and proceed in the same way for each element on that column. We illustrate the method with an example.

Example 84 1. Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 4 & 3 & 2 & 2 \\ 0 & 1 & -3 & 3 \\ 0 & -1 & 3 & 3 \\ 0 & 3 & 1 & 1 \end{bmatrix}.$$

We fix by convenience the first column. Then the possible minors are

$$\begin{vmatrix} 1 & -3 & 3 \\ -1 & 3 & 3 \\ 3 & 1 & 1 \end{vmatrix} \quad \begin{vmatrix} 3 & 2 & 2 \\ -1 & 3 & 3 \\ 3 & 1 & 1 \end{vmatrix} \quad \begin{vmatrix} 3 & 2 & 2 \\ 1 & -3 & 3 \\ 3 & 1 & 1 \end{vmatrix} \quad \begin{vmatrix} 3 & 2 & 2 \\ 1 & -3 & 3 \\ -1 & 3 & 3 \end{vmatrix}$$

Then, our determinant is

$$\begin{aligned}
 |\mathbf{A}| &= (-1)^{1+1}a_{1,1} \begin{vmatrix} 1 & -3 & 3 \\ -1 & 3 & 3 \\ 3 & 1 & 1 \end{vmatrix} + (-1)^{2+1}a_{2,1} \begin{vmatrix} 3 & 2 & 2 \\ -1 & 3 & 3 \\ 3 & 1 & 1 \end{vmatrix} \\
 &\quad + (-1)^{3+1}a_{3,1} \begin{vmatrix} 3 & 2 & 2 \\ 1 & -3 & 3 \\ 3 & 1 & 1 \end{vmatrix} + (-1)^{4+1}a_{4,1} \begin{vmatrix} 3 & 2 & 2 \\ 1 & -3 & 3 \\ -1 & 3 & 3 \end{vmatrix} \\
 &= 4 \begin{vmatrix} 1 & -3 & 3 \\ -1 & 3 & 3 \\ 3 & 1 & 1 \end{vmatrix}
 \end{aligned}$$

as all the other elements of the first column are zero (this is why in this example choosing the first column is convenient!). Now, all we have to do is to use Sarrus rule to compute the determinant of the remaining 3×3 matrix. You can check that

$$|\mathbf{A}| = 4 \times (-60) = -240.$$

2. Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 4 & 3 & 2 & 2 \\ 0 & 1 & 0 & -2 \\ 1 & -1 & 3 & 3 \\ 2 & 3 & 1 & 1 \end{bmatrix}.$$

We can fix by convenience the second row, as it contains some zero elements - and we

know this will simplify our calculations. Then, our determinant is

$$\begin{aligned}
 |\mathbf{A}| &= (-1)^{2+1}a_{2,1} \begin{vmatrix} 3 & 2 & 2 \\ -1 & 3 & 3 \\ 3 & 1 & 1 \end{vmatrix} + (-1)^{2+2}a_{2,2} \begin{vmatrix} 4 & 2 & 2 \\ 1 & 3 & 3 \\ 2 & 1 & 1 \end{vmatrix} \\
 &\quad + (-1)^{2+3}a_{2,3} \begin{vmatrix} 4 & 3 & 2 \\ 1 & -1 & 3 \\ 2 & 3 & 1 \end{vmatrix} + (-1)^{2+4}a_{2,4} \begin{vmatrix} 4 & 3 & 2 \\ 1 & -1 & 3 \\ 2 & 3 & 1 \end{vmatrix} \\
 &= (-1)^{2+2}1 \begin{vmatrix} 4 & 2 & 2 \\ 1 & 3 & 3 \\ 2 & 1 & 1 \end{vmatrix} + (-1)^{2+4}(-2) \begin{vmatrix} 4 & 3 & 2 \\ 1 & -1 & 3 \\ 2 & 3 & 1 \end{vmatrix} \\
 &= 0 + (-2)(-15) \\
 &= 30.
 \end{aligned}$$

Therefore

$$|\mathbf{A}| = 30.$$

2.3.6 Properties of the determinant

The determinant has a number of properties that can help quite substantially with calculations (although beyond dimension 5 we are much better off using dedicated software packages like Matlab - once you have input the matrix $\mathbf{A}$, you can obtain the determinant by typing the command `det(A)`).

We list them in the following (without proof), and then we discuss their use and some very important consequences of these properties.

1. Let $\mathbf{B}$ be obtained from $\mathbf{A}$ by exchanging two rows (or columns). Then

$$|\mathbf{B}| = -|\mathbf{A}|.$$

Example 85 Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 2 & 1 & 3 \\ 2 & 4 & 1 \\ 3 & 2 & 7 \end{bmatrix}$$

with determinant 17. Build the matrix $\mathbf{B}$ by inverting rows 2 and 3, i.e.

$$\mathbf{A} = \begin{bmatrix} 2 & 1 & 3 \\ 3 & 2 & 7 \\ 2 & 4 & 1 \end{bmatrix}$$

then its determinant is -17.

2. Let $\mathbf{B}$ be obtained from $\mathbf{A}$ by multiplying some row (or some column) by the scalar c . Then

$$|\mathbf{B}| = c|\mathbf{A}|.$$

Example 86 Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 2 & 1 & 3 \\ 2 & 4 & 1 \\ 3 & 2 & 7 \end{bmatrix}$$

with determinant 17. Build the matrix $\mathbf{B}$ by multiplying row 1 by 2, i.e.

$$\mathbf{B} = \begin{bmatrix} 4 & 2 & 6 \\ 2 & 4 & 1 \\ 3 & 2 & 7 \end{bmatrix}$$

then its determinant is 34 (2×17).

An immediate consequence of this property is that if $\mathbf{A}$ is a $n \times n$ matrix, then the determinant of $c\mathbf{A}$ is $c^n|\mathbf{A}|$.

3. Let $\mathbf{B}$ be obtained from $\mathbf{A}$ by adding a multiple of one row (or column) to another. Then

$$|\mathbf{B}| = |\mathbf{A}|.$$

Example 87 Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 2 & 1 & 3 \\ 2 & 4 & 1 \\ 3 & 2 & 7 \end{bmatrix}$$

with determinant 17. Build the matrix $\mathbf{B}$ by adding row 1 to row 2 of $\mathbf{A}$, i.e.

$$\mathbf{B} = \begin{bmatrix} 2 & 1 & 3 \\ 2+2 & 1+4 & 3+1 \\ 3 & 2 & 7 \end{bmatrix} = \begin{bmatrix} 2 & 1 & 3 \\ 4 & 5 & 4 \\ 3 & 2 & 7 \end{bmatrix}$$

then its determinant is still 17.

4. Let $\mathbf{A}$ and $\mathbf{B}$ be $n \times n$ matrices. Then

$$|\mathbf{A} \times \mathbf{B}| = |\mathbf{A}||\mathbf{B}|.$$

Example 88 Consider the matrices

$$\mathbf{A} = \begin{bmatrix} 2 & 1 & 3 \\ 2 & 4 & 1 \\ 3 & 2 & 7 \end{bmatrix}$$

with determinant 17, and

$$\mathbf{B} = \begin{bmatrix} 4 & 3 & 2 \\ 1 & -1 & 3 \\ 2 & 3 & 1 \end{bmatrix}$$

with determinant -15. Then

$$\mathbf{M} = \mathbf{A} \times \mathbf{B} = \begin{bmatrix} 26 & 23 & 38 \\ 7 & 2 & 20 \\ 19 & 19 & 25 \end{bmatrix}$$

and its determinant is -255 which is indeed $17 \times (-15)$.

5. $|\mathbf{A}^\top| = |\mathbf{A}|$ (it follows from the co-factor expansion).

6. Last

$$|\mathbf{I}_n| = 1.$$

Now, let us have a look at the implications of these properties.

Firstly. Consider a matrix in which two rows (or columns) are identical, for example

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{11} & a_{12} & a_{13} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}.$$

Then, we could build a new matrix $\mathbf{B}$ by exchanging row 1 and row 2 of $\mathbf{A}$. But then, Property 1 implies

$$|\mathbf{A}| = -|\mathbf{B}|;$$

on the other hand, by construction $\mathbf{A} = \mathbf{B}$ and therefore

$$|\mathbf{A}| = |\mathbf{B}|.$$

Both statements are true but contradictory, unless

$$|\mathbf{A}| = 0.$$

Therefore, if a matrix $\mathbf{A}$ has two identical rows (or columns) then

$$|\mathbf{A}| = 0.$$

Example 89 The matrix

$$\begin{bmatrix} 4 & 2 & 2 \\ 1 & 3 & 3 \\ 2 & 1 & 1 \end{bmatrix}$$

has zero determinant - the second and third columns are identical.

Secondly, let us consider a matrix with one row (or column) equal to a linear combination of other rows (resp. columns), for example

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ v_1 a_{11} + v_2 a_{31} & v_1 a_{12} + v_2 a_{32} & v_1 a_{13} + v_2 a_{33} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}.$$

By Property 3 we know that

$$|\mathbf{A}| = |\mathbf{B}|,$$

for

$$\mathbf{B} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ v_1 a_{11} & v_1 a_{12} & v_1 a_{13} \\ a_{31} & a_{32} & a_{33} \end{bmatrix},$$

for example. But, Property 2 tells us that

$$|\mathbf{B}| = v_1 |\mathbf{C}|,$$

for

$$\mathbf{C} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{11} & a_{12} & a_{13} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}.$$

We have shown above though that $|\mathbf{C}| = 0$ as two rows are identical. Therefore

$$|\mathbf{A}| = 0.$$

In other words, the determinant of a matrix is zero any time the matrix contains a linear combination of rows (or columns).

Example 90 The matrix

$$\begin{bmatrix} 1 & 3 & 9 & 2 \\ 5 & 8 & 4 & 3 \\ 12 & 19 & 9 & 14 \\ 2 & 3 & 1 & 8 \end{bmatrix}$$

has zero determinant: the third row is a linear combination of row 2 and row 4. Specifically $R_3 = 2R_2 + R_4$.

This property of the determinant has a number of very important implications. By definition of rank of a matrix, we know that a matrix with zero determinant has not full rank. If it does not have full rank, then it is singular. If it is singular, then it is not invertible. Stated differently, consider square matrix $\mathbf{A}$ of size $n \times n$. Then the following statements are equivalent.

$$\text{rank}(\mathbf{A}) = n \Leftrightarrow |\mathbf{A}| \neq 0 \Leftrightarrow \mathbf{A} \text{ nonsingular} \Leftrightarrow \mathbf{A} \text{ invertible}$$

$$\text{rank}(\mathbf{A}) < n \Leftrightarrow |\mathbf{A}| = 0 \Leftrightarrow \mathbf{A} \text{ singular} \Leftrightarrow \mathbf{A} \text{ non invertible.}$$

A further implication of the above is that the determinant of a matrix is zero if all elements of one row (or column) are zeros. Indeed, this case can be treated as a linear combination with weights 1 and -1 operated on the same row (column), i.e.

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ 0 & 0 & 0 \\ a_{31} & a_{32} & a_{33} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} - a_{21} & a_{22} - a_{22} & a_{23} - a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}.$$

Finally, we can use Property 4 to obtain the determinant of the inverse matrix. Indeed, as by definition $\mathbf{A}^{-1}$ is such that

$$\mathbf{A}\mathbf{A}^{-1} = \mathbf{I},$$

then

$$|\mathbf{A}\mathbf{A}^{-1}| = |\mathbf{I}|,$$

and

$$|\mathbf{A}\mathbf{A}^{-1}| = |\mathbf{A}||\mathbf{A}^{-1}|.$$

Further, by Property 6

$$|\mathbf{A}\mathbf{A}^{-1}| = |\mathbf{A}||\mathbf{A}^{-1}| = 1,$$

i.e.

$$|\mathbf{A}^{-1}| = \frac{1}{|\mathbf{A}|}.$$

2.3.7 Adjoint Matrix and the Inverse Matrix

The classical adjoint matrix (or adjugate matrix) is the matrix

$$\text{adj} \mathbf{A} = [A_{ij}]^{\top},$$

i.e. the transpose of the matrix containing all the co-factors of $\mathbf{A}$.

Example 91 Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 2 & 3 & -4 \\ 0 & -4 & 2 \\ 1 & -1 & 5 \end{bmatrix}.$$

Then, the co-factors are

$$\begin{aligned} A_{11} &= (-1)^{1+1} \begin{vmatrix} -4 & 2 \\ -1 & 5 \end{vmatrix} & A_{12} &= (-1)^{1+2} \begin{vmatrix} 0 & 2 \\ 1 & 5 \end{vmatrix} & A_{13} &= (-1)^{1+3} \begin{vmatrix} 0 & -4 \\ 1 & -1 \end{vmatrix} \\ A_{21} &= (-1)^{2+1} \begin{vmatrix} 3 & -4 \\ -1 & 5 \end{vmatrix} & A_{22} &= (-1)^{2+2} \begin{vmatrix} 2 & -4 \\ 1 & 5 \end{vmatrix} & A_{23} &= (-1)^{2+3} \begin{vmatrix} 2 & 3 \\ 1 & -1 \end{vmatrix} \\ A_{31} &= (-1)^{3+1} \begin{vmatrix} 3 & -4 \\ -4 & 2 \end{vmatrix} & A_{32} &= (-1)^{3+2} \begin{vmatrix} 2 & -4 \\ 0 & 2 \end{vmatrix} & A_{33} &= (-1)^{3+3} \begin{vmatrix} 2 & 3 \\ 0 & -4 \end{vmatrix} \end{aligned}.$$

This implies that

$$[A_{ij}] = \begin{bmatrix} -18 & 2 & 4 \\ -11 & 14 & 5 \\ -10 & -4 & -8 \end{bmatrix},$$

and consequently the classical adjoint matrix is

$$\text{adj } \mathbf{A} = \begin{bmatrix} -18 & -11 & -10 \\ 2 & 14 & -4 \\ 4 & 5 & -8 \end{bmatrix}.$$

Now, we have all we need to obtain the inverse matrix.

Theorem 26 *Let $\mathbf{A}$ be a $n \times n$ matrix. Then*

$$\mathbf{A}(\text{adj } \mathbf{A}) = |\mathbf{A}|I_n = (\text{adj } \mathbf{A})\mathbf{A}.$$

Consequently, if $|\mathbf{A}| \neq 0$,

$$\mathbf{A}^{-1} = \frac{1}{|\mathbf{A}|} \text{adj } \mathbf{A}.$$

The proof is an application of the product rule, the co-factor rule and the properties of the determinant. So, we look at it to gain practice and understanding.

Proof. By definition of the product between two matrices (see Definition 20) we have

$$\mathbf{B} = \mathbf{A}(\text{adj } \mathbf{A}) = [b_{ij}], \quad b_{ij} = \sum_{j=1}^n a_{kj} A_{ij}, \quad k = 1, \dots, n,$$

where the last term follows from the definition of adjoint matrix. We note that

$$b_{ij} = \sum_{j=1}^n a_{kj} A_{ij} = |\mathbf{A}| \quad \text{if } k = i$$

by co-factors rule. However, if $k \neq i$, $b_{ij} = 0$. To see this, consider the matrix $\mathbf{A}' = [a'_{ij}]$ obtained from $\mathbf{A}$ by replacing the i -th row of $\mathbf{A}$ by its k -th row. Then, by properties of the determinant $|\mathbf{A}'| = 0$ (the matrix has two identical rows). On the other hand, by construction $a'_{ij} = a'_{kj} = a_{kj}$, and $A'_{ij} = A_{ij}$ as both are obtained eliminating the i -th row - which is the only row by which the two matrices differ. Therefore

$$|\mathbf{A}'| = \sum_{j=1}^n a'_{ij} A'_{ij} = \sum_{j=1}^n a_{kj} A_{ij} = 0.$$

Consequently

$$\mathbf{B} = \mathbf{A}(\text{adj } \mathbf{A}) = [b_{ij}], \quad b_{ij} = \sum_{j=1}^n a_{kj} A_{ij} = \begin{cases} |\mathbf{A}| & \text{if } k = i \\ 0 & \text{if } k \neq i \end{cases},$$

that is

$$\mathbf{B} = \mathbf{A}(\text{adj } \mathbf{A}) = |\mathbf{A}|I_n.$$

A similar argument shows that $(\text{adj } \mathbf{A})\mathbf{A} = |\mathbf{A}|I_n$.

Finally,

$$\mathbf{A}(\text{adj } \mathbf{A}) = |\mathbf{A}|I_n \Rightarrow I_n = \frac{1}{|\mathbf{A}|} \mathbf{A}(\text{adj } \mathbf{A});$$

(pre-)multiply both side of the above by the inverse of $\mathbf{A}$, to obtain

$$\begin{aligned}
 I_n &= \frac{1}{|\mathbf{A}|} \mathbf{A}(\text{adj } \mathbf{A}) \\
 \mathbf{A}^{-1} I_n &= \frac{1}{|\mathbf{A}|} \mathbf{A}^{-1} \mathbf{A}(\text{adj } \mathbf{A}) \\
 \mathbf{A}^{-1} &= \frac{1}{|\mathbf{A}|} \mathbf{A}^{-1} \mathbf{A}(\text{adj } \mathbf{A}) \quad \text{by def. of identity matrix} \\
 \mathbf{A}^{-1} &= \frac{1}{|\mathbf{A}|} I_n(\text{adj } \mathbf{A}) \quad \text{by def. of inverse matrix} \\
 \mathbf{A}^{-1} &= \frac{1}{|\mathbf{A}|} \text{adj } \mathbf{A} \quad \text{by def. of identity matrix,}
 \end{aligned}$$

as required. ■

Example 92 Continuing from Example 91. We can calculate the determinant of the given matrix using the co-factor rule fixing the first column

$$|\mathbf{A}| = 2(-18) + 0 + 1(-10) = -46$$

(you can also verify this result using the Sarrus rule). Then, the inverse matrix is

$$\mathbf{A}^{-1} = \frac{1}{|\mathbf{A}|} \text{adj } \mathbf{A} = -\frac{1}{46} \begin{bmatrix} -18 & -11 & -10 \\ 2 & 14 & -4 \\ 4 & 5 & -8 \end{bmatrix} = \begin{bmatrix} 9/23 & 11/46 & 5/23 \\ -1/23 & -7/23 & 2/23 \\ -2/23 & -5/46 & 4/23 \end{bmatrix}.$$

2.3.8 What is a determinant? A geometric interpretation

We start in dimension 2, i.e. two vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^2$ of the form

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}, \quad \mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \end{bmatrix}.$$

and consider the parallelogram $\mathcal{P}$ they span.

To understand this geometrically, we draw the arrow from the origin to the head of the vectors $\mathbf{x}, \mathbf{y}$, referring to Figure 24 - left hand side panel, the points X and Y .

Then we consider the vector $\mathbf{x} + \mathbf{y} = (x_1 + y_1, x_2 + y_2)$, and again we draw the corresponding arrow from the origin to the corresponding point $Z = X + Y$.

If we join the points we obtain a parallelogram with vertices in the origin, X, Y and Z . The vector of the sum is the main diagonal of the parallelogram. This is called the parallelogram law for vector addition. The symmetry of the parallelogram illustrates the commutative law $\mathbf{x} + \mathbf{y} = \mathbf{y} + \mathbf{x}$.

Now. We want to express the area of $\mathcal{P}$ in terms of the coordinates of $\mathbf{x}$ and $\mathbf{y}$. We achieve this, by noting that this area can be obtained from the area of the rectangle with vertices in the origin, $(x_1 + y_1, 0)$, $(x_1 + y_1, x_2 + y_2)$ and $(0, x_2 + y_2)$ - referring to the right hand side panel of

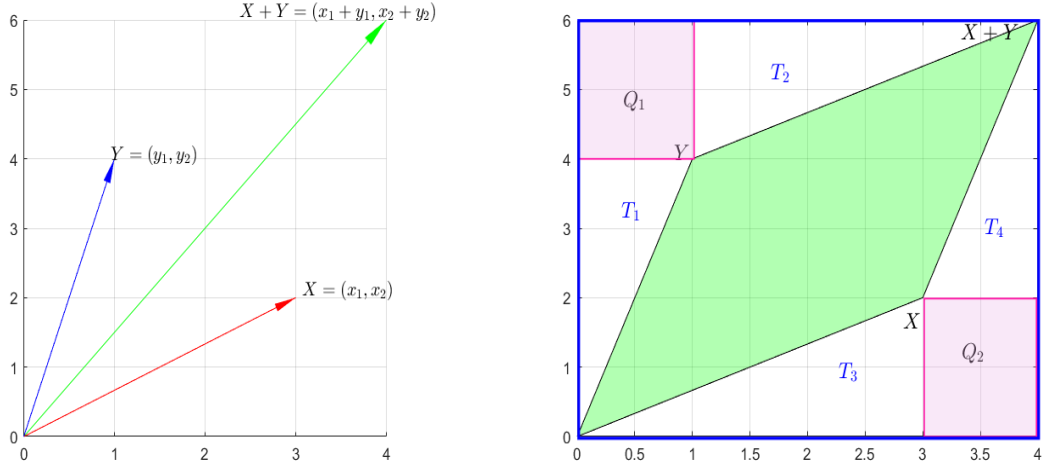


Figure 24: Determinant of a 2×2 matrix and the area of a parallelogram.

Figure 24, this rectangle, call it R , has a perimeter delineated in blue - by subtracting the areas of the two rectangles Q_1 and Q_2 , and the areas of the four triangles T_1, T_2, T_3 and T_4 .

From elementary geometry, we know

$$\begin{aligned} R &= (x_1 + y_1)(x_2 + y_2) \\ Q_1 &= x_2 y_1 = Q_2 \\ T_1 &= \frac{1}{2} y_1 y_2 = T_4 \\ T_2 &= \frac{1}{2} x_1 x_2 = T_3. \end{aligned}$$

Then

$$\begin{aligned} \text{area}(\mathcal{P}) &= (x_1 + y_1)(x_2 + y_2) - 2x_2 y_1 - y_1 y_2 - x_1 x_2 \\ &= x_1 y_2 - x_2 y_1, \end{aligned}$$

which we notice to be the determinant of the 2×2 matrix with row vectors $\mathbf{x}$ and $\mathbf{y}$,

$$\begin{bmatrix} x_1 & y_1 \\ x_2 & y_2 \end{bmatrix}.$$

In the light of the above, what does it mean that if there is a linear combination, then the determinant is zero? Let us have a look at the plot in the left hand panel of Figure 25. We start from the vectors $\mathbf{x}$ - corresponding to the point $X = (x_1, x_2)$ (first red arrow from the right) and $\mathbf{y}$ - corresponding to the point $Y = (y_1, y_2)$ (the first arrow from the left). Suppose you keep the X point fixed and rotate the vector of Y clockwise, until it lies exactly on the straight line starting from the origin and passing through X . The resulting vector is indicated by the blue arrow (first from the right). But this vector is a linear combination of the vector $\mathbf{x}$ (you can check that this line has equation $y = mx$ for $m = x_2/x_1$).

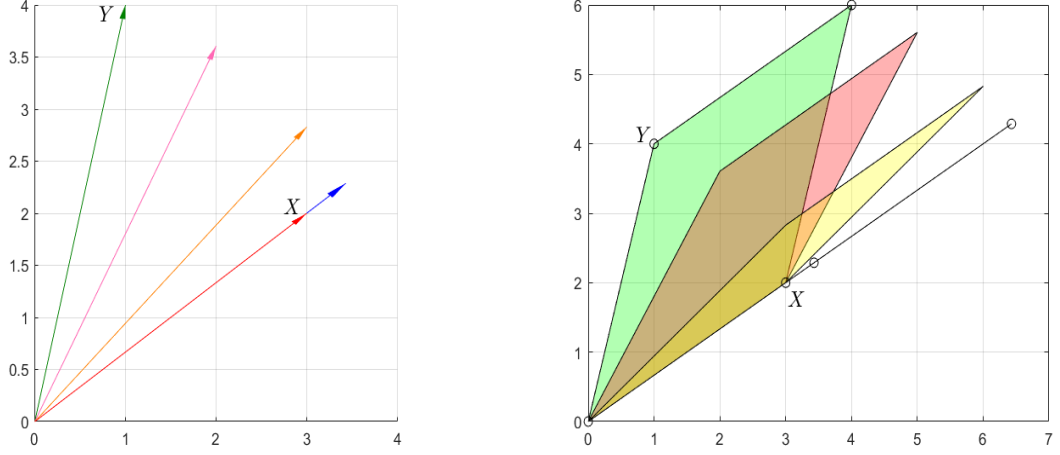


Figure 25: Linear combinations: the area of a parallelogram and the determinant.

What happens to the area of the parallelograms spanned by the vectors $\mathbf{x}$ and $\mathbf{y}$ as $\mathbf{y}$ rotates towards the blue arrow? The area of these parallelograms shrinks until it collapses to zero as shown in Figure 25, right hand panel. In other words, we lose one dimension.

Let us have a look at the corresponding matrix, then. This is given by the coordinates of the two points on the line, i.e.

$$\begin{bmatrix} x_1 & y_1 \\ x_2 & my_1 \end{bmatrix}.$$

Bearing in mind that $x_2 = mx_1$, then

$$\begin{vmatrix} x_1 & y_1 \\ mx_1 & my_1 \end{vmatrix} = 0.$$

Therefore, the area of $\mathcal{P}$ is non-zero as long as $\mathbf{x}$ and $\mathbf{y}$ are linearly independent (not collinear).

Similarly, moving to three dimensions, the determinant will be the volume of the parallelepiped spanned by three vectors $\mathbf{x}, \mathbf{y}$, and $\mathbf{z} \in \mathbb{R}^3$, and in higher dimension it is still the volume of the resulting parallelotope (i.e. parallelepiped in higher dimensions).

2.4 Linear Systems

Consider a system of linear equations of the form

$$\begin{aligned} a_{11}x_1 + a_{12}x_2 + \cdots + a_{1n}x_n &= b_1 \\ a_{21}x_1 + a_{22}x_2 + \cdots + a_{2n}x_n &= b_2 \\ &\vdots \\ a_{m1}x_1 + a_{m2}x_2 + \cdots + a_{mn}x_n &= b_m, \end{aligned}$$

where a_{ij}, b_i are real numbers, to be solved simultaneously.

This can be written in a compact way as

$$\mathbf{Ax} = \mathbf{b} \quad (9)$$

where $\mathbf{A}$ is a $(m \times n)$ matrix with elements a_{ij} , $\mathbf{x}$ is the n -vector of unknowns x_j and $\mathbf{b}$ is a m -vector with elements b_i . The compact form in (9) is known as a linear system.

The problem about solving the given system can be broken down in two steps

1. Existence and number of solutions
2. Calculation of the solution(s).

2.4.1 Existence and number of solutions

The given linear system has zero, one or infinitely many solutions depending on the matrix $\mathbf{A}$ and the vector $\mathbf{b}$. Indeed, we observe that by properties of the inverse matrix, the system (9) can be rewritten as

$$\mathbf{A}^{-1}\mathbf{Ax} = \mathbf{A}^{-1}\mathbf{b};$$

in other words, the solution to the linear system is

$$\mathbf{x} = \mathbf{A}^{-1}\mathbf{b}.$$

The above is telling us under which minimal conditions the solution to the system exists: the matrix $\mathbf{A}$ has to be nonsingular, i.e. square and invertible, i.e. with full rank.

On the other hand, let $\mathbf{a}_{\cdot,j}$ be the $m \times 1$ column vectors forming the matrix $\mathbf{A}$. Then, the system (9) can be rewritten as

$$x_1\mathbf{a}_{\cdot,1} + x_2\mathbf{a}_{\cdot,2} + \cdots + x_n\mathbf{a}_{\cdot,n} = \mathbf{b}.$$

The above shows that in order to find a solution, we must be able to express the vector $\mathbf{b}$ as a linear combination of the columns vectors of $\mathbf{A}$ with weights $x_1, \dots, x_n$. If this is the case, the augmented matrix $\mathbf{A}|\mathbf{b}$, which is obtained by adding the vector $\mathbf{b}$ to the matrix $\mathbf{A}$ after its last column, cannot have full rank as $\mathbf{b}$ is a linear combination, and therefore ‘redundant’. Hence, $\text{rank}(\mathbf{A}|\mathbf{b}) = \text{rank}(\mathbf{A})$.

These results are summarized in the following.

Theorem 27 (Rouché-Capelli) *The system $\mathbf{Ax} = \mathbf{b}$ admits solutions, i.e. is consistent, if and only if $\text{rank}(\mathbf{A}) = \text{rank}(\mathbf{A}|\mathbf{b})$. Otherwise, the system is inconsistent. Moreover, if the system is consistent, it admits $\infty^{n-\text{rank}(\mathbf{A})}$ solutions (i.e. infinitely many solution with $n - \text{rank}(\mathbf{A})$ degrees of freedom), where n is the number of unknowns.*

By degree of freedom we mean the number of variables which can be freely chosen, while the remaining variables are determined once the free ones have been fixed.

Theorem 27 is a very useful tool as it allows us to determine the number of solutions to a linear system, without necessarily solving it. This can help us to avoid spending time on solving systems which cannot be solved. It also helps us to search for all the required solutions in the case the system is consistent.

Example 93 1. Consider the linear system

$$\begin{aligned}x_1 + 2x_2 &= 4 \\ 3x_1 + x_2 &= 2.\end{aligned}$$

In this case

$$\mathbf{A} = \begin{bmatrix} 1 & 2 \\ 3 & 1 \end{bmatrix}$$

which has full rank as the matrix determinant is different from zero (and therefore all columns/rows are linearly independent). This also implies that the augmented matrix

$$\mathbf{A}|\mathbf{b} = \left[\begin{array}{cc|c} 1 & 2 & 4 \\ 3 & 1 & 2 \end{array} \right]$$

has rank 2 (the higher order minor is in fact 2). Consequently the solution exists. As $n = 2$, the solution is also unique. In fact the matrix $\mathbf{A}$ is square and invertible.

2. Consider the system of equations

$$\begin{aligned}x + y + 2z &= 3, \\ x + y + z &= 1, \\ 2x + 2y + 2z &= 2.\end{aligned}$$

The coefficient matrix is

$$\mathbf{A} = \begin{bmatrix} 1 & 1 & 2 \\ 1 & 1 & 1 \\ 2 & 2 & 2 \end{bmatrix},$$

and the augmented matrix is

$$\mathbf{A}|\mathbf{b} = \left[\begin{array}{ccc|c} 1 & 1 & 2 & 3 \\ 1 & 1 & 1 & 1 \\ 2 & 2 & 2 & 2 \end{array} \right].$$

Since both of these have the same rank, namely 2 (note that in both matrices $r_3 = 2r_2$), there exists at least one solution. However, the number of unknowns is 3 (x, y, z), therefore there are infinitely many solutions with 1 degree of freedom. For example, you can fix x equal to an arbitrary value, say $\bar{x}$, and solve the system in function of this value.

3. In contrast, consider the system

$$\begin{aligned}x + y + 2z &= 3, \\ x + y + z &= 1, \\ 2x + 2y + 2z &= 5.\end{aligned}$$

The coefficient matrix is

$$\mathbf{A} = \begin{bmatrix} 1 & 1 & 2 \\ 1 & 1 & 1 \\ 2 & 2 & 2 \end{bmatrix},$$

and the augmented matrix is

$$\mathbf{A}|\mathbf{b} = \left[\begin{array}{ccc|c} 1 & 1 & 2 & 3 \\ 1 & 1 & 1 & 1 \\ 2 & 2 & 2 & 5 \end{array} \right].$$

In this example the coefficient matrix has rank 2 (note that $r_3 = 2r_2$), while the augmented matrix has rank 3; so this system of equations has no solution, i.e. it is inconsistent.

2.4.2 Calculation of the solution

If the system is feasible and the solution is unique, we can find this solution using the inverse matrix. Indeed, as stated above, if $\mathbf{Ax} = \mathbf{b}$, then

$$\mathbf{x} = \mathbf{A}^{-1}\mathbf{b}$$

if $|\mathbf{A}| \neq 0$.

Example 94 Let us consider the linear systems in Example 93.

1. We know already that the system

$$\begin{aligned} x_1 + 2x_2 &= 4 \\ 3x_1 + x_2 &= 2. \end{aligned}$$

has unique solution as

$$|\mathbf{A}| = \begin{vmatrix} 1 & 2 \\ 3 & 1 \end{vmatrix} = -5 \neq 0.$$

Then

$$\mathbf{A}^{-1} = \frac{1}{-5} \begin{bmatrix} 1 & -2 \\ -3 & 1 \end{bmatrix} = \begin{bmatrix} -1/5 & 2/5 \\ 3/5 & -1/5 \end{bmatrix}.$$

Consequently

$$\begin{aligned} \mathbf{x} &= \mathbf{A}^{-1}\mathbf{b} \\ &= \begin{bmatrix} -1/5 & 2/5 \\ 3/5 & -1/5 \end{bmatrix} \begin{pmatrix} 4 \\ 2 \end{pmatrix} = \begin{pmatrix} 0 \\ 2 \end{pmatrix}. \end{aligned}$$

2. From the previous analysis, the system

$$\begin{aligned}x + y + 2z &= 3, \\x + y + z &= 1, \\2x + 2y + 2z &= 2\end{aligned}$$

has infinitely many solutions with 1 degree of freedom, as the third equation is redundant. Thus, we solve

$$\begin{aligned}x + y + 2z &= 3, \\x + y + z &= 1.\end{aligned}$$

Say that we fix $x = \bar{x}$. Then, the above can be rewritten as

$$\begin{aligned}y + 2z &= 3 - \bar{x}, \\y + z &= 1 - \bar{x}.\end{aligned}$$

Now the system is consistent and has unique solution. The inverse of the new matrix of coefficients is

$$\begin{bmatrix} 1 & -2 \\ -1 & 1 \end{bmatrix},$$

therefore

$$\begin{aligned}\begin{pmatrix} y \\ z \end{pmatrix} &= \mathbf{A}^{-1}\mathbf{b} \\ &= -\begin{bmatrix} 1 & -2 \\ -1 & 1 \end{bmatrix} \begin{pmatrix} 3 - \bar{x} \\ 1 - \bar{x} \end{pmatrix} = \begin{pmatrix} -1 - \bar{x} \\ 2 \end{pmatrix}.\end{aligned}$$

According to the value you choose for $\bar{x}$, you obtain a new vector solution. This is why we have infinitely many solutions.

We now observe the following: as $\mathbf{x} = \mathbf{A}^{-1}\mathbf{b}$, then by construction of the inverse matrix

$$\mathbf{x} = \frac{1}{|\mathbf{A}|}(\text{adj}\mathbf{A})\mathbf{b}.$$

Let us now perform the matrix multiplication which returns

$$\begin{aligned}
 \mathbf{x} &= \frac{1}{|\mathbf{A}|} (\text{adj } \mathbf{A}) \mathbf{b} \\
 &= \frac{1}{|\mathbf{A}|} \begin{bmatrix} b_1 A_{11} + b_2 A_{21} + \cdots + b_n A_{n1} \\ b_1 A_{12} + b_2 A_{22} + \cdots + b_n A_{n2} \\ \vdots \\ b_1 A_{1n} + b_2 A_{2n} + \cdots + b_n A_{nn} \end{bmatrix} \\
 &= \begin{pmatrix} \frac{|\mathbf{A}\mathbf{b}_1|}{|\mathbf{A}|} \\ \frac{|\mathbf{A}\mathbf{b}_2|}{|\mathbf{A}|} \\ \vdots \\ \frac{|\mathbf{A}\mathbf{b}_n|}{|\mathbf{A}|} \end{pmatrix}
 \end{aligned}$$

where $\mathbf{A}\mathbf{b}_i$ is the matrix formed by replacing the i -th column of $\mathbf{A}$ by the column vector $\mathbf{b}$. This gives us an alternative way to solve the linear system which is known as Cramer's rule.

Theorem 28 (Cramer) *Let $\mathbf{A}$ be a nonsingular $n \times n$ matrix, and let $\mathbf{b} \in \mathbf{R}^n$. Then the i^{th} coordinate of the vector $\mathbf{x}$ solving*

$$\mathbf{A}\mathbf{x} = \mathbf{b}$$

is

$$x_i = \frac{|\mathbf{A}\mathbf{b}_i|}{|\mathbf{A}|} \quad i = 1, \dots, n$$

where $\mathbf{A}\mathbf{b}_i$ is the matrix formed by replacing the i -th column of $\mathbf{A}$ by the column vector $\mathbf{b}$.

Example 95 Let us solve once again the linear systems in Example 93. This time we use Cramer's rule.

1. By Cramer's rule, the solution to the system is given by

$$x_1 = \frac{\begin{vmatrix} 4 & 2 \\ 2 & 1 \\ 1 & 2 \\ 3 & 1 \end{vmatrix}}{\begin{vmatrix} 1 & 2 \\ 2 & 1 \\ 1 & 2 \\ 3 & 1 \end{vmatrix}} = \frac{0}{-5} = 0$$

and

$$x_2 = \frac{\begin{vmatrix} 1 & 4 \\ 3 & 2 \end{vmatrix}}{\begin{vmatrix} 1 & 2 \\ 3 & 1 \end{vmatrix}} = \frac{-10}{-5} = 2.$$

Therefore $\mathbf{x} = \begin{pmatrix} 0 & 2 \end{pmatrix}^\top$.

2. From the previous analysis, the system can be rewritten as

$$\begin{aligned} y + 2z &= 3 - \bar{x}, \\ y + z &= 1 - \bar{x}. \end{aligned}$$

By Cramer's rule,

$$y = \frac{\begin{vmatrix} 3 - \bar{x} & 2 \\ 1 - \bar{x} & 1 \end{vmatrix}}{\begin{vmatrix} 1 & 2 \\ 1 & 1 \end{vmatrix}} = \frac{1 + \bar{x}}{-1} = -1 - \bar{x},$$

and

$$z = \frac{\begin{vmatrix} 1 & 3 - \bar{x} \\ 1 & 1 - \bar{x} \end{vmatrix}}{\begin{vmatrix} 1 & 2 \\ 1 & 1 \end{vmatrix}} = \frac{-2}{-1} = 2.$$

Therefore, the vector solution is $\begin{pmatrix} \bar{x} & -1 - \bar{x} & 2 \end{pmatrix}^\top$.

Example 96 Let us reconsider once again the problem of determining the coefficients of a linear function $y = a + bx$, which passes through the points (x_1, y_1) and (x_2, y_2) which we met in Section 1.2.1. We need to solve the system

$$\begin{aligned} y_1 &= a + bx_1 \\ y_2 &= a + bx_2, \end{aligned}$$

which can be rewritten as

$$\begin{bmatrix} 1 & x_1 \\ 1 & x_2 \end{bmatrix} \begin{pmatrix} a \\ b \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}.$$

The system is well defined as long as

$$\begin{vmatrix} 1 & x_1 \\ 1 & x_2 \end{vmatrix} = x_2 - x_1 \neq 0.$$

In this case, we can use Cramer's rule to obtain the solution

$$a = \frac{\begin{vmatrix} y_1 & x_1 \\ y_2 & x_2 \end{vmatrix}}{\begin{vmatrix} 1 & x_1 \\ 1 & x_2 \end{vmatrix}} = \frac{y_1x_2 - x_1y_2}{x_2 - x_1} = \frac{y_1x_2 - x_1y_1 + x_1y_1 - x_1y_2}{x_2 - x_1} = y_1 - x_1 \frac{y_2 - y_1}{x_2 - x_1},$$

and

$$b = \frac{\begin{vmatrix} 1 & y_1 \\ 1 & y_2 \end{vmatrix}}{\begin{vmatrix} 1 & x_1 \\ 1 & x_2 \end{vmatrix}} = \frac{y_2 - y_1}{x_2 - x_1}.$$

Consequently

$$y = y_1 + \frac{y_2 - y_1}{x_2 - x_1}(x - x_1).$$

Solving linear systems via the inverse matrix method or via Cramer's rule leads of course to the same result. So we choose either of the two methods. Cramer's rule though is particularly effective in the situation in which we are interested not in the full vector solution, but only in one specific element of the vector solution.

Example 97 Determine the value of z solution of the system

$$\begin{aligned} x + y + z &= 5, \\ x - 2y - 3z &= -1, \\ 2x + y - z &= 3. \end{aligned}$$

As in this case we are only after the value of z , then we can use Cramer's rule to obtain

$$z = \frac{\begin{vmatrix} 1 & 1 & 5 \\ 1 & -2 & -1 \\ 2 & 1 & 3 \end{vmatrix}}{\begin{vmatrix} 1 & 1 & 1 \\ 1 & -2 & -3 \\ 2 & 1 & -1 \end{vmatrix}} = 3.$$

(You can verify as an exercise that the full solution is $(4 \ -2 \ 3)^\top$).

We can also solve parametric systems.

Example 98 Consider the linear system in parametric form

$$\begin{aligned} kx_1 + x_2 &= -k \\ -2x_1 + x_2 + kx_3 &= 0 \\ x_2 + x_3 &= 2 \quad k \in \mathbb{R}. \end{aligned}$$

In this case

$$\mathbf{A} = \begin{bmatrix} k & 1 & 0 \\ -2 & 1 & k \\ 0 & 1 & 1 \end{bmatrix}.$$

This matrix has full rank if $k \neq 2, k \neq -1$. For these values of k the system is consistent and has a unique solution. We can use Cramer's rule to show that the vector solution is

$$\mathbf{x} = \begin{bmatrix} \frac{k}{2-k} \\ -\frac{2k}{2-k} \\ \frac{4}{2-k} \end{bmatrix}.$$

If $k = 2$, instead,

$$\mathbf{A} = \begin{bmatrix} 2 & 1 & 0 \\ -2 & 1 & 2 \\ 0 & 1 & 1 \end{bmatrix}.$$

This matrix has zero determinant as the second row is a linear combination of the first and the third one ($2r_3 - r_1$). However, the augmented matrix $(\mathbf{A}|\mathbf{b})$ has rank 3 as it has at least a minor

of order 3. Consequently, in this case the system is inconsistent and no solution exist. If $k = -1$, on the other hand,

$$\mathbf{A} = \begin{bmatrix} -1 & 1 & 0 \\ -2 & 1 & -1 \\ 0 & 1 & 1 \end{bmatrix}.$$

Also in this case, the matrix contains a linear combination ($r_1 = 0.5(r_2 + r_3)$). However, it has at least a minor of order 2; hence its rank is 2. The augmented matrix is

$$\mathbf{A}|\mathbf{b} = \left[\begin{array}{ccc|c} -1 & 1 & 0 & 1 \\ -2 & 1 & -1 & 0 \\ 0 & 1 & 1 & 2 \end{array} \right],$$

which contains the same linear combination as matrix $\mathbf{A}$. Therefore, the system is feasible and it admits infinitely many solutions with 1 degree of freedom. In fact, set $x_3 = \alpha \in \mathbb{R}$. Then $x_1 = 1 - \alpha$ and $x_2 = 2 - \alpha$.

2.5 Quadratic forms

Quadratic forms are frequent in Finance, especially in portfolio theory. However, they are also very important to prepare the ground for the analysis of minima and maxima of functions in more than one variable, and for constrained optimization.

We begin from the Definition.

Definition 29 *Given an n -dimensional column vector of unknowns/variables $\mathbf{x}$ and a n -dimensional symmetric matrix of scalars $\mathbf{A}$, the quadratic form associated to $\mathbf{x}$ and $\mathbf{A}$ is defined as the real function*

$$q(\mathbf{x}) = \mathbf{x}^\top \mathbf{A} \mathbf{x}.$$

Notice that the dimension of

$$\underset{1 \times 1}{q} = \underset{1 \times n}{\mathbf{x}^\top} \underset{n \times n}{\mathbf{A}} \underset{n \times 1}{\mathbf{x}}$$

is 1×1 , that is q is a number.

If we write the quadratic form explicitly, its expression becomes more “friendly”

$$\begin{aligned}
\mathbf{x}^\top \mathbf{A} \mathbf{x} &= \begin{pmatrix} x_1 & x_2 & \dots & x_n \end{pmatrix} \begin{pmatrix} a_{1,1} & a_{1,2} & \dots & a_{1,n} \\ a_{2,1} & a_{2,2} & \dots & a_{2,n} \\ \dots & \dots & \dots & \dots \\ a_{n,1} & a_{n,2} & \dots & a_{n,n} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{pmatrix} \\
&= \begin{pmatrix} \sum_{i=1}^n a_{i,1}x_i & \sum_{i=1}^n a_{i,2}x_i & \dots & \sum_{i=1}^n a_{i,n}x_i \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{pmatrix} \\
&= (a_{1,1}x_1 + a_{2,1}x_2 + \dots + a_{n,1}x_n)x_1 + \dots + (a_{1,n}x_1 + a_{2,n}x_2 + \dots + a_{n,n}x_n)x_n \\
&= \sum_{i=1}^n \sum_{j=1}^n a_{i,j}x_i x_j.
\end{aligned}$$

Due to the fact that the matrix is symmetric, that is $a_{i,j} = a_{j,i}$, the last sum can be rewritten as

$$q(x_1, \dots, x_n) = \sum_{i=1}^n a_{i,i}x_i^2 + 2 \sum_{i=1}^{n-1} \sum_{j>i}^n \underbrace{a_{i,j}x_i x_j}_{\text{cross-product}}.$$

This last expression should clarify why q is called quadratic form: a quadratic form is just a convenient way to describe a polynomial of degree 2. The expression also tells how we can recover the matrix behind the quadratic form: this matrix has elements on the main diagonal given by $a_{i,i}$, $i = 1, \dots, n$; the elements on the remaining diagonals are the ones generated by the cross-product term.

Example 99 Let $\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$.

1. Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 4 & 0 \\ 0 & 3 \end{bmatrix}.$$

Then the corresponding quadratic form is

$$4x_1^2 + 3x_2^2.$$

2. Consider the matrix

$$\mathbf{A} = \begin{bmatrix} 3 & -2 \\ -2 & 7 \end{bmatrix}.$$

Then the corresponding quadratic form is

$$3x_1^2 - 4x_1x_2 + 7x_2^2.$$

3. Consider the quadratic form

$$x_1^2 - 8x_1x_2 - 5x_2^2.$$

The symmetric matrix such that

$$x_1^2 - 8x_1x_2 - 5x_2^2 = \mathbf{x}^\top \mathbf{A} \mathbf{x}$$

is

$$\mathbf{A} = \begin{bmatrix} 1 & -4 \\ -4 & -5 \end{bmatrix},$$

as

$$x_1^2 - 8x_1x_2 - 5x_2^2 = 1 \times x_1^2 + 2 \times (-4)x_1x_2 - 5x_2^2$$

As most of the times we will be interested in quadratic forms such that $q(\mathbf{x}) \geq 0$, we give the following definition.

Definition 30 Given a n -dimensional symmetric matrix $\mathbf{A}$ and n -dimensional vectors $\mathbf{x} \neq \mathbf{0}$, then

1. $\mathbf{A}$ is positive semi-definite if $q(\mathbf{x}) \geq 0$, for all $\mathbf{x}$;
2. $\mathbf{A}$ is positive definite if $q(\mathbf{x}) > 0$, for all $\mathbf{x} \neq \mathbf{0}$;
3. $\mathbf{A}$ is negative semi-definite if $q(\mathbf{x}) \leq 0$, for all $\mathbf{x}$;
4. $\mathbf{A}$ is negative definite if $q(\mathbf{x}) < 0$, for all $\mathbf{x} \neq \mathbf{0}$;
5. $\mathbf{A}$ is indefinite if $q(\mathbf{x})$ is either positive, zero or negative based on the choice of $\mathbf{x}$.

If a matrix is positive semi-definite, $q(\mathbf{x})$ is always a non-negative number and we can calculate $\sqrt{q(\mathbf{x})}$.

Example 100 1. Consider the quadratic form given in Example 99.1. In this case

$$q(\mathbf{x}) = 4x_1^2 + 3x_2^2 > 0$$

for any $x_1, x_2 \neq 0$. Therefore

$$\mathbf{A} = \begin{bmatrix} 4 & 0 \\ 0 & 3 \end{bmatrix}$$

is positive definite. To understand the meaning of positive definite, let us have a look at the plot of the function

$$z = 4x_1^2 + 3x_2^2$$

in Figure 26. The positive definite matrix originates a quadratic form which is strictly

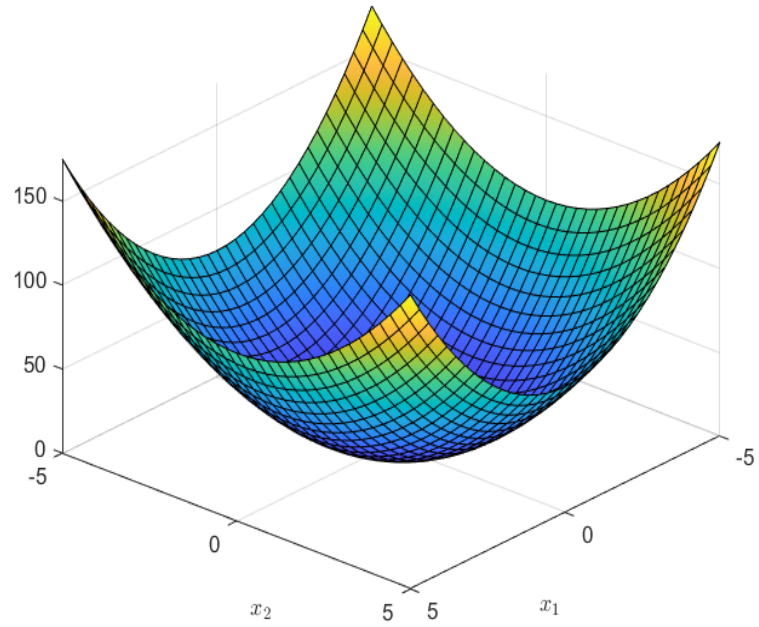


Figure 26: Quadratic form $q(\mathbf{x})$ originated by a positive definite matrix.

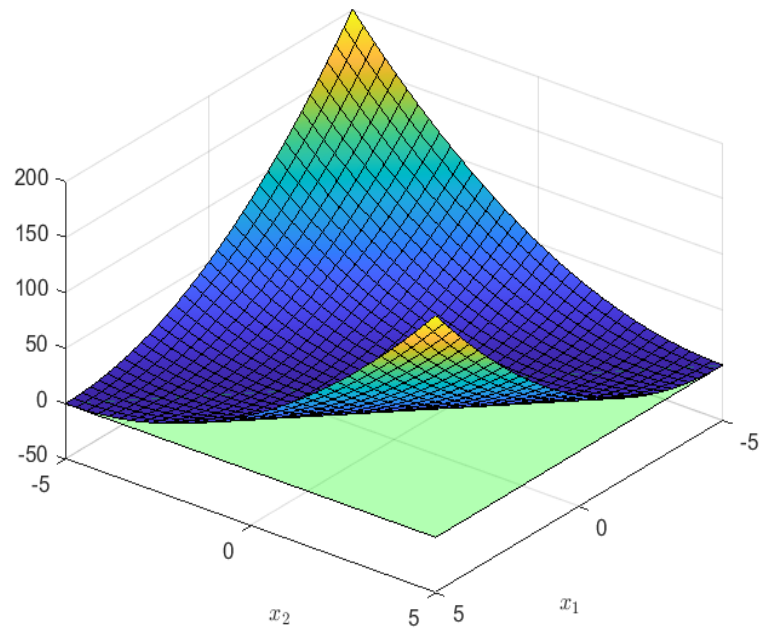


Figure 27: Quadratic form $q(\mathbf{x})$ originated by a positive semi-definite matrix.

positive except at the origin, i.e. a surface which lies above the x_1x_2 plane which goes through the origin. In our case, we have an elliptic paraboloid which is entirely in the positive part of z with the exception of the origin. This is opposed to a positive semi-definite matrix, such as for example

$$\begin{bmatrix} 2 & 2 \\ 2 & 2 \end{bmatrix},$$

which originates the non-negative quadratic form

$$q(\mathbf{x}) = 2x_1^2 + 4x_1x_2 + 2x_2^2$$

plotted in Figure 27. In this case the corresponding surface touches the x_1x_2 plane along a line through the origin, although this surface never goes in the negative z -axis space.

2. Similarly, negative definite (resp. semi-definite) matrices originate quadratic forms which are negative (resp. non positive). Consequently, the surface corresponding to a negative definite matrix lies in the negative part of the z -axis (except at the origin), whilst the surface corresponding to a negative semi-definite matrix touches the x_1x_2 plane along a line through the origin although it never becomes positive. These are the cases illustrated respectively in Figure 28 for

$$q(\mathbf{x}) = -4x_1^2 + 4x_1x_2 - 5x_2^2$$

and in Figure 29 for

$$q(\mathbf{x}) = -2x_1^2 + 4x_1x_2 - 2x_2^2.$$

3. And the indefinite form? As you can imagine, the corresponding surface crosses the x_1x_2 plane along a line through the origin, so it has a portion which is negative, and the rest is positive. Like in Figure 30, which portrays the quadratic form

$$q(\mathbf{x}) = x_1^2 - 8x_1x_2 - 5x_2^2$$

from Example 99.3.

As the previous examples show, it is not always straightforward to determine the sign of $q(\mathbf{x})$, and consequently if the matrix is positive definite or else. We need an alternative approach. This is offered by Sylvester's criterion.

Sylvester's criterion

A symmetric matrix $\mathbf{A}$ is

1. positive definite if and only if all of the leading principal minors are positive. This means that all the following matrices have a positive determinant
 - (a) the upper left 1×1 corner of $\mathbf{A}$,
 - (b) the upper left 2×2 corner of $\mathbf{A}$,
 - (c) the upper left 3×3 corner of $\mathbf{A}$,

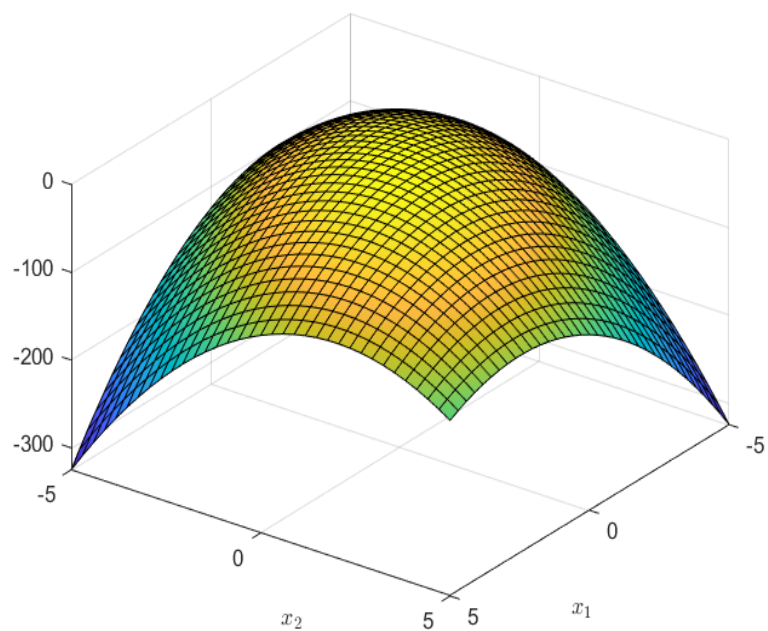


Figure 28: Quadratic form $q(\mathbf{x})$ originated by a negative definite matrix.

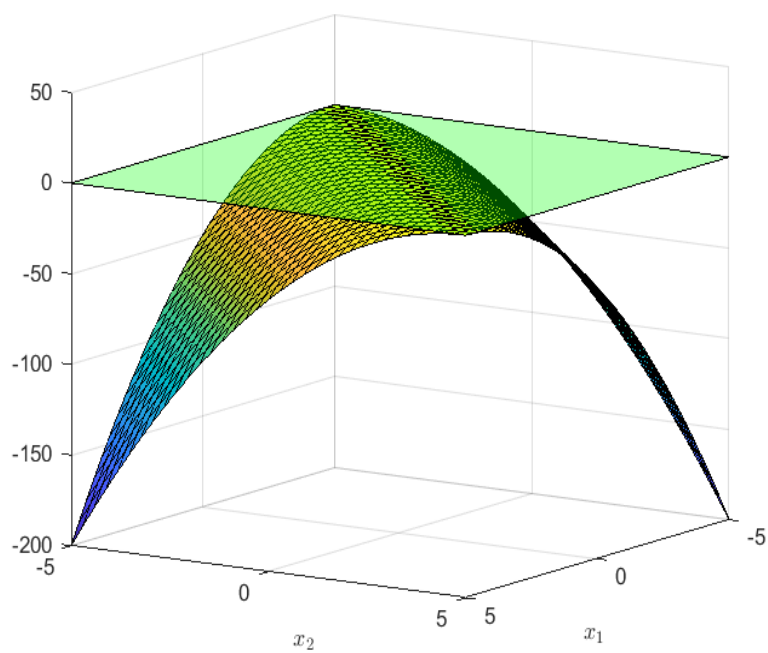


Figure 29: Quadratic form $q(\mathbf{x})$ originated by a negative semi-definite matrix.

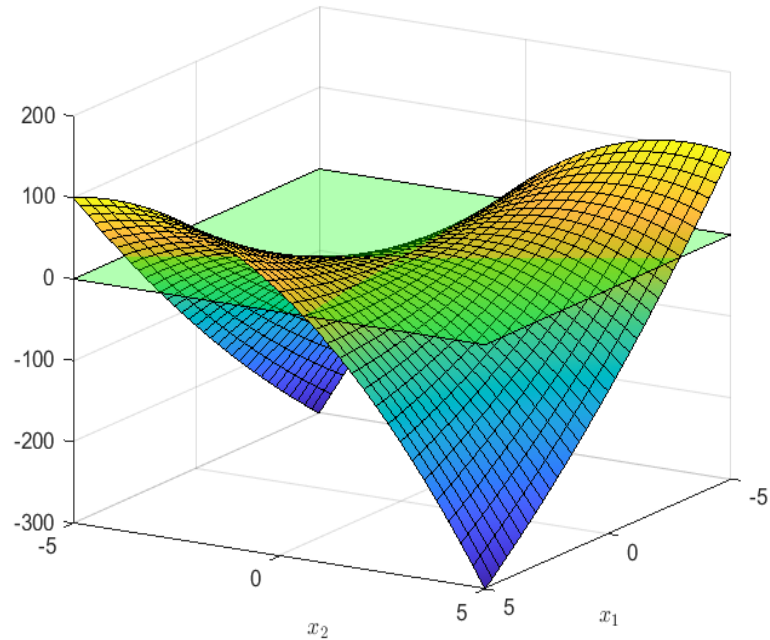


Figure 30: Quadratic form $q(\mathbf{x})$ originated by an indefinite matrix.

$\vdots$

(d) $\mathbf{A}$ itself.

2. positive semi-definite, if and only if all principal minors of $\mathbf{A}$ are nonnegative.
3. negative definite if and only if $-\mathbf{A}$ is positive definite. Notice that, due to the properties of the determinant, it means

the upper left 1×1 corner of $\mathbf{A}$ has negative determinant,

the upper left 2×2 corner of $\mathbf{A}$ has positive determinant,

the upper left 3×3 corner of $\mathbf{A}$ has negative determinant,

$\vdots$

and so on with alternate signs.

4. negative semi-definite if and only if all its principal minors of even dimension are non-negative and all principal minors of odd dimension are non-positive

Example 101 We consider the matrices corresponding to the quadratic forms in Example 100 one by one.

- 1.

$$q(\mathbf{x}) = 4x_1^2 + 3x_2^2 > 0$$

for which

$$\mathbf{A} = \begin{bmatrix} 4 & 0 \\ 0 & 3 \end{bmatrix}.$$

The leading principal minors are $|\mathbf{A}_1| = 4$, $|\mathbf{A}_2| = 12$. Both are positive, therefore the matrix is positive definite.

2.

$$q(\mathbf{x}) = 2x_1^2 + 4x_1x_2 + 2x_2^2$$

for which

$$\mathbf{A} = \begin{bmatrix} 2 & 2 \\ 2 & 2 \end{bmatrix}.$$

The leading principal minors are $|\mathbf{A}_1| = 2$, $|\mathbf{A}_2| = 0$. Which means that the criterium for positive definite matrix is not satisfied. So, we move to check all the principal minors of the matrix. The ones of order 1 are $|a_{11}| = 2 > 0$, $|a_{22}| = 2 > 0$, whilst the one of order 2 is the determinant which have already seen to be 0. Therefore, all principal minor are non-negative, and consequently the matrix is positive semi-definite.

3.

$$q(\mathbf{x}) = -4x_1^2 + 4x_1x_2 - 5x_2^2$$

for which

$$\mathbf{A} = \begin{bmatrix} -4 & 2 \\ 2 & -5 \end{bmatrix}.$$

Then $|\mathbf{A}_1| = -4$, $|\mathbf{A}_2| = 16$. The leading principal minor have alternating signs with the first negative, therefore the matrix is negative definite.

4.

$$q(\mathbf{x}) = -2x_1^2 + 4x_1x_2 - 2x_2^2$$

for which

$$\mathbf{A} = \begin{bmatrix} -2 & 2 \\ 2 & -2 \end{bmatrix}.$$

Then $|\mathbf{A}_1| = -2$, $|\mathbf{A}_2| = 0$. As the criterium for negative definite matrices is not satisfied, we move of to the principal minors. The ones of order 1 are in sequence $-2 < 0$, $-2 < 0$, whilst the one of order 2 is the determinant which is 0. In other words, all principal minors are non-positive, and therefore the matrix is negative semi-definite.

5.

$$q(\mathbf{x}) = x_1^2 - 8x_1x_2 - 5x_2^2$$

for which

$$\mathbf{A} = \begin{bmatrix} 1 & -4 \\ -4 & -5 \end{bmatrix}.$$

Then $|\mathbf{A}_1| = 1$, $|\mathbf{A}_2| = -21$. The principal minors of order 1 are 1, -5 and -21 . Therefore, this form is indefinite as it does not satisfy any of the conditions of the Sylvester's criterion.

Example 102 (A financial application.) A fund manager is trying to construct a diversified portfolio, which contains stocks, bonds and commodity indices. Collecting and analysing monthly data for each index for the past 5 years has revealed the following information regarding their risk profile.

	Stock	Bonds	Commodities
Stock	0.09	-0.0094	-0.0143
Bonds	-0.0094	0.0036	0.002
Commodities	-0.0143	0.002	0.0196

This table can be converted in the so-called variance/covariance matrix

$$\Sigma = \begin{bmatrix} 0.09 & -0.0094 & -0.0143 \\ -0.0094 & 0.0036 & 0.002 \\ -0.0143 & 0.002 & 0.0196 \end{bmatrix}$$

in which the elements on the principal diagonal represent the risk of each individual index, whilst the elements on the off-diagonal represent the interaction between each index. Suppose the fund manager is considering the portfolio with weights $\mathbf{w} = \begin{pmatrix} 0.5 & 0.25 & 0.25 \end{pmatrix}^\top$. The risk profile of the portfolio is described by the quadratic form

$$q(\mathbf{w}) = \mathbf{w}^\top \Sigma \mathbf{w} = 0.0183.$$

This is the variance of the portfolio. However, in finance we usually refer to the volatility/standard deviation of the portfolio which is given by

$$\sqrt{q(\mathbf{w})} = 13.53\%;$$

this quantity is of course well defined if $q(\mathbf{w}) > 0$, i.e. when Σ is positive definite.

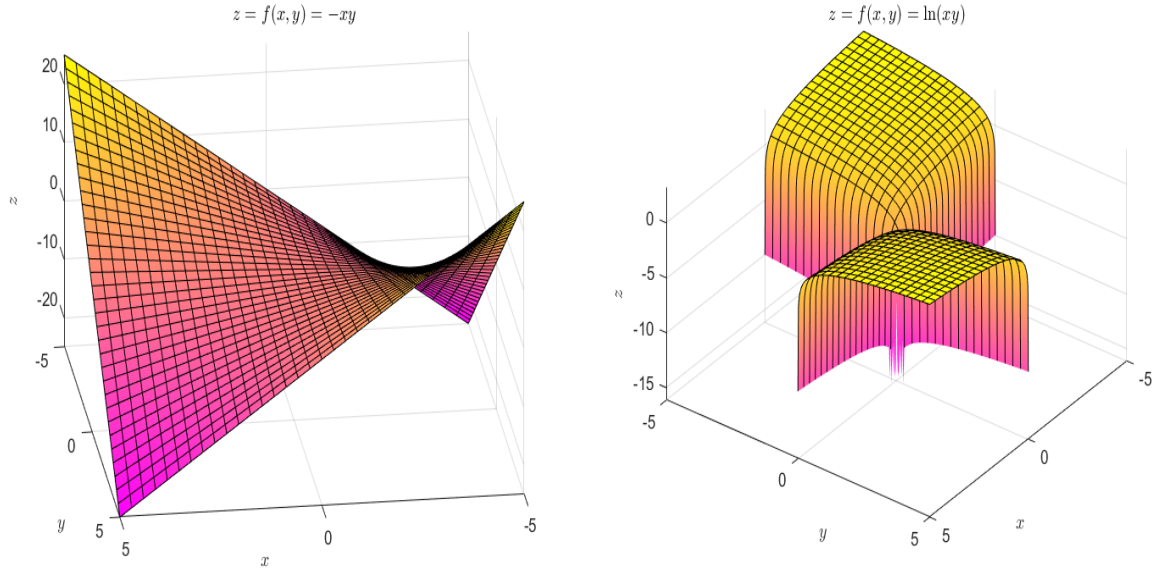


Figure 31: Functions in two variables: surfaces in 3D.

3 Calculus in more than one variable

3.1 Real functions of several variables

A real-valued function of n real variables is a rule that assigns to n real numbers $(x_1, x_2, \dots, x_n)$ a unique real number denoted by $f(x_1, x_2, \dots, x_n)$.

Example 103 1. The function

$$f(x, y) = -xy$$

with domain in $\mathbb{R}^2$ - see Figure 31.

2. The function

$$f(x, y) = \ln(xy)$$

meaningfully defined only for $xy > 0$ - see Figure 31.

3. The function

$$f(x_1, x_2, \dots, x_n) = x_1 + x_2 + \dots + x_n$$

with domain in $\mathbb{R}^n$.

Note: a very useful tool, borrowed from mapmakers, is a contour map on which points of constant elevation are joined to form contour curves, or level curves. The level curves of a function of two variables are the curves with equations $f(x, y) = k$, where k is a constant (in the range of f) - see for example Figure 32. We met already level curves when we looked at indifference curves from a given utility function - see Figure 9.

All the results that we have studied in Chapter 1, such as limits and continuity, for example, can be extended to the case of functions of more than one variable in a straightforward manner. In the following, we focus on differentiation as this plays a special role in Finance and is at the core of constrained optimization, i.e. the typical problem to be solved in the large majority of the areas of Finance.

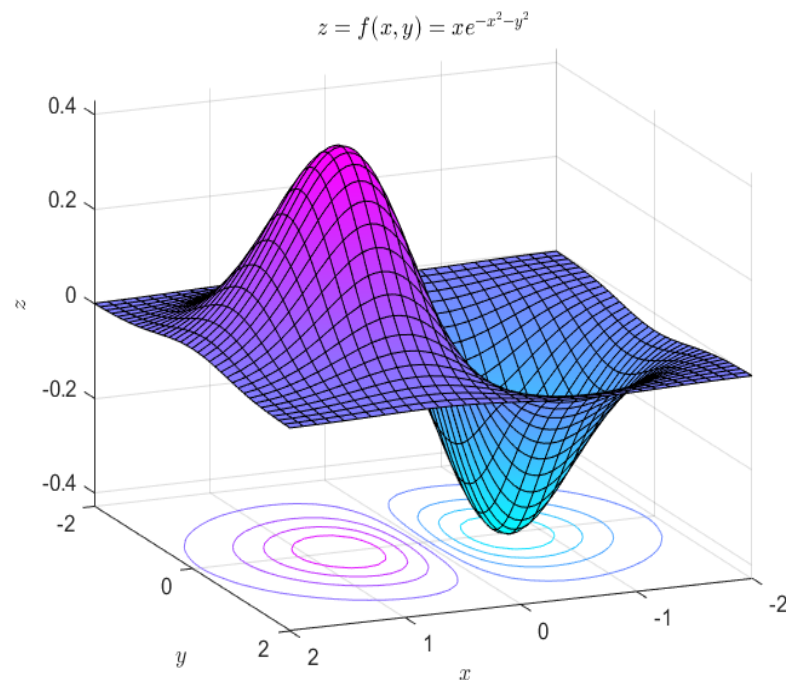


Figure 32: Surfaces and contour map.

Example 104 1. A very important task of financial planners is asset allocation; this process consists in spreading your investments across various asset classes such as stocks (S), bonds (B) and cash/money market securities (M). Thus, the portfolio has form

$$f(w_S, w_B, w_M) = w_S S + w_B B + w_M M$$

for (w_S, w_B, w_M) representing the weights of each asset. The goal of the financial planner is therefore to choose (w_S, w_B, w_M) so that risk is minimized while a certain return is achieved.

Unless otherwise specified, in the following of this section we always refer to the case of a function of two variable, i.e. $z = f(x, y)$, as this will allow us to gain a geometric intuition of the several concepts we will cover. However, the results we present hold also for the more general case of functions of n independent variables.

3.2 Partial differentiation

3.2.1 First order partial derivatives and the gradient

Consider the function $z = f(x, y)$. Suppose we let only x vary, while keeping y fixed, say $y = y_0$. Then, the function can be considered as a function of 1 variable $g(x) = f(x, y_0)$. If this is differentiable, we call its derivative the partial derivative of f with respect to x . Similarly, we can repeat this argument by letting y vary while keeping x fixed and obtain the partial derivative with respect to y .

The above idea is formalized as follows.

Definition 31 *Let f be a function of two variables; its partial derivatives are the functions f_x , f_y defined by*

$$\begin{aligned} f_x(x, y) &= \lim_{h \rightarrow 0} \frac{f(x+h, y) - f(x, y)}{h} \\ f_y(x, y) &= \lim_{h \rightarrow 0} \frac{f(x, y+h) - f(x, y)}{h} \end{aligned}$$

Notation. If $z = f(x, y)$, then its partial derivatives can be denoted as

$$\begin{aligned} f_x(x, y) &= f_x = \frac{\partial f}{\partial x} = \frac{\partial}{\partial x} f(x, y) = \frac{\partial z}{\partial x}. \\ f_y(x, y) &= f_y = \frac{\partial f}{\partial y} = \frac{\partial}{\partial y} f(x, y) = \frac{\partial z}{\partial y}. \end{aligned}$$

All of the above are equivalent.

We usually order the partial derivatives in a column vector called gradient and denoted by the letter ∇ (pronounced nabla)

$$\nabla f(x, y) = \begin{bmatrix} \frac{\partial f}{\partial x} \\ \frac{\partial f}{\partial y} \end{bmatrix}.$$

In the n -dimensional case

$$\nabla f(x_1, x_2, \dots, x_n) = \begin{bmatrix} \frac{\partial f}{\partial x_1} \\ \vdots \\ \frac{\partial f}{\partial x_n} \end{bmatrix}.$$

Rule for finding the partial derivatives of $z = f(\mathbf{x}, \mathbf{y})$. To find f_x , regard y as a constant and differentiate $z = f(x, y)$ with respect to x . Similarly, to find f_y , regard x as a constant and differentiate $z = f(x, y)$ with respect to y

Example 105 1. Let $z = x^2 + 4y^5$. We start with the first order derivative with respect to x , which means to treat y as a constant (remember that x and y are independent)

$$\begin{aligned}\frac{\partial z}{\partial x} &= \frac{\partial z}{\partial x}(x^2) + \frac{\partial z}{\partial x}(4y^5) \quad (\text{rule of the sum}) \\ &= 2x + \frac{\partial z}{\partial x}(4y^5) \quad (\text{power rule}) \\ &= 2x + 0 \quad (\text{derivative of a constant}) \\ &= 2x.\end{aligned}$$

Following the same logic, we now obtain the first order differential with respect to y treating x as a constant, so that

$$\begin{aligned}\frac{\partial z}{\partial y} &= \frac{\partial z}{\partial y}(x^2) + \frac{\partial z}{\partial y}(4y^5) \quad (\text{rule of the sum}) \\ &= 0 + \frac{\partial z}{\partial y}(4y^5) \quad (\text{derivative of a constant}) \\ &= 0 + 4 \times 5y^4 \quad (\text{power rule}) \\ &= 20y^4.\end{aligned}$$

Therefore

$$\nabla z = \begin{bmatrix} \frac{\partial z}{\partial x} \\ \frac{\partial z}{\partial y} \end{bmatrix} = \begin{bmatrix} 2x \\ 20y^4 \end{bmatrix}.$$

2. Let $z = x^3 + x^2y^3 - 2y^2$. Then

$$\frac{\partial z}{\partial x} = 3x^2 + 2xy^3,$$

and

$$\frac{\partial z}{\partial y} = 3x^2y^2 - 4y.$$

Therefore

$$\nabla z = \begin{bmatrix} \frac{\partial z}{\partial x} \\ \frac{\partial z}{\partial y} \end{bmatrix} = \begin{bmatrix} 3x^2 + 2xy^3 \\ 3x^2y^2 - 4y \end{bmatrix}.$$

3. Let $f(x, y, z) = x^2 + y^2 + z^2$. Then

$$\frac{\partial f}{\partial x} = 2x,$$

$$\frac{\partial f}{\partial y} = 2y$$

and

$$\frac{\partial f}{\partial z} = 2z$$

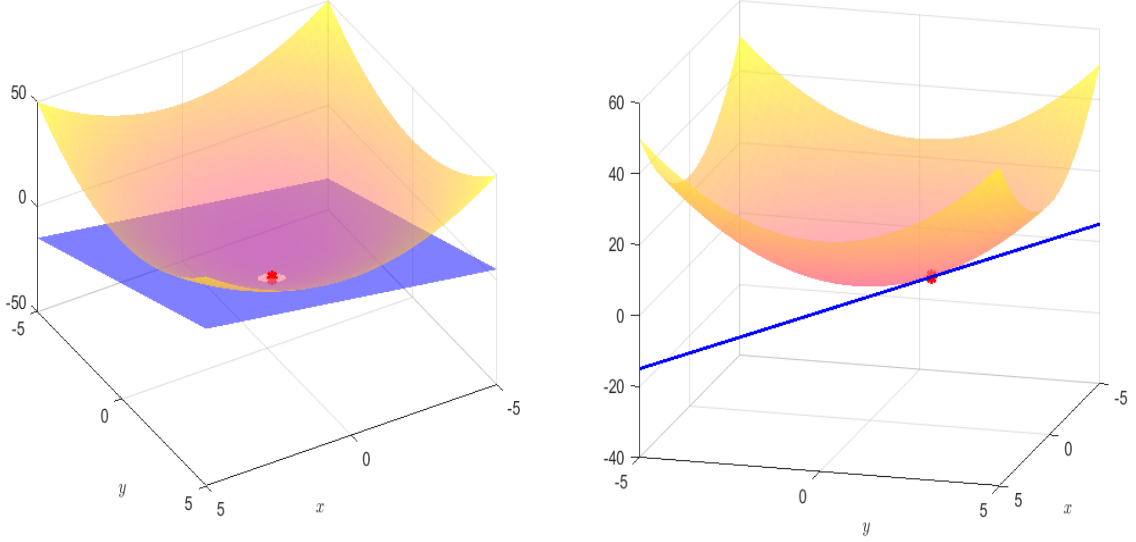


Figure 33: Partial derivatives and the tangent plane.

Therefore

$$\nabla f(x, y, z) = \begin{bmatrix} \frac{\partial f}{\partial x} \\ \frac{\partial f}{\partial y} \\ \frac{\partial f}{\partial z} \end{bmatrix} = \begin{bmatrix} 2x \\ 2y \\ 2z \end{bmatrix}.$$

Partial derivatives: tangent planes and rate of changes In the one dimensional case, the derivative of a function in a point returns the slope of the tangent line to the function in the given point. Similarly, in the multidimensional case the partial derivatives of a function in a point give the slope of the tangent plane to the surface in the given point. This is illustrated in Figure 33, in which we consider the function

$$f(x, y) = x^2 + y^2,$$

the point $(x_0, y_0, f(x_0, y_0))$ and the tangent plane to f in this point with equation

$$g(x, y) = f(x_0, y_0) + \frac{\partial}{\partial x} f(x_0, y_0)(x - x_0) + \frac{\partial}{\partial y} f(x_0, y_0)(y - y_0).$$

Therefore, f_x is the slope of this plane seen from the x axis (i.e. keeping y fixed), whilst f_y is the slope of the plane seen from the y axis (i.e. keeping x fixed).

Partial derivatives can of course be also considered as rates of change: f_x is the rate of change of f with respect to x and f_y is the rate of change of f with respect to y . In this respect, partial derivatives are essential in Economics and Finance to run sensitivity analysis. Sensitivity analysis determines how target output variables are affected based on changes in the input variables. It is also referred to as what-if or simulation analysis. It is usually carried out

by changing one input at the time, leaving all the others unchanged (usually you see the Latin phrase ‘ceteris paribus’ to indicate this aspect). If the target output variable can be expressed via an analytical formula, then partial differentiation helps to perform the sensitivity analysis efficiently.

Example 106 1. Consider a utility function $u(x_1, x_2)$; the partial derivatives

$$\frac{\partial u}{\partial x_1}, \quad \frac{\partial u}{\partial x_2}$$

represent the marginal utilities (MU) with respect to good 1 and good 2, respectively, i.e.

$$\begin{aligned} MU_1 &= \frac{\partial u}{\partial x_1} \\ MU_2 &= \frac{\partial u}{\partial x_2}. \end{aligned}$$

They quantify the change in utility from an increase in consumption of one specific good, or service.

Total differential. For a differentiable function of two variables, $z = f(x, y)$, we define the differentials dx and dy to be independent variables; that is, they can be given any values. Then the differential dz , also called the total differential, is defined by

$$dz = \frac{\partial z}{\partial x} dx + \frac{\partial z}{\partial y} dy.$$

Intuitively speaking, if x changes by a quantity h from x_0 to $x_0 + h$, and y changes by a quantity k from y_0 to $y_0 + k$, then z changes from $z_0 = f(x_0, y_0)$ to $z = f(x_0 + h, y_0 + k)$, and we can approximate the amount of this total change as

$$\underbrace{z - z_0}_{dz} = \frac{\partial}{\partial x} f(x_0, y_0) \underbrace{h}_{dx} + \frac{\partial}{\partial y} f(x_0, y_0) \underbrace{k}_{dy}.$$

Note that the above can be equivalently written as

$$z = f(x_0, y_0) + \frac{\partial}{\partial x} f(x_0, y_0)(x - x_0) + \frac{\partial}{\partial y} f(x_0, y_0)(y - y_0)$$

which is the equation of the tangent plane through (x_0, y_0) seen above.

Example 107 Consider the utility function $U = u(x_1, x_2)$. The total differential is

$$dU = \frac{\partial u}{\partial x_1} dx_1 + \frac{\partial u}{\partial x_2} dx_2.$$

Now, assume we want to increase the consumption of good 1, x_1 , without changing the utility level U . The above relation says that we need to modify the consumption of good 2, x_2 . How? We can use the total differential: as we do not want U to change, then $dU = 0$, which means

$$0 = \frac{\partial u}{\partial x_1} dx_1 + \frac{\partial u}{\partial x_2} dx_2,$$

from which we obtain - using the concept of marginal utility introduced in Example 106 -

$$\frac{dx_2}{dx_1} = -\frac{\partial u/\partial x_1}{\partial u/\partial x_2} = -\frac{MU_1}{MU_2},$$

i.e. the marginal rate of substitution (MRS). The negative sign, tells us that we need to reduce the consumption of good 2. Note that we obtained already the MRS via implicit differentiation in Section 1.5.3: indeed it is the slope of the tangent to the indifference curves.

Implicit differentiation We can now properly formalize the implicit differentiation procedure which we saw in Section 1.5.3. Assume that an equation of the form

$$F(x, y) = 0$$

defines implicitly y as a differentiable function of x , i.e. $y = f(x)$ and therefore $F(x, y) = F(x, f(x)) = 0$. If F is differentiable, we can apply the chain rule to differentiate both sides of the equation $F(x, y) = 0$ with respect to x to obtain

$$F_x \frac{dx}{dx} + F_y \frac{dy}{dx} = 0.$$

As

$$\frac{dx}{dx} = 1,$$

we can rewrite the above as

$$F_x + F_y \frac{dy}{dx} = 0, \tag{10}$$

consequently

$$y' = \frac{dy}{dx} = -\frac{F_x}{F_y},$$

3.2.2 Second order partial derivatives and the Hessian

As partial derivatives are in general functions of more than one variable, we can also consider their partial derivatives. This means that, for $z = f(x, y)$, we are looking at the second order partial derivative of f , denoted as

$$\begin{aligned} f_{xx}(x, y) &= f_{xx} = \frac{\partial^2 f}{\partial x^2} = \frac{\partial^2}{\partial x^2} f(x, y) = \frac{\partial^2 z}{\partial x^2} \\ f_{yy}(x, y) &= f_{yy} = \frac{\partial^2 f}{\partial y^2} = \frac{\partial^2}{\partial y^2} f(x, y) = \frac{\partial^2 z}{\partial y^2} \\ f_{xy}(x, y) &= f_{xy} = \frac{\partial^2 f}{\partial x \partial y} = \frac{\partial^2}{\partial x \partial y} f(x, y) = \frac{\partial^2 z}{\partial x \partial y} \\ f_{yx}(x, y) &= f_{yx} = \frac{\partial^2 f}{\partial y \partial x} = \frac{\partial^2}{\partial y \partial x} f(x, y) = \frac{\partial^2 z}{\partial y \partial x}. \end{aligned}$$

All of the above are equivalent. f_{xy} and f_{yx} are called mixed derivatives. For well behaved functions

$$\frac{\partial^2 z}{\partial x \partial y} = \frac{\partial^2 z}{\partial y \partial x}$$

(Schwarz's theorem - also known as Clairaut or Young theorem).

For convenience, we organize all the second order partial derivatives in a matrix, called the Hessian matrix

$$\mathbf{H}f(x, y) = \begin{bmatrix} \frac{\partial^2 z}{\partial x^2} & \frac{\partial^2 z}{\partial x \partial y} \\ \frac{\partial^2 z}{\partial y \partial x} & \frac{\partial^2 z}{\partial y^2} \end{bmatrix}.$$

For the general case of a function in $\mathbb{R}^n$, we have

$$\mathbf{H}f(x_1, \dots, x_n) = \begin{bmatrix} \frac{\partial^2 f}{\partial x_1^2} & \frac{\partial^2 f}{\partial x_1 \partial x_2} & \cdots & \frac{\partial^2 f}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f}{\partial x_2 \partial x_1} & \frac{\partial^2 f}{\partial x_2^2} & \cdots & \frac{\partial^2 f}{\partial x_2 \partial x_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial x_n \partial x_1} & \frac{\partial^2 f}{\partial x_n \partial x_2} & \cdots & \frac{\partial^2 f}{\partial x_n^2} \end{bmatrix}$$

or, by stating an equation for the coefficients using indices i and j , as we learnt in Section 2.1,

$$\mathbf{H}_{i,j}f = \frac{\partial^2 f}{\partial x_i \partial x_j}.$$

Remark 15 The Hessian matrix is a symmetric matrix due to Schwarz's theorem.

Example 108 1. Let $z = x^2 + 4y^5$. Then

$$\nabla z = \begin{bmatrix} \frac{\partial z}{\partial x} \\ \frac{\partial z}{\partial y} \end{bmatrix} = \begin{bmatrix} 2x \\ 20y^4 \end{bmatrix}$$

see Example 105.1. Now we look into the second order derivatives. Let us first obtain the second order derivatives with respect to x . This means to take the first order derivative with respect to x and differentiate with respect to x once again.

$$\begin{aligned} \frac{\partial^2 z}{\partial x^2} &= \frac{\partial}{\partial x} \left(\frac{\partial z}{\partial x} \right) \\ &= \frac{\partial}{\partial x} (2x) \\ &= 2 \end{aligned}$$

using the standard rules for differentiation. Similarly, we can obtain the second order derivative with respect to y

$$\begin{aligned} \frac{\partial^2 z}{\partial y^2} &= \frac{\partial}{\partial y} \left(\frac{\partial z}{\partial y} \right) \\ &= \frac{\partial}{\partial y} (20y^4) \\ &= 80y^3. \end{aligned}$$

For the mixed derivatives, we proceed as follows. We take the first order differential with respect to y and we differentiate it with respect to x , i.e.

$$\begin{aligned}\frac{\partial^2 z}{\partial x \partial y} &= \frac{\partial}{\partial x} \left(\frac{\partial z}{\partial y} \right) \\ &= \frac{\partial z}{\partial x} (20y^4) \\ &= 0\end{aligned}$$

as y does not depend on x and therefore is equivalent to a constant with respect to changes in x . clearly, we obtain the same result if we take the first order differential with respect to x and differentiate it with respect to y , i.e.

$$\begin{aligned}\frac{\partial^2 z}{\partial y \partial x} &= \frac{\partial}{\partial y} \left(\frac{\partial z}{\partial x} \right) \\ &= \frac{\partial z}{\partial y} (2x) \\ &= 0.\end{aligned}$$

Consequently, the Hessian matrix is

$$\mathbf{H}z = \begin{bmatrix} 2 & 0 \\ 0 & 80y^3 \end{bmatrix}.$$

2. Let $z = x^3 + x^2y^3 - 2y^2$. Then

$$\nabla z = \begin{bmatrix} \frac{\partial z}{\partial x} \\ \frac{\partial z}{\partial y} \end{bmatrix} = \begin{bmatrix} 3x^2 + 2xy^3 \\ 3x^2y^2 - 4y \end{bmatrix}$$

see Example 105.2. Therefore

$$\mathbf{H}z = \begin{bmatrix} 6x + 2y^3 & 6xy^2 \\ 6xy^2 & 6x^2y - 4 \end{bmatrix}.$$

3. Let $f(x, y, z) = x^2 + y^2 + z^2$. Then

$$\nabla f(x, y, z) = \begin{bmatrix} \frac{\partial f}{\partial x} \\ \frac{\partial f}{\partial y} \\ \frac{\partial f}{\partial z} \end{bmatrix} = \begin{bmatrix} 2x \\ 2y \\ 2z \end{bmatrix}$$

see Example 105.3. Therefore

$$\mathbf{H}f(x, y, z) = \begin{bmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}.$$

4. Let $z = f(x, y) = e^x + e^y$ Then

$$\nabla z = \begin{bmatrix} e^x \\ e^y \end{bmatrix},$$

and

$$\mathbf{H}z = \begin{bmatrix} e^x & 0 \\ 0 & e^y \end{bmatrix}.$$

The second order partial derivative can be interpreted as the rate of change of the first order derivative. Also, it measures the convexity/concavity of a function with respect to one of the independent variables, i.e. along the direction of the corresponding axis.

Taylor expansion The Taylor expansion of a function about a given point which we studied in Section 1.5.1 also holds in the case of more than one variable. Thus, for $z = f(x, y)$ and a point (a, b) we have

$$\begin{aligned} f(x, y) &= f(a, b) + f_x(a, b)(x - a) + f_y(a, b)(y - b) \\ &\quad + \frac{1}{2!} (f_{xx}(a, b)(x - a)^2 + 2f_{xy}(a, b)(x - a)(y - b) + f_{yy}(a, b)(y - b)^2) + \dots \end{aligned}$$

The more general case of a function $f(x_1, \dots, x_n)$ is given by

$$\begin{aligned} f(x_1, \dots, x_n) &= f(a_1, \dots, a_n) + \sum_{j=1}^n \frac{\partial f(a_1, \dots, a_n)}{\partial x_j} (x_j - a_j) \\ &\quad + \frac{1}{2!} \sum_{j=1}^n \sum_{k=1}^n \frac{\partial^2 f(a_1, \dots, a_n)}{\partial x_j \partial x_k} (x_j - a_j)(x_k - a_k) + \dots \end{aligned}$$

We can write the above expressions in a more compact way by means of the gradient and the Hessian matrix, i.e. for $\mathbf{x} = (x_1, \dots, x_n)$, the Taylor expansion of $f(\mathbf{x})$ reads

$$f(\mathbf{x}) = f(\mathbf{a}) + (\mathbf{x} - \mathbf{a})^\top \nabla f(\mathbf{a}) + \frac{1}{2!} (\mathbf{x} - \mathbf{a})^\top \mathbf{H}f(\mathbf{a}) (\mathbf{x} - \mathbf{a}) + \dots$$

3.3 Minima and Maxima

Like in the case of functions of 1 variable, we can study the behaviour of functions in more than one variables to determine their minimum and maximum values. We proceed very much along the same line as in the case of 1 variable: first we find the partial derivatives and we seek for the stationary points, i.e. the values of the independent variables for which the partial derivatives are zero. Then, we obtain the Hessian matrix. Finally, we perform the second order test.

In details, and referring to a function of two variables, i.e. $z = f(x, y)$. The first step requires to find the stationary points $\mathbf{x}^* = (x^*, y^*)$ such that

$$\nabla f(\mathbf{x}^*) = 0.$$

This condition is known as the First Order Condition (FOC).

The stationary points can be classified as follows.

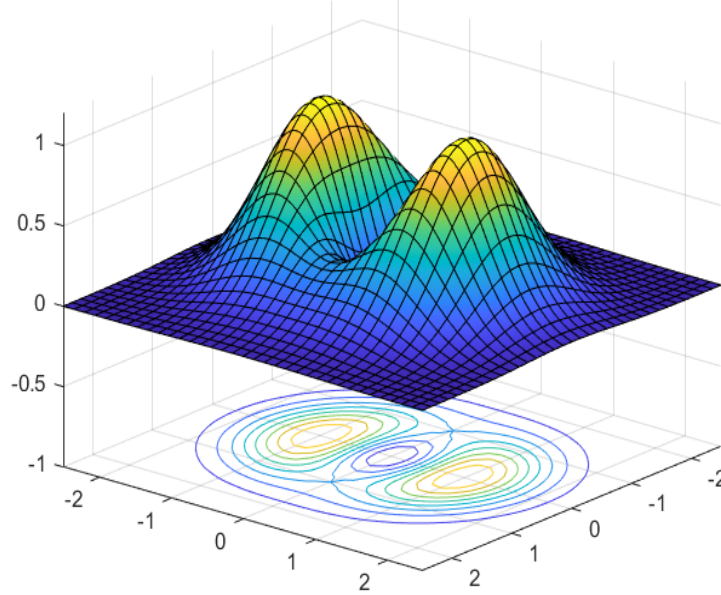


Figure 34: Local minima and maxima.

1. A function of two variables has a local maximum at (x^*, y^*) if $f(x, y) \leq f(x^*, y^*)$ when (x, y) is near (x^*, y^*) . An illustration is offered in Figure 34.
2. If $f(x, y) \geq f(x^*, y^*)$ when (x, y) is near (x^*, y^*) , then (x^*, y^*) is a local minimum value.
3. If the inequalities above hold for all points (x, y) in the domain of f , then f has an absolute maximum (or absolute minimum) at (x^*, y^*) . An illustration is offered in Figures 35 and 36.
4. If none of the above holds, the stationary point (x^*, y^*) is a saddle point. An example of saddle point is illustrated in Figure 37. Specifically, the stationary point in this case is $(0, 0)$; we note that in the x -direction the tangent line has negative sign for $x < 0$ and positive sign for $x > 0$ (red arrow) indicating a point of minimum. However, along the y -direction the tangent line is positive for $y < 0$ and negative for $y > 0$ (purple line) indicating a point of maximum. In this case, the first order test as we know it from Section 1.5.4 would be inconclusive. It is for this reason that in the case of functions of more than one variable, we need necessarily to study the second order partial derivatives.

The implementation of the second order test in the multidimensional case is based on the Hessian matrix. To see this, let us consider the Taylor expansion of $f(\mathbf{x})$ about the stationary point $\mathbf{x}^*$ up to the second order

$$f(\mathbf{x}) = f(\mathbf{x}^*) + \nabla f(\mathbf{x}^*)(\mathbf{x} - \mathbf{x}^*) + \frac{1}{2!}(\mathbf{x} - \mathbf{x}^*)^\top \mathbf{H}f(\mathbf{x}^*)(\mathbf{x} - \mathbf{x}^*).$$

As $\mathbf{x}^*$ is a stationary point, then

$$\nabla f(\mathbf{x}^*) = 0,$$

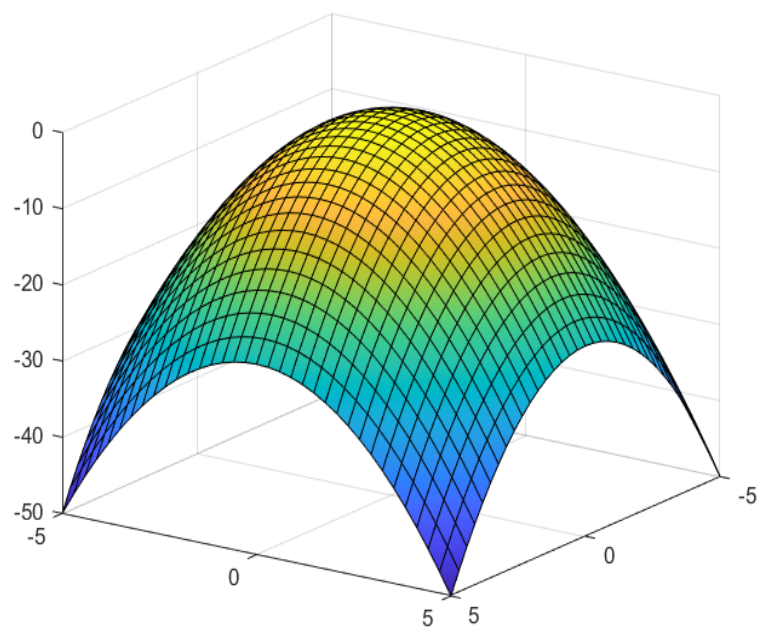


Figure 35: Global maximum.

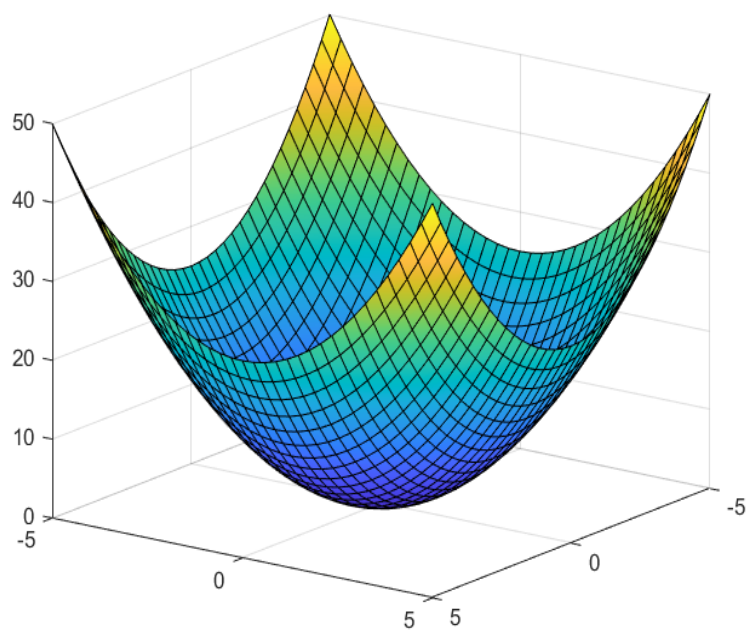


Figure 36: Global minimum.

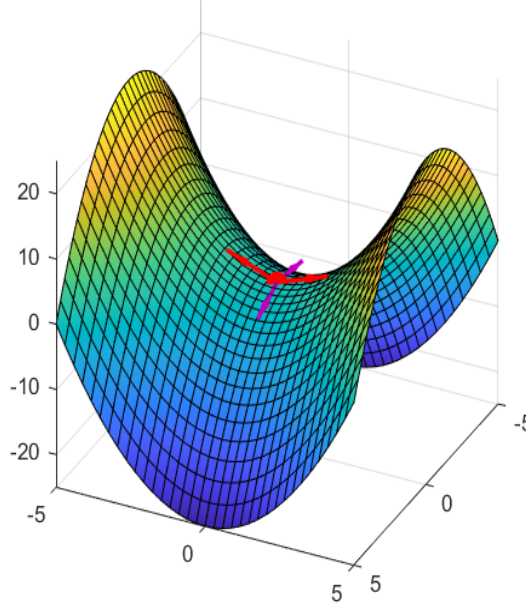


Figure 37: Saddle point.

and therefore

$$f(\mathbf{x}) = f(\mathbf{x}^*) + \frac{1}{2!}(\mathbf{x} - \mathbf{x}^*)^\top \mathbf{H}f(\mathbf{x}^*)(\mathbf{x} - \mathbf{x}^*).$$

Then, for $\mathbf{x}^*$ to be a maximum, we require $f(\mathbf{x}) \leq f(\mathbf{x}^*)$; but from the equation above, this holds only if

$$(\mathbf{x} - \mathbf{x}^*)^\top \mathbf{H}f(\mathbf{x}^*)(\mathbf{x} - \mathbf{x}^*) \leq 0.$$

Similarly, for $\mathbf{x}^*$ to be a minimum, we require $f(\mathbf{x}) \geq f(\mathbf{x}^*)$; but from the equation above, this holds only if

$$(\mathbf{x} - \mathbf{x}^*)^\top \mathbf{H}f(\mathbf{x}^*)(\mathbf{x} - \mathbf{x}^*) \geq 0.$$

However, the form

$$q(\mathbf{x}^*) = (\mathbf{x} - \mathbf{x}^*)^\top \mathbf{H}f(\mathbf{x}^*)(\mathbf{x} - \mathbf{x}^*)$$

is a quadratic form like the ones we have seen in Section 2.5. Hence, we can use Sylvester's criterion on the Hessian matrix to determine if the stationary point $\mathbf{x}^*$ is a maximum, a minimum or a saddle point.

Therefore, the second order test is formalized as follows.

Let $\mathbf{x}^*$ be a stationary point of $f(\mathbf{x})$. Let $\mathbf{H}f(\mathbf{x}^*)$ be the Hessian matrix of f in $\mathbf{x}^*$.

1. If $\mathbf{H}f(\mathbf{x}^*)$ is positive definite, i.e. if its leading principal minors are all positive, then $\mathbf{x}^*$ is a global minimum.
2. If $\mathbf{H}f(\mathbf{x}^*)$ is positive semi-definite, i.e. if its principal minors are all nonnegative, then $\mathbf{x}^*$ is a local minimum.

3. If $\mathbf{H}f(\mathbf{x}^*)$ is negative definite, i.e. $|\mathbf{H}_1| < 0$, $|\mathbf{H}_2| > 0$, $|\mathbf{H}_3| < 0, \dots$, and so on with alternate signs, then $\mathbf{x}^*$ is a global maximum.
4. If $\mathbf{H}f(\mathbf{x}^*)$ is negative semi-definite, i.e. if all its principal minors of even dimension are non-negative and all principal minors of odd dimension are non-positive, then $\mathbf{x}^*$ is a local maximum.
5. If f is a function of 2 variables, and if $\mathbf{H}f(\mathbf{x}^*)$ is indefinite, then $\mathbf{x}^*$ is a saddle point (although this is just a sufficient condition).

The above are known as the Second Order Conditions (SOC).

Example 109 1. Consider the function

$$f(x, y) = 4x^2 + 3y^2.$$

The gradient vector is

$$\nabla f(x, y) = \begin{bmatrix} 8x \\ 6y \end{bmatrix}.$$

Therefore, the FOC returns one stationary point at $\mathbf{x}^* = (x^*, y^*) = (0, 0)$. The Hessian matrix is

$$\mathbf{H}f(x, y) = \begin{bmatrix} 8 & 0 \\ 0 & 6 \end{bmatrix} = 2 \begin{bmatrix} 4 & 0 \\ 0 & 3 \end{bmatrix}.$$

We know from Exercise 100.1 that this matrix is positive definite. Therefore, the given function has a global minimum in the origin.

2. Consider the function $f(x, y) = -xy$. Then, by differentiation we obtain that the gradient vector is

$$\nabla f(x, y) = \begin{bmatrix} -y \\ -x \end{bmatrix}.$$

Therefore, the FOC returns one stationary point at $\mathbf{x}^* = (x^*, y^*) = (0, 0)$. The Hessian matrix is

$$\mathbf{H}f(x, y) = \begin{bmatrix} 0 & -1 \\ -1 & 0 \end{bmatrix}.$$

Note that $|\mathbf{H}_1| = 0$ and $|\mathbf{H}_2| = -2$. Then the stationary point is a saddle point. This function is plotted in Figure 31 - left hand side panel.

3. Consider the function $f(x, y) = \ln(xy)$. The corresponding gradient vector is

$$\nabla f(x, y) = \begin{bmatrix} \frac{1}{x} \\ \frac{1}{y} \end{bmatrix}.$$

Therefore, there are no stationary points. This function is plotted in Figure 31 - right hand side panel.

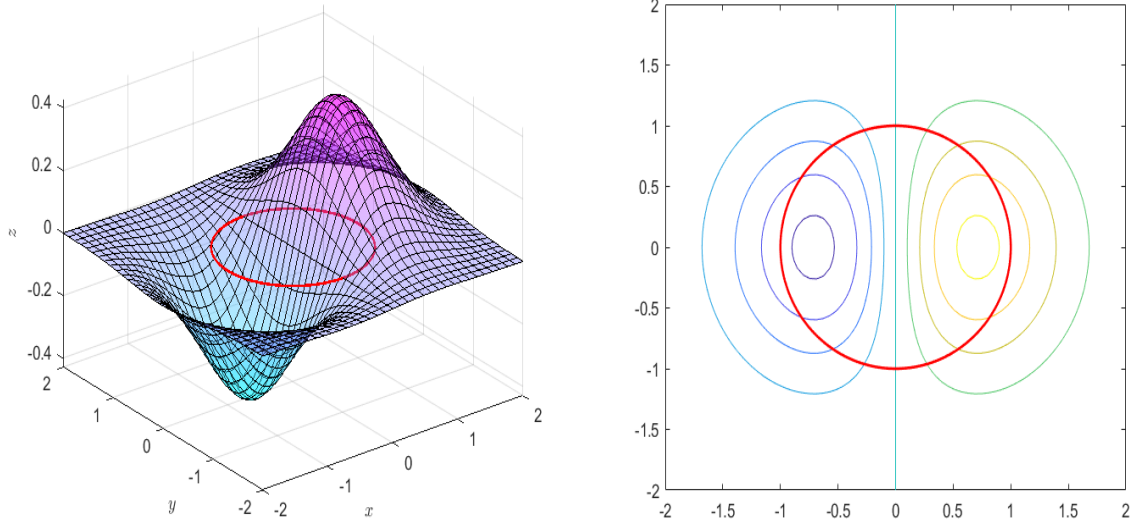


Figure 38: Constrained Optimization. Constraint: $g(x, y) = x^2 + y^2 = 1$ plotted in red. Left hand side panel: 3D representation. Right hand side panel: corresponding contour map.

4 Optimization and Applications in Economics and Finance

In the previous sections we learnt how to study minima and maxima of functions, a process also known as optimization. In the context of Finance and Economics optimization is at the core of many important applications; however, more often than not, it is carried out in presence of constraints. For example, consumers usually have budget limitations, in the sense that they can only use a limited part of resources for purchasing goods. Consequently, they need to choose the best allocation of their funds as to optimize their utility. Similarly, an asset manager has the task of choosing the allocation within the portfolio as to minimize the risk of this portfolio; the manager though has not infinite resources for this scope, and cannot sacrifice too much return on the portfolio returns either. This is another example of constrained optimization: find the optimal asset allocation subject to achieving a certain returns level and a given budget constraint.

In presence of constraints, the first and second order tests no longer work the way we previously saw; indeed, due to the presence of the constraint, the function might no longer achieve its optimum value in correspondence of the ‘natural’ (i.e. unconstrained) minima and/or maxima. This is illustrated in Figure 38: if we are only allowed to stay on the constraint (the red circle), we cannot reach the actual minimum and maximum of the function which lie away from the constraint.

Thus, the scope of this final section is to understand the basics of constrained optimization: how do we find minima/maxima of functions of one or more variables, $\mathbf{x} = x_1, x_2, \dots, x_n$, in presence of constraints on the independent variables. In general terms, the problem is formalized as follows

$$\begin{array}{ll} \max_{\mathbf{x}} f(\mathbf{x}) & \text{or} \quad \min_{\mathbf{x}} f(\mathbf{x}) \\ \text{s.t } g(\mathbf{x}) = c & \text{s.t } g(\mathbf{x}) = c \end{array}$$

according to whether the task is to maximize or to minimize. The function $f(\mathbf{x})$ is the objective function, the function $g(\mathbf{x})$ determines the constraint. We could have more than one constraint. Further, in the following we will consider only equality constraints, i.e. of the form

$$g(\mathbf{x}) = c.$$

However, a more general theory can be developed for inequality constraints of the form

$$g(\mathbf{x}) \leq c,$$

which we do not cover here.

In order to keep the presentation simple and intuitive, we work in the setting of a function of 2 independent variables, $f(x, y)$, and one equality constraint. The main results hold also in case of more variables and constraints.

We have different ways of solving a constrained optimization problem; the choice of the method depends heavily on the nature of the constraint and the number of variables involved. The methods in question are substitution, implicit differentiation (yes: we have met this already), and the more general Lagrange multiplier method - named after the mathematician Joseph-Louis Lagrange.

4.1 Solving by substitution

In certain instances, it is possible to write the constraint in explicit form, i.e. you can explicit y as a function of x just based on the constraint. Then, you can substitute into the objective function, and then proceed with standard optimization (unconstrained) for the new objective function.

Let us illustrate this procedure in an example.

Example 110 Consider a consumer with utility function

$$U(x, y) = xy + 2x,$$

and a budget of £60 for expenditure. The prices of the two goods are $P_x = 4$ and $P_y = 2$. Determine the consumer's optimal bundle and the corresponding value of the utility function. Given the information above, this is equivalent to say that the budget line has equation $4x + 2y = 60$. Therefore, the optimization problem can be formalized as

$$\begin{aligned} \max_{x, y} U(x, y) &= xy + 2x \\ \text{s.t.} \quad 4x + 2y &= 60. \end{aligned}$$

In order to tackle the above problem by substitution, let us explicit y in terms of x in the constraint, i.e.

$$y = 30 - 2x,$$

and then substitute in the utility function

$$U(x, y) = xy + 2x = x(30 - 2x) + 2x = 32x - 2x^2 = u(x).$$

In this way, we reduced the number of independent variables to 1, i.e. x . The optimization problem then can be rewritten as

$$\max_x u(x) = 32x - 2x^2.$$

Then, we can use the results from Calculus in 1 variable: the first order condition reads

$$u'(x) = 32 - 4x = 0$$

which returns $x = 8$ which is the stationary point. Now check the second order condition

$$u''(x) = -4 < 0,$$

which implies that the stationary point is a point of maximum. In order to obtain also the y coordinate, substitute in the constraint to obtain $y = 30 - 2 \times 8 = 14$.

In conclusion, the constrained optimization problem

$$\begin{aligned} \max_{x,y} U(x,y) &= xy + 2x \\ \text{s.t.} \quad 4x + 2y &= 60 \end{aligned}$$

has solution in $(x^*, y^*) = (8, 14)$. The value of the objective function in the optimum point is $U(x^*, y^*) = 128$.

The substitution method is quite simple to implement. However, when functions are more complex, or when there are several constraints to consider, the technique of substitution could become burdensome.

4.2 Solving by implicit differentiation

A potential alternative method uses the technique of implicit differentiation, based on the idea that in the point of optimum the constraint and the objective function should be tangent to each other. Let us consider again the previous Example 110 and solve it now by implicit differentiation - this will also prepare the ground for the more general Lagrange multiplier method.

Example 111 The problem in question is

$$\begin{aligned} \max_{x,y} U(x,y) &= xy + 2x \\ \text{s.t.} \quad 4x + 2y &= 60. \end{aligned}$$

As the constraint defines y as a function of x which we now consider as implicit, we can obtain the slope of the budget line by implicit differentiation, so that

$$2 + y' = 0,$$

i.e.

$$y' = -2.$$

Now, let's obtain the information about the slope of the tangent to the utility function for some value $\bar{u}$ of $U(x, y)$, i.e. set

$$\bar{u} = xy + 2x$$

then differentiate so that

$$y + xy' + 2 = 0$$

from which we obtain

$$y' = \frac{-2 - y}{x}.$$

For tangency we require the slopes to be the same, therefore

$$\frac{-2 - y}{x} = -2$$

which leads to

$$y = 2x - 2.$$

This line is the locus of the points such that the slope of the tangent line to the indifference curves y' is equal to the slope of the budget line -2 . In order to find the optimal bundle we can look for the intersection between the budget line and the locus of points we have just identified, as to place ourselves directly on the ‘frontier’ so that all of our budget is used. The solution leads to $x = 8$ and $y = 14$.

Once again, in presence of more complex functions with more variables, or several constraints, this technique could become cumbersome and actually quite difficult to implement. However, the idea of implicit differentiation helps us to build the Lagrange multiplier method explained in the next section.

4.3 Solving by Lagrange multipliers

4.3.1 An intuitive take

Let us formalize the implicit differentiation method seen in Example 111.

We start from the optimization problem stated in its general form

$$\begin{aligned} \max_{\mathbf{x}} (\min) f(\mathbf{x}) \\ \text{s.t. } g(\mathbf{x}) = c. \end{aligned}$$

The constraint $g(x, y) = c$ defines y implicitly; assuming $g(x, y)$ is differentiable, then by implicit differentiation as in equation (10) we obtain

$$g_x + g_y y' = 0$$

(remember to use the chain rule), from which it follows that

$$y' = -\frac{g_x}{g_y}.$$

This is the ‘slope’ of the tangent to the constraint for every possible value of x and y .

We can repeat the same as above for the objective function, bearing in mind that y is an implicit function of x . Thus, choose an arbitrary value $\bar{z}$ and set

$$\bar{z} = f(x, y),$$

then once again by implicit differentiation

$$0 = f_x + f_y y'.$$

In order to find the point of tangency between the objective function and the constraint, we need to use y' obtained from the constraint, i.e.

$$0 = f_x - f_y \frac{g_x}{g_y}. \quad (11)$$

Now, set

$$\lambda = \frac{f_y}{g_y}, \quad (12)$$

then the expressions (11) and (12) can be rewritten as

$$\begin{aligned} f_x - \lambda g_x &= 0 \\ f_y - \lambda g_y &= 0. \end{aligned}$$

Thus, the point of constrained optimum must satisfy the conditions

$$f_x - \lambda g_x = 0 \quad (13)$$

$$f_y - \lambda g_y = 0 \quad (14)$$

$$g(x, y) = c. \quad (15)$$

Conditions (13)-(14) ensure that the gradient of f is the same as the gradient of g up to a constant λ , which is known as the Lagrange multiplier. Condition (15) ensures that the equality constraint is satisfied.

The above could be seen as equivalent to the first order conditions of a new objective function

$$\mathcal{L}(x, y, \lambda) = f(x, y) - \lambda(g(x, y) - c); \quad (16)$$

the function $\mathcal{L}(x, y, \lambda)$ is known as the Lagrangean function.

As we know from Section 3.3, the first order conditions only identify the stationary point, but they do not classify it, i.e. they do not tell us if the point is a minimum, a maximum or a saddle. We require the second order conditions for this purpose. But we have a constraint. So, the Hessian matrix of f is not going to help us.

On the other hand, we know by the previous argument that y is an implicit function of x defined through the constraint, and therefore the objective function is in reality a function of x only. Set $z = f(x, y)$. We know the first order derivative of this function, in which we can incorporate the constraint condition via y' , i.e.

$$\frac{dz}{dx} = f_x - f_y \frac{g_x}{g_y}.$$

Let us then differentiate the above - using again implicit differentiation and the chain rule - to obtain

$$\begin{aligned} \frac{d^2 z}{dx^2} &= f_{xx} + f_{xy}y' \\ &\quad - (f_{yx} + f_{yy}y') \frac{g_x}{g_y} - f_y \frac{(g_{xx} + g_{xy}y')g_y - g_x(g_{yx} + g_{yy}y')}{g_y^2} \\ &= \frac{1}{g_y^2} ((f_{xx} - \lambda g_{xx})g_y^2 + (f_{yy} - \lambda g_{yy})g_x^2 - 2(f_{xy} - \lambda g_{xy})g_x g_y), \end{aligned}$$

where we used the expression for y' and the definition of the Lagrange multiplier given by equation (12). The stationary point then is a maximum if

$$\frac{d^2 z}{dx^2} < 0,$$

i.e.

$$(f_{xx} - \lambda g_{xx})g_y^2 + (f_{yy} - \lambda g_{yy})g_x^2 - 2(f_{xy} - \lambda g_{xy})g_x g_y < 0.$$

Admittedly the condition above is quite lengthy, but we could remember it more easily if we see it as the determinant of the matrix

$$- \begin{bmatrix} f_{xx} - \lambda g_{xx} & f_{xy} - \lambda g_{xy} & -g_x \\ f_{xy} - \lambda g_{xy} & f_{yy} - \lambda g_{yy} & -g_y \\ -g_x & -g_y & 0 \end{bmatrix}$$

which can be easily verified via the Sarrus rule.

If we look again at the Lagrangean function defined in equation (16), we notice that its Hessian matrix is actually

$$\mathbf{H}\mathcal{L}(x, y, \lambda) = \begin{bmatrix} f_{xx} - \lambda g_{xx} & f_{xy} - \lambda g_{xy} & -g_x \\ f_{xy} - \lambda g_{xy} & f_{yy} - \lambda g_{yy} & -g_y \\ -g_x & -g_y & 0 \end{bmatrix}.$$

Therefore we conclude that the stationary point identified by the first order conditions (13)-(15) is a maximum if

$$\frac{d^2 z}{dx^2} < 0,$$

i.e.

$$\frac{d^2 z}{dx^2} = -|\mathbf{H}\mathcal{L}(x, y, \lambda)| < 0,$$

or, equivalently

$$|\mathbf{H}\mathcal{L}(x, y, \lambda)| > 0.$$

For the minimum, we ask for

$$|\mathbf{H}\mathcal{L}(x, y, \lambda)| < 0.$$

The matrix $\mathbf{H}\mathcal{L}$ is known as the bordered Hessian.

Thus, the procedure to find the constrained optima is as follows.

1. Write the Lagrangean function

$$\mathcal{L}(x, y, \lambda) = f(x, y) - \lambda(g(x, y) - c).$$

2. Impose the FOC to identify the stationary points

$$\nabla \mathcal{L}_{x,y} = 0; \quad -g(x, y) + c = 0.$$

3. Obtain the bordered Hessian

$$\mathbf{H}\mathcal{L}(x, y, \lambda) = \begin{bmatrix} f_{xx} - \lambda g_{xx} & f_{xy} - \lambda g_{xy} & -g_x \\ f_{xy} - \lambda g_{xy} & f_{yy} - \lambda g_{yy} & -g_y \\ -g_x & -g_y & 0 \end{bmatrix},$$

and compute its determinant. If $|\mathbf{H}\mathcal{L}| > 0$, the stationary point is a local constrained maximum. If $|\mathbf{H}\mathcal{L}| < 0$, the stationary point is a local constrained minimum.

Example 112 Let us consider once again Example 110. Let us build the Lagrangean function

$$\mathcal{L}(x, y, \lambda) = xy + 2x - \lambda(4x + 2y - 60).$$

The first order conditions are

$$\begin{aligned}\mathcal{L}_x &= y + 2 - 4\lambda = 0, & \text{i.e. } y &= 4\lambda - 2 \\ \mathcal{L}_y &= x - 2\lambda = 0, & \text{i.e. } x &= 2\lambda \\ 4x + 2y &= 60, & \text{i.e. } 16\lambda &= 64\end{aligned}$$

so that

$$(x^*, y^*) = (8, 14), \quad \lambda^* = 4.$$

Furthermore

$$\mathbf{H}\mathcal{L} = \begin{bmatrix} 0 & 1 & -4 \\ 1 & 0 & -2 \\ -4 & -2 & 0 \end{bmatrix};$$

consequently $|\mathbf{H}\mathcal{L}| = 16 > 0$, and therefore $(x^*, y^*) = (8, 14)$ is the local maximum.

4.3.2 Method of Lagrange multipliers: the full formulation

Although the treatment of the constrained optimization problem via the method of the Lagrange multiplier has been given for the case of 2 variables and 1 equality constraint, the argument essentially is identical for more variables and more equality constraints. We need ‘only’ to slightly generalize the SOC to a formulation in spirit not too dissimilar from the Sylvester criterion. We summarize all of our findings together with this much needed generalization in the following.

Theorem 32 (The method of Lagrange multipliers) *Consider the general problem of constrained optimization*

$$\begin{aligned}\max_{\mathbf{x}} & (\min) f(\mathbf{x}) \\ \text{s.t. } & g_1(\mathbf{x}) = c_1 \\ & g_2(\mathbf{x}) = c_2 \\ & \vdots \\ & g_m(\mathbf{x}) = c_m\end{aligned}$$

for $\mathbf{x} = x_1, x_2, \dots, x_n$ and differentiable functions $f, g_1, \dots, g_m$. Let $\boldsymbol{\lambda} = \lambda_1, \lambda_2, \dots, \lambda_m$, and define the Lagrangean function

$$\mathcal{L}(\mathbf{x}, \boldsymbol{\lambda}) = f(\mathbf{x}) - \lambda_1(g_1(\mathbf{x}) - c_1) - \dots - \lambda_m(g_m(\mathbf{x}) - c_m).$$

1. The optimum $\mathbf{x}^*$ to the given constrained optimization problem is the solution to the first order conditions

$$\begin{aligned}
 \nabla_{\mathbf{x}} \mathcal{L}(\mathbf{x}, \boldsymbol{\lambda}) &= 0 \\
 g_1(\mathbf{x}) &= c_1 \\
 g_2(\mathbf{x}) &= c_2 \\
 &\vdots \\
 g_m(\mathbf{x}) &= c_m
 \end{aligned} \tag{17}$$

where $\nabla_{\mathbf{x}} \mathcal{L}$ denotes the gradient vector containing the first order partial derivatives with respect to $x_1, x_2, \dots, x_n$ of the Lagrangean function.

2. The optimum $\mathbf{x}^*$ is a point of local (constrained) maximum if $|\mathbf{H}\mathcal{L}|$ has the same sign as $(-1)^n$ and the remaining first $n - m - 1$ principal minors of the bordered Hessian $\mathbf{H}\mathcal{L}$ alternate in sign.
3. The optimum $\mathbf{x}^*$ is a point of local (constrained) minimum if $|\mathbf{H}\mathcal{L}|$ and the remaining first $n - m - 1$ principal minors of the bordered Hessian $\mathbf{H}\mathcal{L}$ have the same sign as $(-1)^m$.

It is understood that if none of the above mentioned second order conditions is satisfied, the stationary point cannot be classified.

Example 113 1. Consider the problem

$$\begin{aligned}
 \max_{x,y} \quad & x^2 + y^2 \\
 \text{s.t} \quad & x^2 + xy + y^2 = 3.
 \end{aligned}$$

In order to solve it, let us first build the Lagrangean function

$$\mathcal{L}(x, y, \lambda) = x^2 + y^2 - \lambda(x^2 + xy + y^2 - 3).$$

Then, let us derive the first order conditions

$$\begin{aligned}
 \frac{\partial \mathcal{L}}{\partial x} &= 2x - \lambda(2x + y) = 0 \\
 \frac{\partial \mathcal{L}}{\partial y} &= 2y - \lambda(x + 2y) = 0 \\
 x^2 + xy + y^2 &= 3.
 \end{aligned}$$

Consider the first two conditions: by rearranging we obtain

$$\frac{x}{2x + y} = \frac{y}{x + 2y}$$

which implies $y^2 = x^2$, i.e. $y = |x|$. By substituting into the constraint, we obtain

$$2x^2 + x|x| = 3.$$

If $y = x$, then $x = \pm 1$ and therefore the stationary points are $P_1 = (-1, -1)$, $P_2 = (1, 1)$, as

$$\lambda = \frac{2x}{2x + y},$$

the Lagrange multiplier is therefore $2/3$ in both cases. If $y = -x$, then the stationary points are $P_3 = (-\sqrt{3}, \sqrt{3})$ and $P_4 = (\sqrt{3}, -\sqrt{3})$. The Lagrange multiplier is in both cases 2 .

Note that another possible solution to $y = |x|$ is the origin, i.e. $(0, 0)$; however, this is not an admissible solution as it does not satisfy the constraint.

Let us now consider the second order conditions, in order to be able to classify the four stationary points we have just obtained. To begin with, let us derive the bordered Hessian matrix

$$\mathbf{H}\mathcal{L} = \begin{bmatrix} 2 - 2\lambda & -\lambda & -2x - y \\ -\lambda & 2 - 2\lambda & -x - 2y \\ -2x - y & -x - 2y & 0 \end{bmatrix}.$$

We have 2 independent variables, i.e. $n = 2$, and 1 constraint, i.e. $m = 1$. Thus we need to check only $n - m = 1$ principal minor, which is the determinant of the full matrix. For P_1, P_2 , $\lambda = 2/3$, therefore

$$\mathbf{H}\mathcal{L}|_{P_1} = \begin{bmatrix} \frac{2}{3} & -\frac{2}{3} & 3 \\ -\frac{2}{3} & \frac{2}{3} & 3 \\ 3 & 3 & 0 \end{bmatrix} \quad \mathbf{H}\mathcal{L}|_{P_2} = \begin{bmatrix} \frac{2}{3} & -\frac{2}{3} & -3 \\ -\frac{2}{3} & \frac{2}{3} & -3 \\ -3 & -3 & 0 \end{bmatrix}.$$

The determinant is -24 in both cases (note that the bordered Hessian in P_2 is the same as in P_1 with both row 3 and column 3 changed of sign - see Properties of the determinant). Thus, it is negative, i.e. it has the same sign as $(-1)^m = (-1)^1 = -1 < 0$. Therefore P_1 and P_2 are both local constrained minima.

For P_3, P_4 , $\lambda = 2$, therefore

$$\mathbf{H}\mathcal{L}|_{P_3} = \begin{bmatrix} -2 & -2 & \sqrt{3} \\ -2 & -2 & -\sqrt{3} \\ \sqrt{3} & -\sqrt{3} & 0 \end{bmatrix} \quad \mathbf{H}\mathcal{L}|_{P_4} = \begin{bmatrix} -2 & -2 & -\sqrt{3} \\ -2 & -2 & \sqrt{3} \\ -\sqrt{3} & \sqrt{3} & 0 \end{bmatrix};$$

the determinant is 24 in both cases, i.e. it is positive. This is the same sign as $(-1)^n = (-1)^2 = 1$. Therefore P_3 and P_4 are both local constrained maxima. In conclusion, the given problem has solution in P_3 and P_4 as these are the only (local) constrained maxima; furthermore the optimal value of the objective function is 6 in both points.

2. Consider the problem

$$\begin{aligned} & \min_{x,y,z} x^2 + y^2 + z^2 \\ & \text{s.t } x + y + z = 1. \end{aligned}$$

As above, let us first build the Lagrangean function

$$\mathcal{L}(x, y, \lambda) = x^2 + y^2 + z^2 - \lambda(x + y + z - 1).$$

Then, let us derive the first order conditions

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial x} &= 2x - \lambda &= 0 \\ \frac{\partial \mathcal{L}}{\partial y} &= 2y - \lambda &= 0 \\ \frac{\partial \mathcal{L}}{\partial z} &= 2z - \lambda &= 0 \\ x + y + z &= 1. \end{aligned}$$

The first three conditions imply $\lambda = 2x = 2y = 2z$ i.e. $x = y = z$. By substituting into the constraint, we obtain $3x = 1$, and therefore

$$x = \frac{1}{3}.$$

Consequently the stationary point is

$$P = \left(\frac{1}{3}, \frac{1}{3}, \frac{1}{3} \right), \quad \lambda = \frac{2}{3}.$$

Let us now consider the second order conditions, in order to verify that the stationary point we have just obtained is indeed a minimum. To begin with, let us derive the bordered Hessian matrix

$$\mathbf{H}\mathcal{L} = \begin{bmatrix} 2 & 0 & 0 & -1 \\ 0 & 2 & 0 & -1 \\ 0 & 0 & 2 & -1 \\ -1 & -1 & -1 & 0 \end{bmatrix}.$$

We have 3 independent variables, i.e. $n = 3$, and 1 constraint, i.e. $m = 1$. Thus we need to check $n - m = 2$ principal minor. The first one is the determinant of the full matrix. We use the method of co-factors working on, for example, the first column, so that

$$\begin{aligned} |\mathbf{H}\mathcal{L}| &= (-1)^{1+1} \times 2 \times \begin{vmatrix} 0 & 0 & -1 \\ -1 & -1 & 0 \end{vmatrix} + (-1)^{1+4}(-1) \begin{vmatrix} 0 & 0 & -1 \\ 0 & 2 & -1 \end{vmatrix} \\ &= -12 < 0. \end{aligned}$$

We note that the sign of $|\mathbf{H}\mathcal{L}|$ is the same as the one of $(-1)^m = (-1)^1 = -1 < 0$.

The second principal minor is the one we obtain from the bordered Hessian by removing

the first row and the first column (see Definition 24), i.e. from the matrix. Consequently

$$M_{1,1}(\mathbf{H}\mathcal{L}) = \begin{vmatrix} 2 & 0 & -1 \\ 0 & 2 & -1 \\ -1 & -1 & 0 \end{vmatrix} = -4 < 0.$$

Thus, $M_{1,1}$ has the same sign as $|\mathbf{H}\mathcal{L}|$. By Theorem 32 the point P is indeed a local constrained minimum. The optimal value of the objective function is $1/3$.

3. Consider the problem

$$\begin{aligned} & \max_{x,y,z} xyz \\ & \text{s.t. } x^2 + y^2 = 1 \\ & \quad x + z = 1. \end{aligned}$$

As above, let us first build the Lagrangean function

$$\mathcal{L}(x, y, z, \lambda_1, \lambda_2) = xyz - \lambda_1(x^2 + y^2 - 1) - \lambda_2(x + z - 1).$$

Then, let us derive the first order conditions

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial x} &= yz - 2\lambda_1 x - \lambda_2 = 0 \\ \frac{\partial \mathcal{L}}{\partial y} &= xz - 2\lambda_1 y = 0 \\ \frac{\partial \mathcal{L}}{\partial z} &= xy - \lambda_2 = 0 \\ & x^2 + y^2 = 1 \\ & x + z = 1. \end{aligned}$$

The second and third conditions imply

$$2\lambda_1 = \frac{xz}{y}, \quad \lambda_2 = xy.$$

By substituting into the first condition we obtain

$$y^2 z - x^2 z - xy^2 = 0.$$

From the constraints we know: $y^2 = 1 - x^2$ and $z = 1 - x$. Substitute in the above and rearrange to obtain

$$(1 - x)(1 - x - 3x^2) = 0.$$

The quadratic equation has solutions

$$x = \frac{-1 \pm \sqrt{13}}{6}.$$

Consequently the stationary points are

$$\begin{aligned}
P_1 &= (-0.7676, -0.6409, 1.7676) & (\lambda_1, \lambda_2) &= (1.0585, 0.4920) \\
P_2 &= (-0.7676, 0.6409, 1.7676) & (\lambda_1, \lambda_2) &= (-1.0585, -0.4920) \\
P_3 &= (0.4343, -0.9008, 0.5657) & (\lambda_1, \lambda_2) &= (-0.1364, -0.3912) \\
P_4 &= (0.4343, 0.9008, 0.5657) & (\lambda_1, \lambda_2) &= (0.1364, 0.3912).
\end{aligned}$$

Note that there is a fifth solution to the equation above, i.e. $x = 1$ however this is not admissible: indeed we would have $y = z = 0$ which means that λ_1 is not defined.

Let us now consider the second order conditions, in order to classify the stationary points obtained above. To begin with, let us derive the bordered Hessian matrix

$$\mathbf{H}\mathcal{L} = \begin{bmatrix} -2\lambda_1 & z & y & -2x & -1 \\ z & -2\lambda_1 & x & -2y & 0 \\ y & x & 0 & 0 & -1 \\ -2x & -2y & 0 & 0 & 0 \\ -1 & 0 & -1 & 0 & 0 \end{bmatrix}.$$

We have 3 independent variables, i.e. $n = 3$, and 2 constraint, i.e. $m = 2$. Thus we need to check $n - m = 1$ principal minor, which means we just need to compute the determinant of the bordered Hessian in correspondence of the 4 stationary points obtained from the FOC. The method of co-factors is applied here - or, use Matlab as this is a 5×5 matrix - some software is needed. In either cases, we obtain

$$\begin{aligned}
|\mathbf{H}\mathcal{L}|_{P_1} &= -16.3393 < 0, \text{ and } (-1)^n = (-1)^3 = -1 < 0, \\
|\mathbf{H}\mathcal{L}|_{P_2} &= 16.3393 > 0, \text{ and } (-1)^m = (-1)^2 = 1 > 0, \\
|\mathbf{H}\mathcal{L}|_{P_3} &= 7.3501 > 0, \text{ and } (-1)^m = (-1)^2 = 1 > 0, \\
|\mathbf{H}\mathcal{L}|_{P_4} &= -7.3501 < 0, \text{ and } (-1)^n = (-1)^3 = -1 < 0.
\end{aligned}$$

By Theorem 32, P_1 and P_4 are the candidate constrained maximizers. To choose which one, we need to evaluate the objective function in both points. So, for $f(x, y, z) = xyz$, we obtain

$$P_1 : f(-0.7676, -0.6409, 1.7676) = 0.8696; \quad P_4 : f(0.4343, 0.9008, 0.5657) = 0.2213.$$

Therefore the local constrained maximum is obtained in P_1 and the optimum value of the objective function is 0.8696.

A note on the bordered Hessian In many textbooks the bordered Hessian is written as

$$\mathbf{H}\mathcal{L}(x, y, \lambda) = \begin{bmatrix} 0 & -g_x & -g_y \\ -g_x & f_{xx} - \lambda g_{xx} & f_{xy} - \lambda g_{xy} \\ -g_y & f_{xy} - \lambda g_{xy} & f_{yy} - \lambda g_{yy} \end{bmatrix}.$$

However, for the case of 2 independent variables and 1 constraint this formulation is completely equivalent to the one provided in these lecture notes. Indeed, as we are required to compute determinants, by properties of the determinant (see Section 2.2), we know that the determinant does not change if we exchange first column 1 with column 2, then column 2 with column 3, then row 1 with row 2, and finally row 2 with row 3.

For the more general case, we need to adjust the second order conditions in Theorem 32 as follows.

2. The optimum $\mathbf{x}^*$ is a point of (constrained) local maximum if the last $n - m$ leading principal minors of the bordered Hessian $\mathbf{H}\mathcal{L}$ alternate in sign, with $|\mathbf{H}\mathcal{L}|$ having the same sign as $(-1)^n$.
- 3 The optimum $\mathbf{x}^*$ is a point of (constrained) local minimum if $|\mathbf{H}\mathcal{L}|$ and the last $n - m$ leading principal minors of the bordered Hessian $\mathbf{H}\mathcal{L}$ have the same sign as $(-1)^m$.

4.3.3 Meaning of the Lagrange multiplier: the shadow price

Consider again the constrained optimization problem

$$\begin{array}{ll} \max_{\mathbf{x}} f(\mathbf{x}) & \text{or} \quad \min_{\mathbf{x}} f(\mathbf{x}) \\ \text{s.t } g(\mathbf{x}) = c & \text{s.t } g(\mathbf{x}) = c \end{array}.$$

Suppose the solution of this problem is given by the point (x^*, y^*) . In general, (x^*, y^*) depend on c , i.e. $x^* = x^*(c)$ and $y^* = y^*(c)$. We assume that these solutions are differentiable functions of c . The associated value of $f(x, y)$ is then also a function of c , i.e.

$$f^*(c) = f(x^*(c), y^*(c))$$

is the (optimal) value function for the problem. Of course, the associated value of the Lagrange multiplier also depends on c , in general, so we write $\lambda(c)$.

Now. Take the total differential of $f^*(c)$

$$df^*(c) = df(x^*(c), y^*(c)) = f_x(x^*(c), y^*(c))dx^* + f_y(x^*(c), y^*(c))dy^*.$$

Due to the first order conditions, the above can be rewritten as

$$\begin{aligned} df^*(c) &= \lambda g_x(x^*(c), y^*(c))dx^* + \lambda g_y(x^*(c), y^*(c))dy^* \\ &= \lambda (g_x(x^*(c), y^*(c))dx^* + g_y(x^*(c), y^*(c))dy^*) \\ &= \lambda dc, \end{aligned}$$

where the last equality is due to the fact that in (x^*, y^*)

$$g(x^*(c), y^*(c)) = c$$

and therefore

$$dg(x^*(c), y^*(c)) = g_x(x^*(c), y^*(c))dx^* + g_y(x^*(c), y^*(c))dy^* = dc.$$

Thus, the Lagrange multiplier $\lambda = \lambda(c)$ is the rate at which the optimal value of the objective function changes with respect to changes in the constraint constant c . In particular, if dc is a small change in c , then

$$f^*(c + dc) - f^*(c) \approx \lambda(c)dc. \quad (18)$$

In the context of economic and financial applications, c often denotes the available stock of some resource, and $f(x, y)$ denotes utility or profit. Then, $\lambda(c)dc$ measures the approximate change in utility or profit that can be obtained from dc units more. Therefore, λ is known as the shadow price of the resource - in this case the available total amount c .

If $f^*(c)$ is the maximum profit when the resource input is c , then equation (18) says that λ quantifies the approximate increase in profit per unit increase in the resource.

Example 114 Consider once more Example 112. We obtained the optimal bundle at $(x^*, y^*) = (8, 14)$ with $\lambda^* = 4$ and $U(x^*, y^*) = U(8, 14) = 128$ (see also Example 110). What would happen if the constraint is relaxed as the available budget becomes 61 instead of 60? We could recalculate the solution - we would obtain $(x^{**}, y^{**}) = (8.125, 14.25)$ and $U(x^{**}, y^{**}) = U(8.125, 14.25) = 132.0313$. Alternatively, we could look at the Lagrange multiplier: $\lambda^* = 4$. Based on equation (18) we would approximate the improvement in the utility function to be $\lambda^* \times 1$, i.e. 4, which would bring the optimal utility value from 128 to $128 + 4 = 132$. Not a bad approximation!

5 A simple overview of Ordinary Differential Equations (ODEs)

A differential equation is an equation for an unknown function that contains not only the function but also its derivatives. In general, the unknown function may depend on several variables and the equation may include various partial derivatives - this is what is called a Partial Differential Equation (PDE). Differential equations for a function of a single real variable are called Ordinary Differential Equations (ODE).

A general form for ODEs is

$$F(x, y, y', \dots, y^{(n)}) = 0,$$

where F is a given function of $n + 2$ variables and $y = y(x)$ is an unknown function of a real variable x . The maximal order n of the derivative $y^{(n)}$ in the above equation is called the order of the ODE.

Differential equations arise in Finance in many applications. ODEs generally are key to determine the characteristic function of affine models for stochastic volatility and interest rates. PDEs instead arise in problems concerning the pricing and hedging of non-linear derivatives contracts such as options. It has to be noted that in the majority of these applications, solutions to O/PDEs can only be found numerically. Nevertheless, knowledge of the basics to solve simple ODEs can prove very useful for tackling certain problems arising for example in Stochastic Analysis. These are the cases of first order ODEs - linear and quasi linear - and are the ones we focus on in this section.

5.1 First order linear ODEs

Consider the ODE of the form

$$y' + a(x)y = b(x) \tag{19}$$

where a and b are given functions of x , defined on a certain interval I . This equation is called linear because it depends linearly on y and y' . In the case in which $b(x) = 0$, then the ODE reduces to

$$y' + a(x)y = 0$$

which is called homogeneous - this indicates that the ODE does not have terms which do not contain y or y' . In contrast, the ODE (19) is called inhomogeneous.

Example 115 1. The ODE

$$y' = y$$

is a homogeneous first order linear ordinary differential equation. The above can be alternatively written as

$$\frac{dy}{dx} = y$$

or, equivalently

$$dy = ydx.$$

2. The ODE

$$y = -xy'$$

is again a homogeneous first order linear ordinary differential equation. It could also be interpreted as a first order PDE, by writing the above as

$$\begin{aligned}y + x \frac{dy}{dx} &= 0 \\ ydx + xdy &= 0.\end{aligned}$$

The last equation above indicates that we are looking for a function $F(x, y)$ such that $F_x(x, y) = y$ and $F_y(x, y) = x$.

The above examples suggest that in order to solve differential equations, we need to find primitive functions. In other words, differential equations can be considered as shorthand notation for... integrals! Like with integrals, we might have a family of solutions due to the presence of a ‘constant of integration’, and we could have ‘particular’ solutions if sufficient initial conditions are specified to identify the value of such constant. This consideration leads to the following formalization.

Let $(x, y) \in \mathbb{R}^2$, and consider a function f with domain $D \in \mathbb{R}^2$ such that

$$y' = f(x, y). \quad (20)$$

Then.

Definition 33 *A real valued function $y(x)$ defined on an interval $I \in \mathbb{R}$, is called a particular solution of (20) if $y(x)$ is differentiable at any $x \in I$, the point $(x, y(x))$ belongs to D for any $x \in I$ and the identity $y'(x) = f(x, y(x))$ holds for all $x \in I$.*

The family of all particular solutions of (20) is called the general solution. The graph of a particular solution is called an integral curve of the equation. Any integral curve is contained in the domain D .

The ODEs in Example 115 show that direct application of integration might not be the best way forward. In the following, we explore a few methods, which might have to be used in combination with each other according to the given circumstances.

5.1.1 Solution of homogenous linear ODEs: Separation of variables

Consider a separable ODE, that is, an ODE of the form

$$y' = f(x)g(y). \quad (21)$$

Clearly, homogenous first order linear ODEs are examples of separable ODEs. The class of separable ODEs however is more general.

Any separable equation can be solved by means of the following theorem.

Theorem 34 (The method of separation of variables) *Let $f(x)$ and $g(y)$ be continuous functions on open intervals I and J , respectively, and assume that $g(y) \neq 0$ on J . Let $F(x)$ be a primitive function of $f(x)$ on I and $G(y)$ be a primitive function of $1/g(y)$ on J . Then a function y defined on some subinterval of I , solves the differential equation (21) if and only if it satisfies the identity*

$$G(y(x)) = F(x) + C,$$

for all x in the domain of y , and C a real constant.

Example 116 Consider the ODEs from the previous Example 115.

1. The ODE

$$y' = y$$

is separable; in this case we have $f(x) = 1$ and $g(y) = y$. Indeed we can literally separate all terms in x and y on the two sides of the equality as

$$\frac{dy}{y} = dx.$$

Now, we integrate both sides, i.e.

$$\int \frac{dy}{y} = \int dx.$$

From which we obtain

$$\ln y = x + C_1, \text{ so that } G(y) = \ln y, F(x) = x$$

i.e.

$$y = Ce^x, \quad C = e^{C_1}.$$

According to the value of the constant of integration C , we obtain a different particular solution. Hence, we have obtained a family of solutions parametrized by C . This family of solutions is represented in Figure 39 - top left hand side panel.

2. The ODE

$$y = -xy'$$

is separable as well, as it can be written as

$$\frac{dy}{y} = -\frac{dx}{x}.$$

Integrating both sides returns

$$\ln y = -\ln x + C_1,$$

i.e.

$$y = \frac{C}{x}, \quad C = e^{C_1}.$$

The family of solutions is represented in Figure 39 - top right hand side panel.

Example 117 Another example of separable ODE is given by

$$y' = \sqrt{|y|},$$

which can be solved following the same reasoning to obtain

$$2\sqrt{y} = x + C \text{ for } y > 0 \quad 2\sqrt{-y} = -(x + C) \text{ for } y < 0.$$

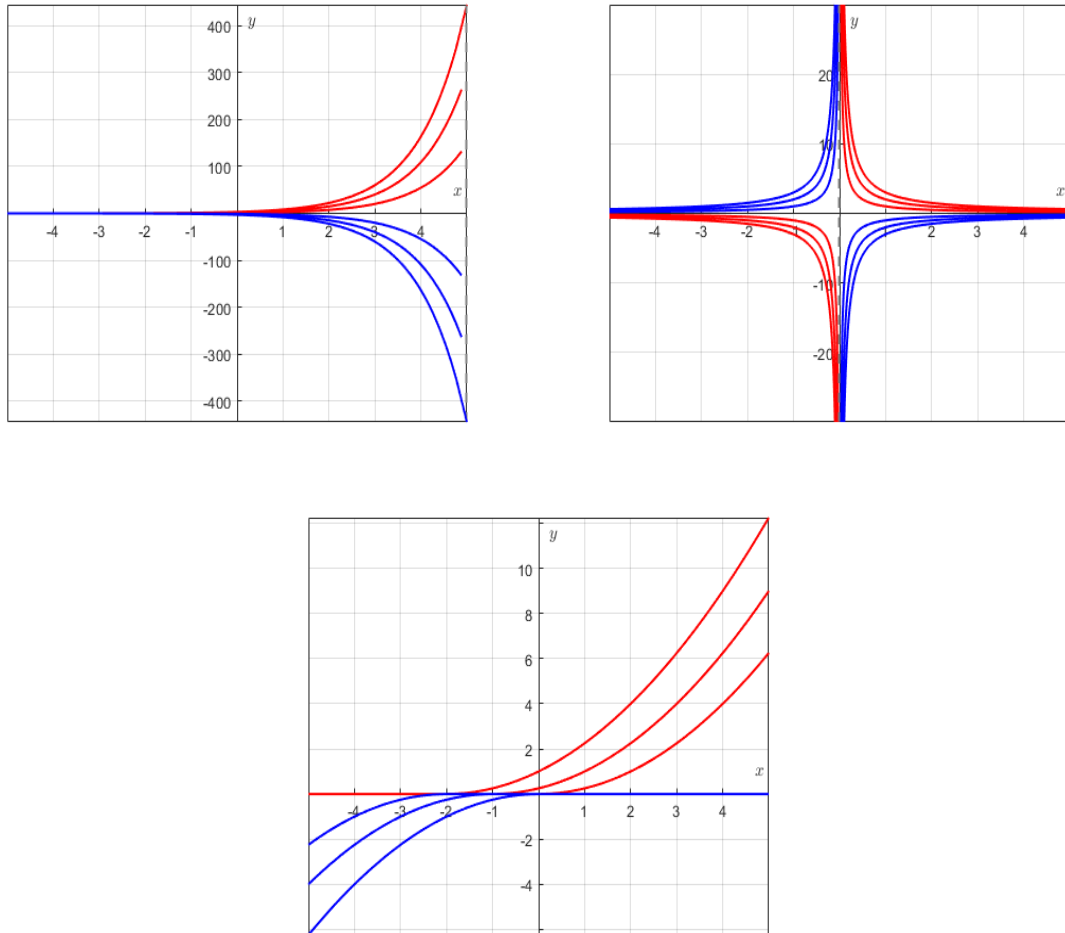


Figure 39: Integral curves of the family of solutions to the ODEs in Example 115–116 and Example 117.

The family of solutions is therefore given by

$$y = \begin{cases} \frac{1}{4}(x+C)^2 & x > -C \\ 0 & x = -C \\ -\frac{1}{4}(x+C)^2 & x < -C \end{cases}$$

The corresponding integral curves are given in Figure 39 - bottom panel. This ODE albeit separable is obviously not linear.

5.1.2 Solution of inhomogenous linear ODEs: integrating factor

The solution to the more general case of inhomogenous ODEs requires more structure. A useful technique is based on the so-called integrating factor.

An integrating factor is a function $q(x)$ by which an ordinary differential equation can be multiplied in order to make it integrable. Finding such a function is not always trivial and requires plenty of practice. A good starting point in the case of linear ODEs is the solution to the corresponding homogenous ODE.

Let us consider again the inhomogenous linear ODE (19)

$$y' + a(x)y = b(x);$$

its corresponding homogenous ODE is

$$y' + a(x)y = 0.$$

By separating variables the above can be rewritten as

$$\frac{dy}{y} = -a(x)dx,$$

consequently

$$\ln y = -A(x) + C_1, \quad A(x) = \int a(x)dx.$$

The solution to the homogenous equation is therefore

$$\hat{y} = Ce^{-A(x)}$$

(observe that this solution is denoted $\hat{y}$ to remind ourselves that this is NOT the solution to the original inhomogenous ODE).

In such a case, the integrating factor can be obtained as

$$q(x) = \frac{\hat{y}(0)}{\hat{y}(x)}$$

i.e.

$$q(x) = e^{A(x)}.$$

We use this integrating factor to build an auxiliary function $F(x)$ defined as

$$F(x) = y(x)q(x).$$

If we have chosen $q(x)$ correctly, we should be able to obtain $y(x)$ from $F(x)$. Of course at this stage we do not know the form of $F(x)$, but we can obtain its derivative, which might be easier to integrate than y' . Let us see.

$$\frac{dF}{dx} = a(x)F(x) + e^{A(x)}y';$$

now substitute the expression of y' to obtain

$$\frac{dF}{dx} = a(x)F(x) + e^{A(x)}(b(x) - a(x)y),$$

which implies

$$dF(x) = e^{A(x)}b(x)dx.$$

Now we can integrate both sides to obtain

$$e^{A(x)}y(x) = \int e^{A(x)}b(x)dx,$$

i.e.

$$y(x) = e^{-A(x)} \int e^{A(x)}b(x)dx.$$

We observe that if we multiply the given ODE

$$y' + a(x)y = b(x)$$

by the integrating factor $e^{A(x)}$, we obtain

$$e^{A(x)}y' + e^{A(x)}a(x)y = e^{A(x)}b(x).$$

Keeping in mind the chain rule for differentiation, we recognize on the left hand side the derivative of $e^{A(x)}y(x)$. Therefore we integrate. This is indeed the scope of the integrating factor.

As a final note: we need to get used to the idea of solutions which might not be in closed form, as the above integral might not be solved analytically. This is not necessarily an issue, as we can use suitable numerical schemes to obtain values out of this expression anyhow.

Example 118 Consider the ODE

$$y' + \frac{y}{x} = e^{x^2}$$

in the domain $x > 0$. Let us build the integrating factor by first solving the homogeneous equation

$$y' + \frac{y}{x} = 0$$

which is separable as it can be written as

$$\frac{y'}{y} = -\frac{1}{x}$$

i.e.

$$\frac{dy}{y} = -\frac{dx}{x}$$

with solution $\hat{y}(x) = C/x$. The integrating factor is therefore $q(x) = x$, and the auxiliary function is $F(x) = xy(x)$ with differential

$$dF(x) = y(x)dx + x \left(e^{x^2} - \frac{y(x)}{x} \right) dx.$$

By integration, we obtain the solution

$$y(x) = \frac{1}{2x}(e^{x^2} + C).$$

5.2 First order quasi-linear ODEs

Quasi-linear ODEs have the form

$$a(x, y) + b(x, y)y' = 0$$

and they are called quasi-linear because they are linear with respect to y' but not necessarily linear with respect to y , as the functions $a(x, y)$ and $b(x, y)$ are not necessarily linear. All we are asking is that these functions are continuously differentiable. This also implies that these ODEs are not necessarily separable, although homogenous.

Example 119 1. The ODE

$$2xy + (x^2 + y^2)y' = 0$$

is quasi-linear and non separable.

2. The ODE

$$y' = -\frac{y}{x + 4x^2y}$$

is quasi-linear and non separable.

These ODEs are more complex to solve due to the lack of separability of the variables, which makes the identification of the integrating factor more tricky. However, we could try to rewrite the ODE as if it were a PDE and try then to identify the total differential of some solution function $F(x, y)$. To be more clear, let us proceed in steps as follows.

Start from the general ODE

$$a(x, y) + b(x, y)y' = 0$$

and rewrite it as

$$a(x, y) + b(x, y)\frac{dy}{dx} = 0,$$

which leads to

$$a(x, y)dx + b(x, y)dy = 0.$$

If the solution function $F(x, y)$ is such that

$$F_x(x, y) = \frac{\partial F}{\partial x} = a(x, y), \quad F_y(x, y) = \frac{\partial F}{\partial y} = b(x, y),$$

then we can recover F from the (partial) integration of the functions $a(x, y)$ and $b(x, y)$. Note though that this is possible only if

$$F_{xy} = \frac{\partial a}{\partial y} = \frac{\partial b}{\partial x} = F_{yx},$$

as we need to ensure that the Schwarz's theorem is satisfied. In this case we say that the differential form

$$dF(x, y) = a(x, y)dx + b(x, y)dy$$

is exact, and inexact otherwise. This is summarised in the following result.

Theorem 35 (The Poincaré lemma) *Let a, b be continuously differentiable functions such that $a_y = b_x$. Then the differential form $a(x, y)dx + b(x, y)dy$ is exact.*

Thus, exact forms can be solved by integration. What about the inexact ones? We need to find a suitable integrating factor which renders the quasi-linear ODE exact, i.e. integrable. Unfortunately, this might not be always immediate. Let us reconsider Example 119.

Example 120 1. We have seen before that the ODE

$$2xy + (x^2 + y^2)y' = 0$$

is quasi-linear and non separable. Let us rewrite the ODE as

$$2xydx + (x^2 + y^2)dy = 0.$$

Then, we would have

$$a(x, y) = 2xy, \quad b(x, y) = x^2 + y^2.$$

We note that

$$a_y = 2x = b_x,$$

thus, this is an exact form which we can solve by integration. We can proceed as follows. As $F_x = a$, then

$$F(x, y) = \int a(x, y)dx + g(y).$$

In our case

$$F(x, y) = 2y \int xdx + g(y) = x^2y + g(y).$$

Now, we find $g(y)$ by matching the partial derivative F_y with $b(x, y)$, i.e.

$$F_y = x^2 + g'(y) = x^2 + y^2 = b(x, y).$$

This implies that

$$g(y) = \frac{y^3}{3}.$$

Finally, the solution we are looking for is

$$x^2y + \frac{y^3}{3} = C.$$

2. The ODE

$$y' = -\frac{y}{x + 4x^2y}$$

is quasi-linear and non separable. However, this is not exact. Let us verify this statement. First we write the ODE in ‘differential’ format

$$-ydx + (x + 4x^2y)dy = 0,$$

so that $a(x, y) = y$ and $b(x, y) = x + 4x^2y$. Therefore, the Poincaré conditions are clearly violated as

$$a_y = -1 \neq 1 + 8xy = b_x.$$

However, we could manipulate the ODE as follows

$$-ydx + x^2 \left(\frac{1}{x} + 4y \right) dy = 0,$$

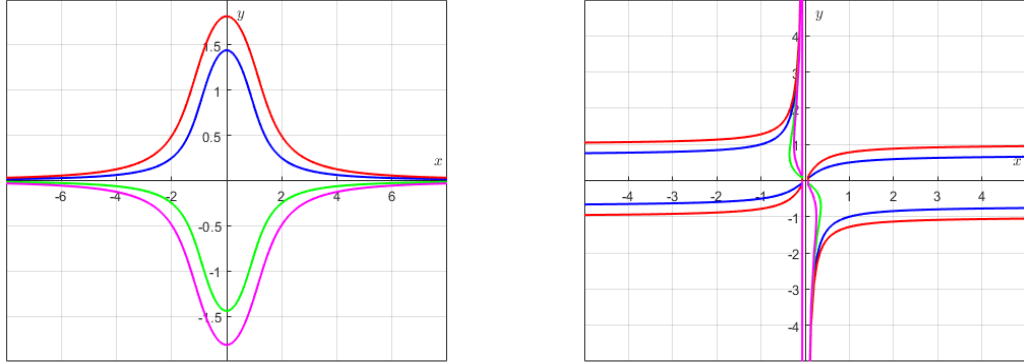


Figure 40: Integral curves of the family of solutions to the ODEs in Example 120.

and observe that

$$\frac{\partial}{\partial x} \frac{1}{x} + 4y = -\frac{1}{x^2}.$$

This seems to suggest that a good integrating factor could be $1/x^2$, as in this case the ODE would then read

$$-\frac{y}{x^2} dx + \left(\frac{1}{x} + 4y \right) dy = 0,$$

and

$$\frac{\partial}{\partial y} \left(-\frac{y}{x^2} \right) = -\frac{1}{x^2}.$$

In other words, we have now an exact form for $a(x, y) = -y/x^2$ and $b(x, y) = 1/x + 4y$. We can now integrate using the same procedure as in the previous example to obtain

$$\frac{y}{x} + 2y^2 = C.$$

The general solutions for both ODEs are given in Figure 40.

Part II**Probability and Statistics**

6 The basic framework: the probability space. Describing a random experiment

Imagine a random experiment like the toss of a coin or the prices of securities traded in the market in the next period of time. Imagine that we want to explore the features of this random experiment in order to make appropriate and informed decisions. These features could be: the expected price of the security tomorrow, or its volatility; the characteristics of the tails of the price distribution (if for example you need to calculate some risk measure such as VaR, or shortfall expectation).

In order to be able to do all this, we need appropriate tools describing the random experiment in such a way that we can extract all this information, i.e. we need a mathematical model of the random experiment. This is represented by the so-called *probability space*.

Definition 36 (Probability space) *We denote the probability space by the triplet*

$$\Theta := (\Omega, \mathbb{F}, \mathbb{P}).$$

A probability space can be considered as a mathematical model of a random experiment.

This definition is telling us that the probability space is made up of three building blocks, which we are going to explore one by one.

The first piece of the probability space is Ω , which represents our *sample space*, i.e. the set of all possible outcomes of random experiment.

Example 121 Let the random experiment be defined as: choose a number from the unit interval $[0, 1]$. Then $\Omega = \{\omega : 0 \leq \omega \leq 1\} = [0, 1]$.

Example 122 Assume now that the random experiment you are interested into is the evolution of a stock price over an infinite time horizon, when only 2 states of nature can occur, i.e. up or down. Then Ω = the set of all infinite sequences of ups and downs $= \{\omega : \omega_1\omega_2\omega_3\dots\}$, where ω_n is the result at the n -th period.

The second object you need in order to have a probability space is $\mathbb{F}$ which is called σ -algebra. The σ -algebra of a random experiment can be interpreted as the collection of all possible histories of the random experiment itself. Formally, it is defined as follows.

Definition 37 (σ -algebra) *Given a set Ω , a collection $\mathbb{F}$ of subsets of Ω is a σ -algebra if*

1. $\emptyset \in \mathbb{F}$
2. $A \in \mathbb{F}$ implies $A^c \in \mathbb{F}$
3. $\{A_m\} \in \mathbb{F}$ implies $\bigcup_{m=1}^{\infty} A_m \in \mathbb{F}$ (infinite union).

Example 123 1. $\mathbb{F} = \{\emptyset, \Omega\}$ is a σ -algebra

2. Consider some event $A \subset \Omega$. Then the σ -algebra generated by A is $\mathbb{F} = \{\emptyset, \Omega, A, A^c\}$.

3. Consider the sample space defined above for the evolution of the stock price in a 2-state economy, i.e. Ω = the set of infinite sequences of ups and downs, and define

$$\begin{aligned} A_U &= \{\omega : \omega_1 = U\} \\ A_D &= \{\omega : \omega_1 = D\}. \end{aligned}$$

The σ -algebra generated by these two sets is

$$\mathbb{F}^{(1)} = \{\emptyset, \Omega, A_U, A_D\}.$$

Now consider the sets

$$\begin{aligned} A_{UU} &= \{\omega : \omega_1 = U, \omega_2 = U\} \\ A_{UD} &= \{\omega : \omega_1 = U, \omega_2 = D\} \\ A_{DU} &= \{\omega : \omega_1 = D, \omega_2 = U\} \\ A_{DD} &= \{\omega : \omega_1 = D, \omega_2 = D\}. \end{aligned}$$

Then

$$\begin{aligned} \mathbb{F}^{(2)} = \{ & \emptyset, \Omega, A_{UU}, A_{UD}, A_{DU}, A_{DD}, A_{UU}^c, A_{UD}^c, A_{DU}^c, A_{DD}^c, A_U, A_D, \\ & A_{UU} \cup A_{DU}, A_{UU} \cup A_{DD}, A_{DU} \cup A_{UD}, A_{UD} \cup A_{DD} \} \end{aligned}$$

is the corresponding σ -algebra.

Every σ -algebra has a set of properties that will be useful in the future.

Theorem 38 *The σ -algebra has the following properties*

1. $\Omega \in \mathbb{F}$.
2. $\{A_m\} \in \mathbb{F}$ implies $\bigcap_{m=1}^{\infty} A_m \in \mathbb{F}$.

An important σ -algebra is the Borel σ -algebra (on $\mathbb{R}$), which is the smallest σ -algebra containing all open sets (or, equivalently, all closed sets) on the real line.

The last element of our probability space is represented by the symbol $\mathbb{P}$. This is called *probability measure*, and you can consider it as a sort of ‘metrics’, that measures the likelihood of a specific event or story of the random experiment.

Definition 39 *A probability measure $\mathbb{P}$ is a set function $\mathbb{P} : \mathbb{F} \rightarrow [0, 1]$ such that*

1. $\mathbb{P}(\Omega) = 1$
2. For any sequence of disjoint events $\{A_m\}$, $\mathbb{P}(\bigcup_{m=1}^{\infty} A_m) = \sum_{m=1}^{\infty} \mathbb{P}(A_m)$.

Based on this definition, you can show that

$$\begin{aligned}\mathbb{P}(\emptyset) &= 0; \\ \mathbb{P}(A \cup B) &= \mathbb{P}(A) + \mathbb{P}(B); \\ \mathbb{P}(A^c) &= 1 - \mathbb{P}(A).\end{aligned}$$

Moreover, we can define independent events: two events, A and B , are independent if and only if $\mathbb{P}(A \cap B) = \mathbb{P}(A)\mathbb{P}(B)$.

Example 124 Consider the previous example of the evolution of the stock price over an infinite time horizon, so that $\Omega = \{\omega : \omega_1\omega_2\omega_3\ldots\}$, and $A_U = \{\omega : \omega_1 = U\}$, $A_D = \{\omega : \omega_1 = D\}$. Assume that the different up/down movements at each time step are independent, and let

$$\mathbb{P}(A_U) = p; \mathbb{P}(A_D) = q = 1 - p.$$

Then

$$\mathbb{P}(A_{UU}) = p^2; \mathbb{P}(A_{UD}) = \mathbb{P}(A_{DU}) = pq; \mathbb{P}(A_{DD}) = q^2.$$

Further, $\mathbb{P}(A_{UU}^c) = 1 - p^2$; similarly, you can calculate the probability of each other set in $\mathbb{F}^{(2)}$. Moreover, if $A_{UUU} = \{\omega : \omega_1 = U, \omega_2 = U, \omega_3 = U\}$, you can calculate that $\mathbb{P}(A_{UUU}) = p^3$. And so on. Hence, in the limit you can conclude that the probability of the sequence $UUU\ldots$ is zero. The same applies for example to the sequence $UDUD\ldots$; in fact this sequence is the intersection of the sequences $U, UD, UDU, \ldots$. From this example, we can conclude that every single infinite sequence of movements in Ω has probability zero.

In the previous example, we have shown that

$$\mathbb{P}(\text{every movement is up}) = 0;$$

this implies that this event is sure not to happen. Similarly, since the above is true, we are sure to get at least one down movement in the sequence, although we do not know exactly when in the sequence. Because of this fact, and the fact that the infinite sequence $UUU\ldots$ is in the sample space (which means that still is a possible outcome), mathematicians have come up with a somehow strange way of saying: we will get at least one down movement *almost surely*.

Definition 40 Let $(\Omega, \mathbb{F}, \mathbb{P})$ be a probability space. If $A \subset \mathbb{F}$ is such that

$$\mathbb{P}(A) = 1,$$

we say that the event A occurs *almost surely* (a.s.).

Now, in order to introduce the next definition, consider the following simple example. Assume that you want to measure the length of a room, and assume you express this measure in meters and centimeters. It turns out that the room is 4.30m. long. Now assume that you want to change the reference system and express the length of the room in terms of feet and inches. Then, the room is 14ft. long. But in the process of switching from one reference system to the other, the room did not change: it did not shrink; it did not expand. The same applies to events and probability measures. The idea is given in the following.

Definition 41 (Absolutely continuous/equivalent probability measure) *Given two probability measures $\mathbb{P}$ and $\mathbb{P}^*$ defined on the same σ -algebra $\mathbb{F}$, then.*

i) $\mathbb{P}$ is absolutely continuous with respect to $\mathbb{P}^*$, i.e. $\mathbb{P} \ll \mathbb{P}^*$, if $\mathbb{P}(A) = 0$ whenever, $\mathbb{P}^*(A) = 0 \forall A \in \mathbb{F}$.

ii) If $\mathbb{P} \ll \mathbb{P}^*$ and also $\mathbb{P}^* \ll \mathbb{P}$, then $\mathbb{P} \sim \mathbb{P}^*$, i.e. $\mathbb{P}$ and $\mathbb{P}^*$ are equivalent measures. Thus, for $\mathbb{P} \sim \mathbb{P}^*$ the following statements are equivalent.

- $\mathbb{P}(A) = 0 \Leftrightarrow \mathbb{P}^*(A) = 0$ (same null sets)
- $\mathbb{P}(A) = 1 \Leftrightarrow \mathbb{P}^*(A) = 1$ (same a.s. sets)
- $\mathbb{P}(A) > 0 \Leftrightarrow \mathbb{P}^*(A) > 0$ (same sets of positive measures)

Example 125 As random experiment of reference, consider rolling a die. Suppose we have three different ‘metrics’ to measure the probability of the outcome of this experiment as described in the following table.

Table 6: Die outcomes and probabilities in Example 125.

z		1	2	3	4	5	6
$\mathbb{P}$	$p_Z(z)$	1/6	1/6	1/6	1/6	1/6	1/6
$\tilde{\mathbb{P}}$	$\tilde{p}_Z(z)$	1/21	2/21	3/21	4/21	5/21	6/21
$\mathbb{P}^*$	$p_Z^*(z)$	0	1/3	0	1/3	0	1/3

In the above, the first row shows a metrics under which the die is fair; the second row returns a metrics which calculates probabilities as proportional to the outcomes; the third row instead uses a metrics which assigns equal probabilities only to even numbers. Using the given definition we notice the following.

$$\mathbb{P}(A) = 0 \Rightarrow \mathbb{P}^*(A) = 0 \text{ but } \mathbb{P}^*(A) = 0 \not\Rightarrow \mathbb{P}(A) = 0,$$

and therefore $\mathbb{P} \not\sim \mathbb{P}^*$, i.e. the two probabilities are not equivalent. The probabilities $\mathbb{P}$ and $\tilde{\mathbb{P}}$ are instead equivalent.

As there is a conversion factor that helps you to switch between meters and feet, so that $4.30\text{m} = 14\text{ft}$, there is also a conversion factor between probability measures. The idea can be shown in relatively simple terms once we know better random variables and their related distribution. So, for now we take a leave from changes of measures and we will discuss it again later in this section.

7 Random variables

So far, we have considered random events, like an up or down movement in the stock price over the next period of time, and the likelihood of such events to occur, as described by the probability measure. The next step in which you might be interested is to ‘quantify’ the outcome of the random event, for example you might want to know how much the stock price is going to change if an up or down movement is going to occur in the next time period. In order to do this, you need the idea of *random variable*.

Definition 42 (Random variable) Let $(\Omega, \mathbb{F}, \mathbb{P})$ be a probability space. A random variable X is a function $X : \Omega \rightarrow \mathbb{R}$ such that $\{\omega \in \Omega \mid X(\omega) \leq x\} \in \mathbb{F} \quad \forall x \in \mathbb{R}$.

Note that if B is any subset of the Borel σ -algebra $\mathbb{B}$, i.e. B is a set of the form $B = (-\infty, x] \quad \forall x \in \mathbb{R}$, then Definition 42 implies that $X^{-1}(B) \in \mathbb{F} \quad \forall x \in \mathbb{R}$. In other words, any random variable is a numerical quantity whose value is determined by the random experiment of choosing some $\omega \in \Omega$.

Example 126 Consider once again the random experiment of the evolution of the stock price over an infinite time horizon in a 2-state economy, described in Example 123.3. Let us define the stock prices by the formula

$$\begin{aligned} S_0(\omega) &= 4; \\ S_1(\omega) &= \begin{cases} 8 & \text{if } \omega_1 = up \\ 2 & \text{if } \omega_1 = down \end{cases} \\ S_2(\omega) &= \begin{cases} 16 & \text{if } \omega_1 = \omega_2 = up \\ 4 & \text{if } \omega_1 \neq \omega_2 \\ 1 & \text{if } \omega_1 = \omega_2 = down. \end{cases} \end{aligned}$$

$S_0(\omega), S_1(\omega), S_2(\omega)$ are random variables, assigning a numerical value to each sequence of up and down movements in the stock price at each time period. Example 124 tells us how to calculate the probability that the random variable S takes any of these values; for example

$$\begin{aligned} \mathbb{P}(S_1(\omega) = 8) &= \mathbb{P}(A_U) = p; \\ \mathbb{P}(S_2(\omega) = 4) &= \mathbb{P}(A_{DU} \cup A_{UD}) = 2pq. \end{aligned}$$

The above Example shows that we can associate to any random variable another function measuring the likelihood of the outcomes. This is what we call the law of X . Precisely, by law of X we mean a probability measure on $(\mathbb{R}, \mathbb{B})$, $L_X : \mathbb{B} \rightarrow [0, 1]$ such that

$$L_X(B) = \mathbb{P}(X \in B) \quad \forall B \in \mathbb{B}.$$

In general, we prefer to speak in terms of *distribution of a random variable*; this is a function $F_X : \mathbb{R} \rightarrow [0, 1]$ defined as

$$F_X(a) = \mathbb{P}(X \leq a) = \mathbb{P}(\omega : X(\omega) \leq a).$$

This is the law of X for any set B of the form $B = (-\infty, a]$, i.e. $F_X(a) = L_X(-\infty, a]$. In some special cases, we can describe the distribution function of a random variable X in even more details. The first case is the case of a *discrete random variable*, like the one introduced in Example 126, which assigns lumps of mass to events. For this random variable, we can express the distribution function as

$$F_X(a) = \mathbb{P}(X \leq a) = \sum_{X \leq a} p_X(x),$$

where $p_X(x)$ is the probability mass function of X .

Example 127 The binomial random variable $X \sim B(n, p)$ is a discrete random variable defined by the probability mass function

$$p_X(k) = \binom{n}{k} p^k (1-p)^{n-k}.$$

If instead the random variable X spreads the mass continuously over the real line, then we have a *continuous random variable* and

$$F_X(a) = \mathbb{P}(X \leq a) = \int_{-\infty}^a f_X(x) dx, \quad (22)$$

where $f_X(x)$ denotes the *density function* of X .

Example 128 The normal (or Gaussian) random variable $X \sim N(\mu, \sigma^2)$ is a continuous random variable defined by the density function

$$f_X(x) = \frac{e^{-\frac{(x-\mu)^2}{2\sigma^2}}}{\sigma\sqrt{2\pi}}.$$

Note. It is often convenient to work in terms of the standard normal random variable $Z \sim N(0, 1)$; this is possible via the standardization of the random variable X by setting

$$Z = \frac{X - \mu}{\sigma};$$

from the properties of the normal distribution it follows that Z indeed has zero mean and unit variance. Alternatively, you can write $X = \mu + \sigma Z$. It follows from standardization that

$$\begin{aligned} \mathbb{P}(X \leq x) &= \mathbb{P}(\mu + \sigma Z \leq x) \\ &= \mathbb{P}\left(Z \leq \frac{x - \mu}{\sigma}\right) \end{aligned}$$

or, equivalently, by substitution

$$\begin{aligned} \mathbb{P}(X \leq x) &= \int_{-\infty}^x \frac{e^{-\frac{(y-\mu)^2}{2\sigma^2}}}{\sigma\sqrt{2\pi}} dy \\ &= \int_{-\infty}^{\frac{x-\mu}{\sigma}} \frac{e^{-\frac{z^2}{2}}}{\sqrt{2\pi}} dz. \end{aligned}$$

The Gaussian distribution is one of the fundamental building block in Quantitative Finance, especially as it serves as the basis for constructing other distributions which better capture the stylised empirical features of financial markets. Examples of distributions which can be recovered via transformations of the Gaussian distribution are the log-normal distribution and the chi-square distribution (central and non central). Specifically, X is log-normal distributed with parameters μ, σ^2 if $X = e^Y$, for $Y \sim N(\mu, \sigma^2)$. The chi-square distribution instead originates from the sum of n Gaussian distributed random variables. If these have all zero mean and unit variance, we obtain the central chi-square random variable with n degrees of freedom. If instead the Gaussian random variables which are summed have non-zero mean and non-unit variance, we obtain the non-central chi-square.

Another useful way of constructing distributions via the Gaussian one is to use mean-variance mixtures. Formally, a normal variance-mean mixture with mixing probability density g is the continuous probability distribution of a random variable Y of the form

$$Y = \alpha + \beta V + \sigma \sqrt{V} X$$

for $\alpha, \beta \in \mathbb{R}$, $\sigma > 0$, and random variables X and V are independent, $X \sim N(0, 1)$, and V is continuously distributed on the positive half-axis with probability density function g . The conditional distribution of Y given V is thus a normal distribution with mean $\alpha + \beta V$ and variance $\sigma^2 V$. Classical examples are the t-distribution (for $\beta = 0$), the Variance Gamma distribution and the Normal Inverse Gaussian distribution. In these cases, the mixing random variable V is given by the inverse gamma, the gamma, and the inverse Gaussian respectively.

We note the following: it is more useful to remember how certain random variables are constructed rather than knowing directly their density function. By knowing the construction, we can deduce many properties directly from the Gaussian distribution.

Remark 16 (A matter of notation) From equation (22), we see that we could write the density function as

$$f_X(x) = \frac{dF_X}{dx} = \frac{d\mathbb{P}(\omega)}{dx} \quad \forall x \in \mathbb{R}.$$

The *expectation* $\mathbb{E}$ of a random variable X on $(\Omega, \mathbb{F}, \mathbb{P})$ is then defined as

$$\mathbb{E}(X) = \int_{\Omega} X(\omega) d\mathbb{P}(\omega) = \int_{-\infty}^{\infty} x dF_X(x).$$

The expectation returns the mean of the distribution; you might be also interested in the dispersion around the mean, this feature is described by the variance of a random variable. Further features that characterize the distribution of a random variable are the skewness (degree of asymmetry) and the kurtosis (behaviour of the tails). These features are described by the moments (from the mean) of a random variable which can be recovered via the moment generating function of the random variable itself. This is explored in details in the next section.

7.1 Moments, Cumulants and the Generating Functions - a reminder

The Moment Generating Function (MGF) is defined as follows.

$$M_X(k) = \mathbb{E}(e^{kX}) = \int_{-\infty}^{\infty} e^{kx} dF_X(x).$$

As the name of this function suggests, it can be used to study the distribution features of random variables as it helps to quickly recover the relevant moments. In more details, we distinguish between.

- The raw moments of the random variable X , which are defined as

$$\mu_n = \mathbb{E}(X^n),$$

and which can be recovered by direct differentiation of the MGF as

$$\mu_n = \left. \frac{\partial^n M_X}{\partial k^n}(k) \right|_{k=0}.$$

Proof. Let us use the Taylor expansion of the exponential function to rewrite the MGF as

$$\begin{aligned} M_X(k) &= \mathbb{E} \left(\sum_{n=0}^{\infty} \frac{(kX)^n}{n!} \right) \\ &= \sum_{n=0}^{\infty} \frac{k^n}{n!} \mathbb{E}(X^n) \\ &= 1 + k\mathbb{E}(X) + \frac{1}{2}k^2\mathbb{E}(X^2) + \dots \end{aligned}$$

First order differentiation of the above with respect to k returns

$$\frac{\partial M_X}{\partial k}(k) = \mathbb{E}(X) + k\mathbb{E}(X^2) + \dots;$$

second order differentiation of the above with respect to k returns

$$\frac{\partial^2 M_X}{\partial k^2}(k) = \mathbb{E}(X^2) + k\mathbb{E}(X^3) + \dots,$$

and so on. The required results follows setting $k = 0$ in the above. ■

- The central moments of the random variable X , defined as

$$m_n = \mathbb{E}[(X - \mathbb{E}(X))^n] \quad n \geq 2.$$

The central moments can be easily obtained from the cumulants which are defined as

$$c_n = \left. \frac{\partial^n \ln M_X}{\partial k^n}(k) \right|_{k=0}.$$

Proof. First order differentiation of $\ln M_X(k)$ returns

$$\frac{\partial \ln M_X}{\partial k}(k) = \frac{1}{M_X(k)} \frac{\partial M_X}{\partial k}(k);$$

setting $k = 0$ returns the first order cumulant which coincides with the mean. Second order differentiation returns

$$\frac{\partial^2 \ln M_X}{\partial k^2}(k) = -M_X(k)^{-2} \left(\frac{\partial M_X}{\partial k}(k) \right)^2 + \frac{1}{M_X(k)} \frac{\partial^2 M_X}{\partial k^2}(k);$$

setting $k = 0$ returns $-\mathbb{E}(X)^2 + \mathbb{E}(X^2)$, i.e. the variance. The higher order cumulants can be obtained following the same procedure. ■

In particular, the following identities hold.

$$\begin{aligned} c_1 &= \mathbb{E}(X) \\ c_2 &= \mathbb{E}[(X - \mathbb{E}(X))^2] = \text{Var}(X) = m_2 \\ c_3 &= \mathbb{E}[(X - \mathbb{E}(X))^3] = m_3 \\ c_4 &= \mathbb{E}[(X - \mathbb{E}(X))^4] - 3\text{Var}^2(X) = m_4 - 3m_2^2. \end{aligned}$$

The skewness and the kurtosis of the distribution of a random variable are measured respectively by the index of skewness

$$\gamma_1 = \frac{c_3}{c_2^{3/2}}$$

and by the index of excess kurtosis

$$\gamma_2 = \frac{c_4}{c_2^2}.$$

The distribution of a random variable is said to be asymmetric to the left if $\gamma_1 < 0$, symmetric if $\gamma_1 = 0$, and asymmetric to the right otherwise. It is platikurtic if $\gamma_2 < 0$, mesokurtic if $\gamma_2 = 0$ and leptokurtic otherwise.

As the following example shows, the Gaussian distribution is both symmetric and mesokurtic. If any other random variable is, for example, leptokurtic, it means that the tails of its distribution are thicker than the ones of the Gaussian random variable. Therefore, ‘extreme’ events are more likely than in the Gaussian case.

Example 129 Few (and very important, as we will use them throughout the entire year) examples of random variables.

1. As already seen before, the normal (or Gaussian) random variable $X \sim N(\mu, \sigma^2)$ is a continuous random variable defined by the density function

$$f_X(x) = \frac{e^{-\frac{(x-\mu)^2}{2\sigma^2}}}{\sigma\sqrt{2\pi}}.$$

You can easily show that

$$\mathbb{E}(X) = \mu; \text{Var}(X) = \sigma^2; c_3 = 0 = c_4; \quad M_X(k) = e^{k\mu + \frac{k^2\sigma^2}{2}}.$$

2. The Poisson random variable is an example of discrete random variable. More precisely, a Poisson random variable $N \sim \text{Poi}(\lambda)$, with rate λ has probability mass

$$p_N(n) = \frac{e^{-\lambda}\lambda^n}{n!}$$

from which it follows that

$$\mathbb{E}(N) = \lambda = \text{Var}(N) = c_3 = c_4; \quad M_N(k) = e^{\lambda(e^k - 1)}.$$

Hence the Poisson distribution is asymmetric to the right and leptokurtic.

3. Assume $X \sim \Gamma(\alpha, \lambda)$, $\alpha > 0$. Then X is a non-negative random variable which follows a gamma distribution; its density function is given by

$$f(x) = \frac{1}{\Gamma(\alpha)} \lambda^\alpha x^{\alpha-1} e^{-\lambda x},$$

where $\Gamma(\alpha)$ is the gamma function, which is defined as

$$\Gamma(\alpha) = \int_0^\infty x^{\alpha-1} e^{-x} dx,$$

and has the property that

$$\Gamma(\alpha) = \Gamma(\alpha - 1)(\alpha - 1).$$

This means that

$$\Gamma(\alpha) = (\alpha - 1)!$$

where α is a positive integer. The MGF of X is

$$M_X(k) = \frac{1}{\Gamma(\alpha)} \lambda^\alpha \int_0^\infty x^{\alpha-1} e^{-x(\lambda-k)} dx.$$

Set $y = x(\lambda - k)$, then

$$M_X(k) = \frac{1}{\Gamma(\alpha)} \lambda^\alpha \int_0^\infty \left(\frac{y}{\lambda - k} \right)^{\alpha-1} \frac{e^{-y}}{\lambda - k} dy = \left(\frac{\lambda}{\lambda - k} \right)^\alpha.$$

Note that if $\alpha = 1$, then X follows an exponential distribution with rate λ . Using the MGF you can show that the gamma random variable has mean $\mu = \alpha/\lambda$ and variance $\nu = \alpha/\lambda^2$. The parameter α is the shape parameter, whilst λ is the scale parameter. You can also show that the gamma random variable is asymmetric to the right and leptokurtic.

Exercise 1 Verify the results presented in the previous example.

Exercise 2 The binomial random variable $X \sim B(n, p)$ is a discrete random variable defined by the probability mass function

$$p_X(k) = \binom{n}{k} p^k (1-p)^{n-k}.$$

1. Show that the MGF of the binomial random variable is

$$(pe^h + 1 - p)^n.$$

2. Using the MGF given above, deduce that X has mean equal to np and variance equal to $np(1-p)$.
3. Study the behaviour of the indices of skewness and excess kurtosis when the number of trials n increases to infinity.

Example 130 Let $\{Z_i\}_{i=1,\dots,n} \sim N(0, 1)$ be a sequence of i.i.d. random variables. The (central) chi-square random variable with n degrees of freedom is defined as

$$Y = \sum_{i=1}^n Z_i^2.$$

We note that

$$\begin{aligned} \mathbb{E}\left(e^{kZ^2}\right) &= \int_{\mathbb{R}} e^{kz^2} \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}} dz \\ &= \int_{\mathbb{R}} \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}(1-2k)} dz \\ &= (\sqrt{1-2k})^{-1} \int_{\mathbb{R}} \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} dx \\ &= (1-2k)^{-1/2}, \end{aligned}$$

where the penultimate equality follows by substituting $x = z\sqrt{1-2k}$. Therefore, using independence of the random variables Z_i

$$\begin{aligned} \mathbb{E}\left(e^{kY}\right) &= \mathbb{E}\left(e^{k\sum_{i=1}^n Z_i^2}\right) \\ &= \prod_{i=1}^n \mathbb{E}\left(e^{kZ_i^2}\right) \\ &= (1-2k)^{-n/2} \end{aligned}$$

where the last equality follows from the fact that the Z_i are identically distributed. Thus, we have deduced the MGF of the chi-square distribution with n degrees of freedom.

Now assume a sequence of i.i.d. random variables $\{X_i\}_{i=1,\dots,n} \sim N(\lambda_i, 1)$ and build the random variable

$$Y = \sum_{i=1}^n X_i^2.$$

Then Y follows a non-central chi-square distribution with n degrees of freedom and parameter of non centrality $\lambda = \sum_{i=1}^n \lambda_i^2$. Following the same steps as above we can first prove that

$$\mathbb{E}\left(e^{kX_i^2}\right) = (1-2k)^{-1/2} e^{k\lambda_i^2/(1-2k)},$$

and then deduce that the MGF of Y is $(1-2k)^{-n/2} e^{k\lambda/(1-2k)}$. All the moments can then be deduced by direct differentiation.

Exercise 3 Assume a sequence of i.i.d. random variables $\{X_i\}_{i=1,\dots,n} \sim N(\mu_i, \sigma^2)$ and build the random variable

$$Y = \sum_{i=1}^n X_i^2.$$

Show that this random variable is equivalent to a rescaled non-central chi square random variable with n degrees of freedom, parameter of non-centrality $\lambda = \sum_{i=1}^n \lambda_i^2$, for $\lambda_i = \mu_i/\sigma$. The rescaling factor is σ^2 . (Hint: derive the MGF).

Example 131 Let us consider the case of the random variable X which is log-normal distributed with parameters μ, σ^2 . This means that $X = e^Y$, for $Y \sim N(\mu, \sigma^2)$. The mean can be deduced easily using the MGF of Y as

$$\mathbb{E}(Y) = \mathbb{E}(e^Y) = e^{\mu + \frac{\sigma^2}{2}}.$$

The higher order moments of the log-normal random variable can be deduced in the same manner.

Moment generating functions suffer the disadvantage that the integrals which define them may not always be finite.

Example 132 A Cauchy random variable X has density function

$$f(x) = \frac{1}{\pi(1+x^2)} \quad x \in \mathbb{R}.$$

Hence the MGF of X is given by

$$M_X(k) = \int_{-\infty}^{\infty} \frac{e^{kx}}{\pi(1+x^2)} dx.$$

This is an improper integral of the 1st kind which does not converge unless $k = 0$ (which of course is a ‘nonsense’...) In fact, if you perform the convergence test, you obtain that

$$\begin{aligned} \lim_{x \rightarrow \infty} \frac{\frac{e^{kx}}{\pi(1+x^2)}}{\left(\frac{1}{x}\right)^\alpha} &\stackrel{\alpha=2}{=} \frac{1}{\pi} \lim_{x \rightarrow \infty} e^{kx} = \begin{cases} 0 & \text{if } k < 0 \\ \infty & \text{if } k > 0, \end{cases} \\ \lim_{x \rightarrow -\infty} \frac{\frac{e^{kx}}{\pi(1+x^2)}}{\left(\frac{1}{x}\right)^\alpha} &\stackrel{\alpha=2}{=} \frac{1}{\pi} \lim_{x \rightarrow -\infty} e^{kx} = \begin{cases} 0 & \text{if } k > 0 \\ \infty & \text{if } k < 0. \end{cases} \end{aligned}$$

Hence, the MGF of a Cauchy random variable does not exist.

Characteristic functions are another class of functions equally useful and whose finiteness is guaranteed (so you might actually want to say that they are more useful than MGFs...).

Definition 43 The characteristic function of X is the function $\phi_X : \mathbb{R} \rightarrow \mathbb{C}$ defined by

$$\phi_X(u) = \mathbb{E}(e^{iuX})$$

for $i = \sqrt{-1}$.

This is a common transformation and is often called the Fourier transform of the density f of X if this quantity exists. In this case

$$\phi_X(u) = \int e^{iux} dF_X(x) = \int e^{iux} f_X(x) dx.$$

The characteristic function of a random variable has several nice properties. Firstly it always exists and it is finite (in $\mathbb{L}^1$): note that

$$\phi_X(u) = \mathbb{E}(e^{iuX}) = \mathbb{E}(\cos(uX) + i \sin(uX)),$$

hence²

$$|\cos(uX) + i \sin(uX)| := \sqrt{\cos(uX)^2 + \sin(uX)^2} = 1.$$

Then

$$|\mathbb{E}(e^{iuX})| \leq \mathbb{E}(|e^{iuX}|) = 1.$$

Moreover.

- There is bijection between cumulative distribution functions/probability density function and characteristic functions: two distinct probability distributions never share the same characteristic function.
- if X and Y are independent random variables, $\phi_{X+Y}(u) = \phi_X(u) \phi_Y(u)$;
- if $a, b \in \mathbb{R}$ and $Y = aX + b$, then $\phi_Y(u) = e^{iub} \phi_X(au)$.

Examples of characteristic functions * Calculations of integrals involving complex numbers are not always pleasant; usually you should know about contour integration... but for our purposes you can get away with only knowing about analytic continuation.

Analytic continuation provides a way of extending the domain over which a complex function is defined. Let us start from a *complex function* f (like the characteristic function); this function is *complex differentiable* at z_0 and has derivative A if and only if

$$f(z) = f(z_0) + A(z - z_0) + o(z - z_0), \quad \forall z \in \mathbb{C}.$$

A complex function is said to be *analytic* on a region D if it is complex differentiable at every point in D (i.e. has no singularities, i.e. points at which the function ‘blows up’ or becomes degenerate). Now, let f_1 and f_2 be analytic functions on domains D_1 and D_2 respectively, with $D_1 \subset D_2$, such that $f_1 = f_2$ on $D_1 \cap D_2$. Then f_2 is called the *analytic continuation* of f_1 to D_2 . Moreover, if it exists, the analytic continuation of f_1 to D_2 is unique.

Consider now the MGF M_X of some random variable X ; we can say that the function

$$\mathcal{M}_X(z) = \int_{-\infty}^{\infty} f(x) e^{zx} dx \quad z \in \mathbb{C}$$

is the analytic continuation of M_X to the complex plane, if it respects the condition above. Then, the characteristic function of X , ϕ_X , is the restriction of $\mathcal{M}_X$ to the imaginary axis, i.e.

$$\phi_X(u) = \mathcal{M}_X(iu)$$

Example 133 1. Let $X \sim N(0, 1)$. The characteristic function is

$$\phi_X(u) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{iux - \frac{x^2}{2}} dx.$$

Now consider the real valued function

$$M_X(k) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{kx - \frac{x^2}{2}} dx = e^{\frac{k^2}{2}},$$

²Note that this is the complex square of the complex number $z = \cos(uX) + i \sin(uX)$, and you can interpret the notation as a norm.

i.e. the MGF of X . Since $\mathbb{R} \cap \mathbb{C} \neq \emptyset$, then M_X has analytic continuation on the complex plane given by

$$\mathcal{M}_X(z) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{zx - \frac{x^2}{2}} dx = e^{\frac{z^2}{2}} \quad z \in \mathbb{C}.$$

Therefore, by analytic continuation

$$\phi_X(u) = \mathcal{M}_X(iu) = e^{-\frac{u^2}{2}}.$$

2. Let X be a Poisson random variable with rate u . You can apply the same argument as above (i.e. analytic continuation) to show that

$$\phi_X(u) = \mathbb{M}_X(iu) = e^{\lambda(e^{iu} - 1)}.$$

3. Consider now the gamma distribution. Analytic continuation implies that

$$\phi_X(u) = \left(\frac{\lambda}{\lambda - iu} \right)^\alpha.$$

4. Assume X is a Cauchy random variable, i.e.

$$f(x) = \frac{1}{\pi(1+x^2)}.$$

We cannot use the analytic continuation argument because the function is not analytic (can you spot why?). The CF exists anyhow and is equal to

$$\phi_X(u) = e^{-|u|}.$$

However, to prove the result you need to use contour integration and the residue theorem - which are beyond the scope of this module.

Similarly to the MGF, the characteristic function can be used to obtain the raw moments and the central moments of a random variable. In particular, the following hold.

$$\begin{aligned} i^n \mu_n &= \left. \frac{\partial^n \phi_X}{\partial u^n}(u) \right|_{u=0}, \\ i^n c_n &= \left. \frac{\partial^n \ln \phi_X}{\partial u^n}(u) \right|_{u=0} \end{aligned}$$

(the proof follows the same steps as in the case in which the MGF is used).

7.2 Changing measures

In the following we conclude our excursus in changes of measures. As discussed in the previous sections, we have different metrics to quantify, or ‘measure’, the likelihood of random events. How do we switch from one to the other? How do obtain our ‘conversion factor’?

Suppose the random variable X has density function $f_X(\cdot)$ under the measure $\mathbb{P}$ and $\tilde{f}_X(\cdot)$ under an equivalent measure $\tilde{\mathbb{P}}$. Let $A := \{X \leq x\}$. Then.

$$\begin{aligned}\tilde{\mathbb{P}}(A) &= \int_A \tilde{f}_X(x) dx \\ &= \int_A \tilde{f}_X(x) \frac{f_X(x)}{f_X(x)} dx \\ &= \int_A \frac{\tilde{f}_X(x)}{f_X(x)} f_X(x) dx \\ &= \mathbb{E}(\gamma(X)1_A).\end{aligned}$$

for $\gamma(x) = \tilde{f}_X(x)/f_X(x)$. This random variable $\gamma(x)$ is the conversion factor and is known as the Radon-Nikodym derivative; if we look at it very carefully, we recognize a likelihood ratio. Let us illustrate the concept in the following Example.

Example 134 Consider a closed interval $[a, b]$, for $0 \leq a \leq b \leq 1$ and consider the experiment of choosing a number from this interval. Suppose we have the following two metrics on the interval $[a, b]$, $0 \leq a \leq b \leq 1$

$$\begin{aligned}\mathbb{P}(\text{the number chosen is in } [a, b]) &= \mathbb{P}[a, b] := b - a \\ \mathbb{P}^*(\text{the number chosen is in } [a, b]) &= \mathbb{P}^*[a, b] := b^2 - a^2.\end{aligned}$$

We could be more specific and say that

$$\begin{aligned}\mathbb{P}[a, b] &= \int_a^b dx; \\ \mathbb{P}^*[a, b] &= \int_a^b 2x dx.\end{aligned}$$

The last equation is implying that $\gamma(x) = 2x$. Indeed, in this example we have that our random variable X follows the standard Uniform distribution, so that $f_X(x) = 1$, whilst $\tilde{f}_X(x) = 2x$.

The above example shows that the distribution of the random variable can change under a measure change.

Changes of probability are the cornerstone of Mathematical Finance as they allow us to price derivative instruments by defining the framework of risk neutral valuation; they also allow us to simplify significantly the complexity of the calculations for the price and the greeks of some derivative instruments via the change of numeraire. Changes of probability are also fundamental in Computational Finance, as they underpin variance reduction techniques such as Importance Sampling. The random variable $\gamma(x)$ in this context takes the name of importance weight.

8 Conditional expectation

At the beginning of this section, we talked about the problem of setting up a mathematical model of a random experiment, in order to support our decision process. Specifically, we talked about informed decisions, and we have seen that information in the probability space is captured by the σ -algebra. Then, in the previous section, we have seen how to quantify a random event by using random variables.

Now, assume as always that some random experiment is performed, whose outcome is some $\omega \in \Omega$. Imagine that we are given some information, in the form of a sub- σ -algebra of $\mathbb{F}$ say $\mathbb{G}$, about this possible outcome, not enough to know the precise value of ω , but enough to narrow down the possibilities. Then, we can use this information to estimate, although not precisely, the value of the random variable $X(\omega)$. Such an estimate is represented by the *conditional expectation* of X given $\mathbb{G}$.

In order to understand the definition of conditional expectation, we need to familiarize first with the indicator function. Precisely, we use the notation 1_A for

$$1_A(\omega) = \begin{cases} 1 & \text{if } \omega \in A \\ 0 & \text{otherwise} \end{cases}$$

Hence 1_A is a random variable which follows a Bernoulli distribution, taking values 1 with probability $\mathbb{P}(A)$, and 0 with probability $\mathbb{P}(A^c)$. Hence $\mathbb{E}(1_A) = \mathbb{P}(A)$. Properties of the indicator function are listed below.

1. $1_A + 1_{A^c} = 1_{A \cup A^c} = 1_\Omega = 1$;
2. $1_{A \cap B} = 1_A 1_B$.

Now, we are ready for the following.

Definition 44 (Axiomatic definition-Kolmogorov) *Let $(\Omega, \mathbb{F}, \mathbb{P})$ be a probability space and X a random variable with $\mathbb{E}|X| < \infty$. Let $\mathbb{G}$ be a sub σ -algebra of $\mathbb{F}$. Then the random variable $Y = \mathbb{E}(X | \mathbb{G})$ is the conditional expectation of X with respect to $\mathbb{G}$ if*

- i) $Y \in \mathbb{G}$
- ii) $\mathbb{E}|Y| < \infty$
- iii) $\forall A \in \mathbb{G} : \mathbb{E}(Y 1_A) = \mathbb{E}(X 1_A)$, i.e. $\int_A Y d\mathbb{P} = \int_A X d\mathbb{P}$.

The idea is that, if X and $\mathbb{G}$ are somehow connected, we can expect the information contained in $\mathbb{G}$ to reduce our uncertainty about X . In other words, we can better predict X with the help of $\mathbb{G}$. In fact, Definition 44 is telling us that, although the estimate of X based on $\mathbb{G}$ is itself a random variable, the value of the estimate $\mathbb{E}[X | \mathbb{G}]$ can be determined from the information in $\mathbb{G}$ (property 1). Further, Y is an unbiased estimator of X (property 3 with $A = \Omega$).

Example 135 Consider once again the stock price evolution described in Example 126. Suppose you are told that the outcome of the first stock price movement is ‘up’. You can now use this information to estimate the value of S_2

$$\mathbb{E}(S_2(\omega) | \text{up}) = 12p + 4.$$

In this case, $\mathbb{G} = A_U$. Similarly,

$$\mathbb{E}(S_2(\omega) | \text{down}) = 3p + 1,$$

and $\mathbb{G} = A_D$.

Exercise 4 Question: what is

$$\mathbb{E}(S_2(\omega) | \mathbb{G} = A_{UD})?$$

Theorem 45 *The conditional expectation has the following properties.*

1. $\mathbb{E}[\mathbb{E}(X | \mathbb{G})] = \mathbb{E}[X]$, i.e. $\mathbb{E}[Y] = \mathbb{E}[X]$.
2. If $\mathbb{G} = \{\emptyset, \Omega\}$ (smallest σ -algebra), $\mathbb{E}[X | \mathbb{G}] = \mathbb{E}[X]$.
3. If $\mathbb{G} = \mathbb{F}$, $\mathbb{E}[X | \mathbb{G}] = X$.
4. If $X \in \mathbb{G}$, $\mathbb{E}[X | \mathbb{G}] = X$.
5. If $Z \in \mathbb{G}$, then $\mathbb{E}[ZX | \mathbb{G}] = Z\mathbb{E}[X | \mathbb{G}] = ZY$.
6. Let $\mathbb{G}_0 \subset \mathbb{G}$, $\mathbb{E}[\mathbb{E}(X | \mathbb{G}) | \mathbb{G}_0] = \mathbb{E}[X | \mathbb{G}_0]$.
7. Let $\mathbb{G}_0 \subset \mathbb{G}$, $\mathbb{E}[\mathbb{E}(X | \mathbb{G}_0) | \mathbb{G}] = \mathbb{E}[X | \mathbb{G}_0]$.
8. If X is independent of $\mathbb{G}$, then $\mathbb{E}[X | \mathbb{G}] = \mathbb{E}[X]$.
9. **Jensen’s inequality.** $g(\mathbb{E}(X | \mathbb{G})) \leq \mathbb{E}(g(X) | \mathbb{G})$ for all convex functions g .

Example 136 Let $X_1, X_2, \dots$ be identically distributed random variables with mean μ , and let N be a random variable taking values in the non-negative integers and independent of the X_i . Let $S = X_1 + X_2 + \dots + X_N$. Suppose that you need to calculate the mean of the random variable S . Well, calculations here are not straightforward, since S is a random sum: the number of terms in the sum, in fact, depends on N , which is not known. For this reason you cannot apply the additive property of the expectation. The conditional expectation, though, gives you the way forward: assume that you know the value of N , then, given this knowledge, the mean of S is easy to obtain.

$$\begin{aligned} \mathbb{E}(S | N) &= \mathbb{E}(X_1 + X_2 + \dots + X_N | N) \\ &= \mathbb{E}\left(\sum_{i=1}^n X_i \middle| N = n\right) \\ &= \mu N. \end{aligned}$$

Now, what you have obtained depends on N , so all what you have left to do is to average (again), this time with respect to N , and remember property 1.

$$\mathbb{E}(S) = \mathbb{E}[\mathbb{E}(S | N)] = \mathbb{E}(\mu N) = \mu \mathbb{E}(N)$$

Example 137 Consider a Normal mean-variance mixture $Y = \mu + \beta V + \sigma\sqrt{V}Z$, for $Z \sim N(0, 1)$ and V some non-negative continuous random variable. We can now obtain the moment generating function using the properties of the conditional expectation. The key observation is that conditioned on the value that V can take, Y follows a Gaussian distribution with given mean and variance.

$$\begin{aligned}\mathbb{E}\left(e^{kY}\right) &= \mathbb{E}\left(\mathbb{E}\left(e^{k(\mu+\beta V+\sigma\sqrt{V}Z)}\middle|V\right)\right) \\ &= e^{k\mu}\mathbb{E}\left(e^{(k\beta+k^2\sigma^2/2)V}\right) \\ &= e^{k\mu}\mathbb{E}\left(e^{hV}\right)\end{aligned}$$

for $h = k\beta + k^2\sigma^2/2$. Hence, by conditioning we reduce the problem of obtaining the MGF for Y to the MGF of V . Once we define the distribution of the mixing random variable V , we know its MGF and from there we can deduce the MGF of Y . For example, suppose $V \sim \text{Gamma}(\alpha, \lambda)$. From Example 129, we know that

$$\mathbb{E}\left(e^{hV}\right) = \left(\frac{\lambda}{\lambda - h}\right)^\alpha,$$

therefore

$$\mathbb{E}\left(e^{kY}\right) = e^{k\mu} \left(\frac{\lambda}{\lambda - k\beta - k^2\sigma^2/2}\right)^\alpha.$$

Similar calculations lead to the MGF of the Normal inverse Gaussian distribution (NIG) and the t-distribution (although for the latter one, we need to know more about transcendental functions such as the Bessel function of the third kind).

Properties of the conditional expectation *: Proof. One by one.

1. Check point 3 in definition 44 for $A = \Omega$ (remember that $\Omega \in \mathcal{G}$...): $\mathbb{E}[Y1_\Omega] = \mathbb{E}[X1_\Omega]$ but $1_\Omega = 1$.
2. Check point 3 in definition 44. For $A = \emptyset$, we have

$$\int_{\emptyset} Y d\mathbb{P} = \int_{\emptyset} X d\mathbb{P} = 0$$

For $A = \Omega$

$$\mathbb{E}[X1_\Omega] = \mathbb{E}[X]$$

$$\mathbb{E}[\mathbb{E}(X)1_\Omega] = \mathbb{E}[X]$$

in virtue of property 1. Hence both sides return $\mathbb{E}[X]$.

3. Verify definition 44 on X for $\mathcal{G} = \mathcal{F}$.

- (a) $X \in \mathcal{F}$ because it is $\mathcal{F}$ -measurable by definition of random variable.
- (b) $\mathbb{E}|X| < \infty$ by assumption (axiomatic definition).
- (c) $\mathbb{E}(X1_A) = \mathbb{E}(X1_A) \forall A \in \mathcal{G}$.

In this case you have available the entire “history” of X . Hence you know everything and therefore there is no uncertainty left.

4. If $X \in \mathcal{G}$, then we go back to the same situation as depicted in (3).
5. We prove this property for the simple case of an indicator function; hence, assume $Z = 1_B$ for some $B \in \mathcal{G}$; then condition 3 in definition 44 reads

$$\forall A \in \mathcal{G} \quad \mathbb{E}(ZX1_A) = \mathbb{E}(X1_A1_B) = \mathbb{E}(X1_{A \cap B}).$$

But $\forall A \cap B \in \mathcal{G}$, condition 3 implies

$$\mathbb{E}(X1_{A \cap B}) = \mathbb{E}(Y1_{A \cap B}) = \mathbb{E}(Y1_A1_B) = \mathbb{E}(ZY1_A).$$

The extension to the case of a more general random variable relies on the construction of a random variable as the limit of the sum of indicator functions. However, this is out of the grasp of this unit.

6. Let $Y = \mathbb{E}[X|\mathcal{G}]$ and $Z = \mathbb{E}[X|\mathcal{G}_0]$. If $A \in \mathcal{G}_0$, then $\mathbb{E}(Z1_A) = \mathbb{E}(X1_A)$, but since $\mathcal{G}_0 \subset \mathcal{G}$, $A \in \mathcal{G}$ as well, and by definition $\mathbb{E}(Y1_A) = \mathbb{E}(X1_A)$. Therefore $\mathbb{E}(Z1_A) = \mathbb{E}(Y1_A) \forall A \in \mathcal{G}_0$.
7. Let $Z = \mathbb{E}[X|\mathcal{G}_0]$, then $Z \in \mathcal{G}_0$. Since $\mathcal{G}_0 \subset \mathcal{G}$, it follows that $Z \in \mathcal{G}$. Therefore $\mathbb{E}[Z|\mathcal{G}] = Z$.
8. $\forall A \in \mathcal{G} : \mathbb{E}(X1_A) = \mathbb{E}(X) \mathbb{E}(1_A) = \mathbb{E}[\mathbb{E}(X)1_A]$.
9. A (differentiable) function is convex if it lies above all of its tangents i.e.

$$g(x) \geq g(y) + g'(y)(x - y).$$

Set $x = X$ and $y = \mathbb{E}(X|\mathcal{G})$ and substitute in the above expression. Then take again conditional expectation with respect to $\mathcal{G}$ to obtain

$$\mathbb{E}(g(X)|\mathcal{G}) \geq g(\mathbb{E}(X|\mathcal{G})) + g'(\mathbb{E}(X|\mathcal{G}))(\mathbb{E}(X|\mathcal{G}) - \mathbb{E}(X|\mathcal{G}))$$

due to the measurability of $\mathbb{E}(X|\mathcal{G})$. The statement then follows.

■

Exercise 5 We define the conditional covariance between two random variables X and Y given a σ -algebra $\mathbb{F}$ by

$$\text{Cov}(X, Y|\mathbb{F}) = \mathbb{E}[(X - \mathbb{E}(X|\mathbb{F}))(Y - \mathbb{E}(Y|\mathbb{F}))|\mathbb{F}].$$

Show the law of total covariance.

$$\text{Cov}(X, Y) = \mathbb{E}[\text{Cov}(X, Y|\mathbb{F})] + \text{Cov}[\mathbb{E}(X|\mathbb{F}), \mathbb{E}(Y|\mathbb{F})].$$

Now deduce the law of total variance.