

F O U N D A T I O N S
O F
D A T A
S C I E N C E

Foundations of Data Science (Inf2-FDS)
Lecture Notes

2025

David Sterratt and Kobi Gal
with some material based on notes for Inf2B
by Hiroshi Shimodaira, Iain Murray and Steve Renals

19th January 2026

Contents

I Data: ethics, collection, representation, wrangling, exploration, visualisation and descriptive statistics	1
1 Introduction	3
1.1 What is data science?	3
1.2 Examples of data science at work	5
1.3 Data science careers and skills	5
2 Data	7
2.1 Data and metadata	7
2.2 Tabular data and variables	8
2.3 Working with tabular data	10
2.4 Data wrangling	10
2.5 Merging tabular data using database-style joins	11
2.6 Split-apply-combine	12
3 Descriptive statistics	15
3.1 Introduction to descriptive statistics	15
3.2 Distributions of numeric variables	15
3.3 Sample and population mean	16
3.4 Sample and population median	17
3.5 Variance and standard deviation	18
3.6 Standardised variables	19
3.7 Quantiles	20
3.8 The mode	20
4 Introduction to data ethics	21
4.1 Introduction to ethics	21
4.2 Data protection and privacy	22
4.3 Bias	23
5 Exploratory data analysis, data communication and visualisation	25
5.1 The importance of visualisation for exploring and communicating data	25
5.2 Visualisation for Exploratory Data Analysis	25
5.3 Visualisation for data communication	27
5.4 How we create visualisations	28
5.5 Principles of visualisation	29
5.6 Effective and ineffective visualisations	35
6 Data collection and statistical relationships	37
6.1 Collecting data	37
6.2 Obtaining data already hosted online	40
6.3 Correlation	41
6.4 Limitations of statistical relationships	44
6.5 Causal Reasoning in Graphs	45

II	Linear models	47
7	Linear Regression	49
7.1	Regression as prediction	49
7.2	Linear regression	49
7.3	Visual diagnostics and transformations	55
7.4	Numerical diagnostics	57
8	Multiple regression	59
8.1	The principle of multiple regression	59
8.2	Interpreting multiple regression coefficients and metrics	60
8.3	Interaction terms and nonlinear fits	62
8.4	Interpreting and refining multiple regressions on many variables	62
9	Mathematics of multiple regression	67
9.1	Derivation of coefficients in multiple regression	67
9.2	Interpreting the equation for the coefficients	69
10	Dealing with high dimensions – PCA	71
10.1	The principle of Principal Components Analysis (PCA)	71
10.2	Principle of finding principal components	74
10.3	PCA and regression	77
10.4	Derivation of PCA	79
III	Introduction to Machine Learning	85
11	Supervised learning: Classification with Nearest neighbours	87
11.1	Classification	87
11.2	Nearest neighbour classification	89
11.3	Evaluation	91
12	k-NN, hyperparameters, metrics, cross-validation	95
12.1	k -Nearest neighbour classification	95
12.2	Metrics	98
12.3	k -Nearest neighbour regression	99
12.4	Limitations of supervised machine learning and cross-validation	100
13	Unsupervised learning: K-means	105
13.1	Clustering, unsupervised and supervised learning	105
13.2	Types of clustering	108
13.3	K -means	108
13.4	Evaluation and application of K -means clustering	112
IV	Statistical inference	115
14	Randomness, sampling and simulation	117
14.1	Introduction to statistical inference	117
14.2	Sampling, statistics and simulations	119
14.3	Distributions of statistics of small samples from probability distributions	122
14.4	The distribution of the sample mean of large samples	124
15	Estimation	127
15.1	Point estimation	127
15.2	Estimation bias and variance	129
15.3	Standard error	131

16 Confidence intervals	135
16.1 Principle of confidence intervals	135
16.2 Definition of confidence intervals	138
16.3 Method of estimating confidence interval for the mean of a large sample	138
16.4 Bootstrap estimation of confidence intervals	140
16.5 Interpretation of confidence intervals	143
16.6 Confidence intervals for the mean from small samples	143
17 Hypothesis testing and p-values	147
17.1 Principle of hypothesis testing	147
17.2 p -values	149
17.3 Testing for goodness of fit to a model	151
17.4 Issues in hypothesis testing	153
18 A/B testing	155
18.1 The principle of A/B Testing	155
18.2 Increasing certainty in A/B testing	157
18.3 Large sample theory of A/B testing	157
18.4 Issues in A/B testing	159
18.5 Comparing groups with numeric responses	160
18.6 The theoretical method of testing for differences between groups	161
18.7 Quantifying the effect size of differences between two numeric samples	162
18.8 Paired data – to appear	163
18.9 Relationship between hypothesis testing and confidence intervals	163
19 Statistical inference and regression	167
19.1 Inference about linear regression coefficients with the bootstrap	167
19.2 Understanding stats package linear regression output	169
19.3 Sampling theory inference about linear regression coefficients	171
19.4 Derivation of standard error of estimator for slope coefficient	171
V The Maximum Likelihood Principle and Regression	173
20 Logistic regression	175
20.1 Principle of logistic regression	175
20.2 Interpretation of logistic regression coefficients	177
20.3 Multiple logistic regression and confidence intervals	178
20.4 Logistic regression as a classifier	179
21 Maximum likelihood and generalised linear regression	183
21.1 The principle of maximum likelihood	183
21.2 Maximum likelihood principle applied to a simple example	184
21.3 Maximum likelihood estimation of linear regression coefficients	186
21.4 Maximum likelihood estimation of logistic regression coefficients	190
21.5 Generalised linear models	191
21.6 From maximum likelihood to Bayesian inference	191
22 Ethical issues with supervised learning	193
22.1 Fairness in machine learning	193
22.2 Overview of equality legislation	194
22.3 Case Study: The effect of gender on credit scoring	194
VI Project skills	199
23 Software engineering for data science	201
23.1 Reproducible research	201
23.2 Notebooks versus programs	202

23.3 Data and code management	203
24 Writing skills	205

About these lecture notes

These lecture notes are written for the University of Edinburgh Course [Informatics 2 – Foundations of Data Science](#). Ultimately we aim that they will give comprehensive (though perhaps not exhaustive) coverage of the *material* covered in the lectures.

Although the notes are designed to fit in with the Foundations of Data Science course, we're making them an [open educational resource](#) available under a Creative Commons licence. The development of this course has benefited from other open educational resources. Since the notes are fairly well developed (though still not complete or error free!) it seems right to give back to the world.

How to use these notes

The process of learning involves more than just reading the material, and we strongly encourage students on Foundations of Data Science to take part in the learning activities that are designed to go with these notes: the lectures themselves, “comprehension questions”, computer labs, tasks and workshops, and the coursework. You should look at the [course website for details of all of these activities](#). In these notes we've linked to the labs, which are also openly available.

In these notes, we've highlighted activities you should do, recommended activities, examples and non-examinable extra information using the boxes shown below.

Essential reading

This reading is essential – you should do it, ideally before reading the chapter, or going to the lecture.

Related Workshop

You should attend the workshops linked to from these boxes.

Related Python Lab

You should do the lab Jupyter notebooks mentioned in these boxes. The notebooks contain exercises. We encourage you to try to work out the answers to the exercises, but if you get stuck, you can reveal the hints and solutions contained in the notebooks.

Important

You should pay particular attention to these boxes.

Recommended reading

We recommend you read the materials mentioned in these boxes.

 Example

We recommend you follow the examples and worked examples in these boxes to deepen your understanding.

 Exercise

We suggest you try these exercises to help check and deepen your understanding.

 Warning

Pay attention to these boxes, which warn you of a common mistake, misconception or pitfall.

 Further reading (not examinable)

You may find this reading interesting, and it could deepen your understanding, but it's not examinable.

 Note (not examinable)

These notes cover concepts that are either of historical interest, or relate the concepts in this course to concepts beyond the course.

Here's one example warning:

 These notes may contain errors

There are doubtless mistakes in the notes, and we welcome reports of errors and suggestions for improvements.

The philosophy of Foundations of Data Science

It is often said that 80% of the time spent on a data science project is in preparing the data, and the quality of the data is crucial to the outcome of any data science or machine learning project. The philosophy of Foundations of Data Science is to allow students to acquire the foundational skills of data, before moving on to other courses (such as [Machine Learning](#)), which build on these foundations.

As stated in the [Learning Outcomes](#), we intend that on completion of this course, the student will be able to:

- Describe and apply good practices for storing, manipulating, summarising, and visualising data.
- Use standard packages and tools for data analysis and describing this analysis, such as Python and LaTeX.
- Apply basic techniques from descriptive and inferential statistics and machine learning; interpret and describe the output from such analyses.
- Critically evaluate data-driven methods and claims from case studies, in order to identify and discuss a) potential ethical issues and b) the extent to which stated conclusions are warranted given evidence provided.
- Complete a data science project and write a report describing the question, methods, and results.

Although we present (fairly informally) the mathematical foundations of methods such as linear regression, the focus of FDS is more on understanding the principles and the correct application of methods rather than the derivation or development of methods.

Much of data science is about story telling and explanation, so we consider that models such as linear regression are helpful not only for prediction, but also for explanation. We therefore cover topics such as the interpretation of linear and logistic regression coefficients.

We first approach statistical inference from the point of statistical simulations and sampling theory, inspired by the computational approach to statistical inference in the [Berkeley Data 8 course](#), and its associated textbook. Given the ubiquity of using theoretical distributions to generate confidence intervals and undertake hypothesis testing, we also introduce some common statistical distributions and tests.

The structure of the notes

We have organised these notes into 6 separate parts, each containing multiple chapters:

Part I: [Data: ethics, collection, representation, wrangling, exploration, visualisation and descriptive statistics](#)

Part II: [Linear models](#)

Part III: [Introduction to Machine Learning](#)

Part IV: [Statistical inference](#)

Part V: [The Maximum Likelihood Principle and Regression](#)

Part VI: [Project skills](#)

This is a slightly different way to the structure envisaged in the [FDS course descriptor](#), but all the topics indicated there are included. Interleaved with these topics are topics focusing on real-world implications (often using case studies), critical thinking, working and writing skills.

The order of the chapters represents one ideal ordering of the material. For logistical reasons the order of the lectures in a particular year may differ from this order. The material on statistical inference early comes later in the course, so that students have met the relevant concepts of probability in [Discrete Maths and Probability](#).

Influences

Material on supervised learning and clustering mostly comes from the work of Steve Renals, Iain Murray and Hiroshi Shimodaira on the previous Informatics course Inf2B.

We are indebted to other openly-available data science courses, in particular [Harvard's CS109](#), in particular the concept of the data science life cycle, and the [Berkeley Data 8 course](#), with its emphasis on programming statistical simulations for inference.

[Gelman and Nolan's \(2017\)](#) book *Teaching statistics – a bag of tricks* has provided the basis for presentation of some of the concepts and lecture demos.

Resources

There is no one textbook that covers all this course. However, we will refer to the following books and resources at points throughout the course:

Modern Mathematical Statistics with Applications (Devore and Berk, 2012): *Modern Mathematical Statistics with Applications* is an introduction to inferential statistics that follows on from the Discrete Maths and Probability course. It's freely available as a PDF from the library – see the Library Resources link. You can also purchase a print copy for £25 via the page for the book that you reach via the Library.

The Big Book of Dashboards: Visualizing Your Data Using Real-World Business Scenarios (Wexler et al., 2017)
The first chapter of this book provides a concise introduction to principles of visualisation.

An Introduction to Data Ethics (Vallor, 2018) This is an introduction to data ethics from the perspective of *virtue ethics*. It presents case studies to help you develop your ethical sensitivity and ethical reasoning.

Computational and Inferential Thinking (Adhikari et al., 2020): The online textbook *Computational and Inferential Thinking* takes a computational, non-mathematical approach to statistical inference, and the statistical inference sections of FDS have drawn inspiration from their approach. Please note that Python code in the book uses a bespoke python library instead of pandas, which we'll be using on this course, so don't pay too much attention to the details of the code.

***Becoming a critical thinker: for your university studies and beyond* (Ivory, 2021)** This book, available online in the University Library, is an excellent general introduction to critical thinking. We recommend the whole book, but if you have little time, read Chapters 4–6 (on arguments, evidence and communication). This reading will help you to understand the sort of thinking, writing and referencing that we would like to see in the Project.

Notation

The following is the default notation used in the course. Note that different versions may be used sometimes for the ease of readability.

x	A scalar
$\mathbf{x}$	A column vector : $\mathbf{x} = (x_1, x_2, \dots, x_D)^T$, where D is the dimension of vector
$\mathbf{x}^T$	Transpose of vector $\mathbf{x}$, meaning a row vector if $\mathbf{x}$ is a column vector
x_i or x_d	i th sample (scalar), or d th element of vector $\mathbf{x}$
$\sum_{i=1}^N x_i$	Summation, i.e., $x_1 + x_2 + \dots + x_N$
$\prod_{i=1}^N x_i$	Product, i.e., $x_1 x_2 \dots x_N$
$\mathbf{x}_i$	i th sample (vector): $\mathbf{x}_i = (x_{i1}, x_{i2}, \dots, x_{iD})^T$
$\ \mathbf{x}\ $	Euclidean norm or L^2 norm: $\ \mathbf{x}\ = \sqrt{\sum_{d=1}^D x_d^2} = \sqrt{\mathbf{x}^T \mathbf{x}}$, also known as magnitude
$\ \mathbf{x}\ _1$	L^1 norm, i.e. $\sum_{d=1}^D x_d $
$\mathbf{u} \cdot \mathbf{v}$	dot (inner or scalar) product of $\mathbf{u}$ and $\mathbf{v}$, i.e., $\mathbf{u} \cdot \mathbf{v} = \ \mathbf{u}\ \ \mathbf{v}\ \cos \vartheta = \mathbf{u}^T \mathbf{v} = \sum_{d=1}^D u_d v_d$
$\mathbf{A}$	A matrix: $\mathbf{A} = (a_{ij})$. If $\mathbf{A}$ is a M -by- N matrix, $\mathbf{A} = (\mathbf{a}_1, \dots, \mathbf{a}_N)$
A_{ij}	Element of matrix $\mathbf{A}$ at i th row and j th column, i.e. a_{ij}
$\mathbf{I}$ or $\mathbf{I}_d$	Identity matrix of size d by d (ones on diagonal, zeros off)
$\mathbf{A}^T$	Transpose of matrix $\mathbf{A}$, i.e. $\mathbf{A}^T = (a_{ji})$
$\mathbf{A}^{-1}$	Inverse of matrix $\mathbf{A}$, i.e. $\mathbf{A}^{-1} \mathbf{A} = \mathbf{A} \mathbf{A}^{-1} = \mathbf{I}$
$\{\mathbf{x}_i\}_1^n$	A set of samples: $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$
$\bar{x}$	Sample mean of x
s^2, s_x^2 , or $\text{Var}(x)$	Sample variance of x
s_{xy}	Sample covariance of x and y . NB: $s_{xx} = s_x^2$
μ	Population mean
σ^2	Population variance
Σ	Population covariance matrix (cf. summation operator $\sum$)
σ_{ij}	(i, j) -element of a covariance matrix, i.e. $\Sigma = (\sigma_{ij})$. NB: $\sigma_{ii} = \sigma_i^2$
$E[X]$	Expected value of the random variable X
$V[X]$	Expected variance of the random variable X
$\exp(x)$	The natural exponential function of x , i.e. e^x
$\ln(x)$	The natural logarithm of x
$x \propto y$	x is proportional to y

This notation is also used by *Modern Mathematical Statistics with Applications* (Devore and Berk, 2012).

Part I

**Data: ethics, collection, representation,
wrangling, exploration, visualisation and
descriptive statistics**

Chapter 1

Introduction

1.1 What is data science?

Data science is all about exploring, explaining, and inferring insights from vast amounts of data. The end-goal of this course is to provide students with tools they require to be great data scientists, the lucky practitioners that develop and/or use data-scientific technologies! So let's first define what a data scientist does.

What is a data scientist? Data scientists are able to turn raw data into understanding, insight and knowledge. Some of the activities that data scientists do on a daily basis include the following:

- Find, collect, check and clean data.
- Explore data and infer knowledge and insights.
- Explain and interpret these inferences.
- Communicate inferences about the data to others.
- Make prediction about future trends.

This is why data scientists need to employ interdisciplinary tools, from statistics, artificial intelligence and machine learning. We will be addressing all of these different aspects in the course, and have closely aligned our learning goals with the varied skill set that data scientists need to be successful in the field.

The data science process We can see data science as a process, comprising interconnected steps (Figure 1.1). We will expand on all the steps of the data science cycle in this course.

The data explosion: volume, variety and velocity Data is always around us, and it always has been. But the development of technology has made it possible to collect and store vast amounts of data. Just to get some perspective, each day on Earth we generate over 500 million tweets (how many of these are generated by bots?), 294 billion emails, 4 million gigabytes of Facebook data, 65 billion WhatsApp messages and 720,000 hours of new content added daily on YouTube.

In 2018, the total amount of data created, captured, copied and consumed in the world was 33 zettabytes (ZB) – the equivalent of 33 trillion gigabytes. This grew to 59ZB in 2020 and is predicted to reach a mind-boggling 175ZB by 2025. One zettabyte is 8,000,000,000,000,000,000 bits (Vopson, 2021).

Parallel to this explosion of data generation, there is also consistent improvement in computational methods; we are now able to mechanise some aspects of data analysis that data scientists care about. So, while the amount of digital data in the world doubles every two years, so does the computational power that lies at our disposal to analyse it. The speed of computation, which was commonly considered to double every two years, may have plateaued.

How does data science relate to other fields? “Data Science”, “Statistics”, “Machine Learning”: these are some of the terms used in the scientific literature and popular media. Let's try to sort through the confusion.

Data science is a multidisciplinary field which uses scientific methods, processes, and systems in a range of forms.

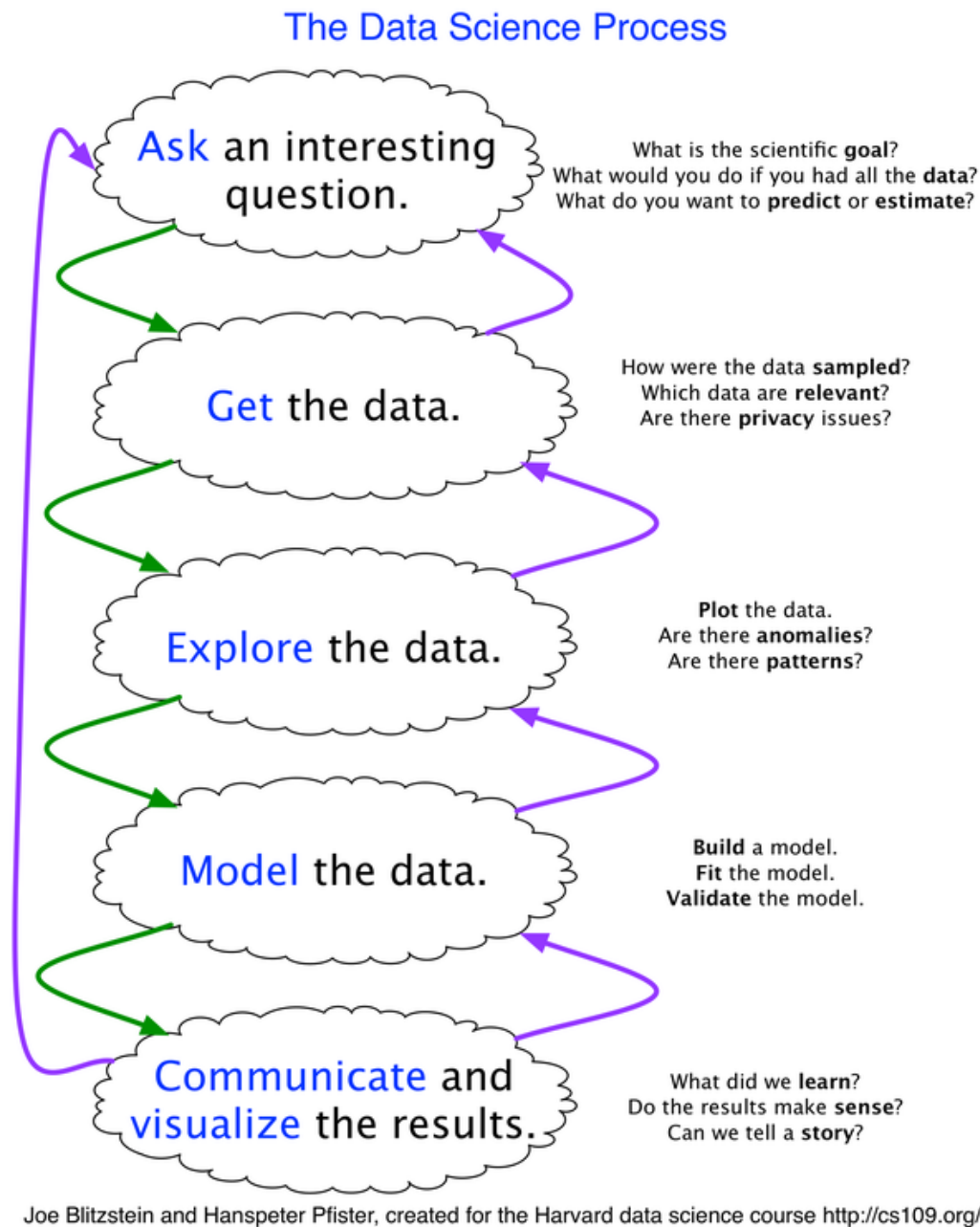


Figure 1.1: The data science process. Credit: Joe Blitzstein and Hanspeter Pfister, created for the [Harvard data science course CS109](#). This figure is not under the [CC BY-SA](#) licence covering the rest of the notes.

Statistics contains two main areas: (1) summarising data (descriptive statistics) and (2) quantifying uncertainty in data and determining if the data supports hypotheses (inferential statistics). Both statisticians and data scientists care about analysing and explaining data. Indeed, some of the methods used by data scientists are taken from statistics, and we shall learn them in this course.

Machine learning is all about algorithms that are able to learn from data to make predictions. Data science also uses some tools from machine learning, and we will study some of these in the course.

Before using methods from statistics and machine learning, data scientists need to do a lot of work on processing the data itself and check that the data makes sense. Data scientists often deal with huge databases, which is why they rely heavily on computational methods for their analysis.

Data scientists begin by exploring the data and formalise questions that can be asked on the data. They care about explaining how (and possibly why) trends in the data arise. They need to communicate quantitative and qualitative arguments to the public. Often, they also care about supporting practitioners in the field (e.g., alerting teachers to struggling students based on their interactions with educational software). The value of a machine learning algorithm is measured by its performance on unseen data. But the value of a data science analysis also needs to consider the human in the loop. We're (happily) not at the stage when these complex tasks can be replaced with computer algorithms. There is a lot of room for skill and creativity that cannot be automated.

In all these fields, it's important to know where the data comes from, how it's been collected, and to be able to reason about whether a dataset has been collected ethically, or if the project we are undertaking is ethical.

1.2 Examples of data science at work

Let's consider two examples of data science at work. Both of these examples are meant to highlight the inherent biases that unfortunately exist in society, and that also arise in the datasets that we generate.

- Example of gender tropes in films. In this example we will see the gender bias reflected in script writing in movies.
- Case study: COMPAS. In this example we will describe racial bias in the criminal justice system.

1.3 Data science careers and skills

As well as being crucial to undertaking scientific work in academia, industry and the public sector is increasingly applying and developing data science methods, in a wide range of areas: for example environmental monitoring, energy networks, healthcare, manufacturing and the finance sector.

Demand is high for data professionals. For example, a 2021 UK Government report estimated that there are potentially 178,000 data specialist roles to be filled in the UK compared to at most 10,000 data scientists coming from UK Universities (Harriss et al., 2023).

As well as bringing about social benefits, the increased automation that data science methods help are likely to bring about social changes, such as the reduction in certain types of jobs. There are also potential harms that applying data science methods might cause: for example loss of privacy due to data breaches or linking data and algorithmic decision-making systems that discriminate. It is also important that people affected by automated systems can understand how decisions are being made about them.

Therefore, the skills needed for data analysis and modelling, are both technical, such as data visualisation and programming; and non-technical, such as communication, creative thinking, and data ethics (Harriss et al., 2023). This course aims to give you the foundations of all of these areas.

Chapter 2

Data

2.1 Data and metadata

What is data? Since the 17th Century, the word **data** has referred to “things known or assumed as facts, and made the basis of reasoning or calculation” (OED). Data according to this definition has existed for 1000s of years, for example in the form of handwritten tables of scientific measurement, census records and financial accounts such as those found in Ancient Egypt (Figure 2.1). This data is all collected by humans, and is typically written down. The amount of data it is possible to collect is therefore low, and any calculations performed on the data had to be done by hand.

Since the mid-20th century, the word has also referred to “the quantities, characters or symbols on which operations are performed by computers and other automatic equipment, and which may be stored and transmitted in the form of electrical signals, records or magnetic, optical or mechanical recording media” (OED). Thus, the definition of data has expanded. Digital data comes in many forms, for example images, sound files, emails and documents, scientific databases and medical records. Some of these forms of data are collected by machine or automatically – for example, a digital camera sets creates the set of numbers that describe an image; a web server records the details of the every visit to a web page. The amount of data it has been possible to collect has increased massively, as has our ability to process that data.

Nevertheless, in this age of big data, humans are still collecting smaller datasets. For example, opinion pollsters might survey 1000 people in a poll, or ecologists might measure the height and weight of a sample of 100 squirrels (Figure 1.1).

Structured and unstructured data There is an important distinction is between structured data and unstructured data. In **structured data**, the data is organised according to a clear structure. An example of structured data is a table describing attributes of Star Wars characters with columns such as name, “height”, “weight” and eye colour. Structured data needn’t be a table: for example a bibliographic database, with

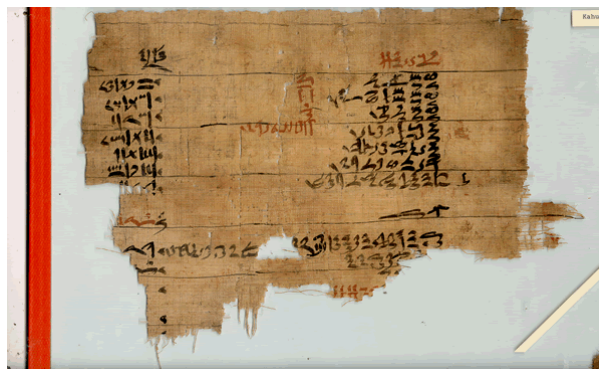


Figure 2.1: Ancient Egyptian grain account from between 2025–1700 BC found in the Lahun Papyri. Acquisition UC32189, Digital Egypt for Universities, University College London, <https://www.ucl.ac.uk/museums-static/digitalegypt/lahun/papyri.html>. Courtesy of the Petrie Museum of Egyptian and Sudanese Archaeology, UCL.

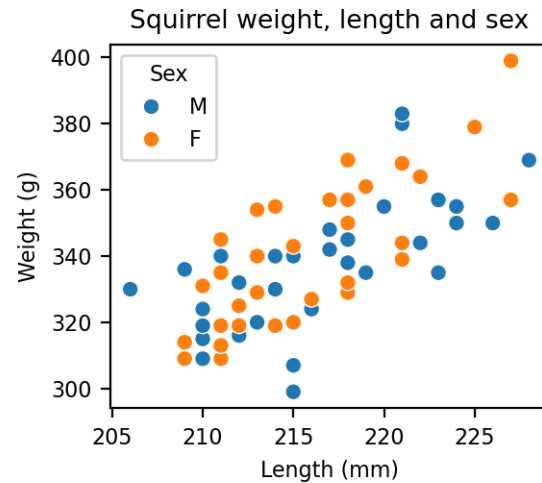


Figure 2.2: **Left:** Red squirrel, *Sciurus vulgaris*. **Credit:** Peter Trimming / CC BY 2.0). **Right:** Weight in grams versus length in millimetres of a sample of squirrels recorded in the winters of 1985 and 1986 in coniferous woods in North Belgium (Wauters and Dhondt, 1989).

Table 2.1: Characteristics of some imaginary squirrels.

Name	Weight (g)	Length (mm)	Sex	Age
Jakub	320	211.0	Male	Under 1 year
Fiona	342	222.0	Female	1–2 years
Cameron	330	215.0	Male	2+ years

different types of entries for books and academic articles. It is easy to query structured data: for example, find all characters in Star Wars with blue eyes.

Unstructured data does not have a predefined structure. Text and images are common examples of unstructured data. It is typically much harder to extract information from unstructured data than from structured data. For example, in a block of text written about Star Wars characters, would find it hard to identify all characters with blue eyes.

Metadata We contain terabytes of data files, but they are useless if we know what they *mean*. **Metadata** – literally “about data” – is information describing the data files it accompanies. The metadata may be a description in a text file (often called README) or it may be in a structured format. Whenever a dataset is created, stored and shared, the data meaning of the data should be recorded. It is surprising how often the meaning of datasets is not clearly described. The metadata should describe how the data were collected, including if there were any legal or ethical considerations, the format of the data files and the meaning of variables in the data files.

2.2 Tabular data and variables

Tabular data In this course, we focus on a common class of structured data, **tabular data**. Tabular data is data arranged in a table. The meaning of each cell in the table is determined by the column and/or row labels. Table 2.1 shows an example of tabular data.

Tidy Data Tidy data is a subset of tabular data (Wickham, 2014). In a tidy dataset, each column corresponds to a **variable** (or **attribute**) and each row corresponds to an **instance** or **observation** possessing each of the variables. The terms **data matrix** or **long form data** are synonyms for tidy data.

For example, suppose we capture squirrels and measure their weight (in grams) and length (in millimetres), and note down their sex. The variables represented in the columns would be “Name” (assuming we’ve named the squirrels we are observing), “Weight (g)”, “Length (mm)” and “Sex”. Each row contains the value of these variables.

Table 2.2: Messy data with values as columns. In the messy table (top), the rows correspond to one variable (“Sex”) and the columns “Under 1 year old”, “1 to 2 years old” and “More than 2 years old” are actually the values of “Age”, a categorical variable. This format is easy to read, but it is not tidy, as the values of a variable are in the column names. We can rearrange to a tidy format (bottom) in which there is one column for each of the variables “Sex” and “Age”, and a column “Count”.

Messy data	Sex	Under 1 year old	1 to 2 years old	More than 2 years old
	Male	4	7	2
	Female	6	9	1

Tidied data	Sex	Age	Count
	Male	Under 1 year old	4
	Female	Under 1 year old	6
	Male	1 to 2 years old	7
	Female	1 to 2 years old	9
	Male	More than 2 years old	2
	Female	More than 2 years old	1

Typically, tidy data is **multivariate**, i.e. it has multiple variables. In the special case of two variables we refer to it as **bivariate**, and when we consider only one variable, we call it **univariate**.

Messy data We call data that is not in the tidy format messy. There are a number of ways in which data can be messy, for example:

- The values of variables could be column headings, as shown in Table 2.2. Here the messy format is easy for humans to read, but is often more difficult to process. Functions such as `melt` in Pandas can be used to rearrange this messy format into a tidy format.
- A single observational unit could be stored in multiple tables. For example, data on the CO2 emissions for countries over time is stored in per-continent files, e.g. `annual_co2_europe.csv`, `annual_co2_asia.csv`... To tidy the data here, we would load in each file, give it a column “Continent” with the value for each row derived from the file name, and then concatenate these files.

Messy data may actually be easier for the human eye to comprehend, but is not so easy to manipulate using software tools, so for the moment we will assume we have tidy data.

Types of variables Variables can be of various types:

- **Numerical variables** are quantities that can be measured (continuous) or counted (discrete). For example, weight and length are continuous numeric variables. In contrast, the number of babies a squirrel gives birth to in a year is a discrete variable, since we cannot have a fractional number of babies. Continuous variables often have a physical dimension (e.g. weight or length) and it is important to quote them with their unit.
- **Categorical variables** can take on one of a number of values. For example, the sex of the squirrel is a categorical variable, since it can be “Male” or “Female”. We can have more than two categories, for example a fruit might be “Apple”, “Orange”, “Lemon”, “Grapefruit” etc.
- **Ordinal variables** can take on one of a number of categories, but those categories are ordered. For example, we might estimate the age of a squirrel as “Under 1 year old”, “1 to 2 years old” or “More than 2 years old”. There are only three categories, but we can order them from youngest to oldest.
- **String variables** contain string information such as names or a comment from a survey form.

Representing categorical variables The human-readable format of categorical variables is as a string. However, for many of the algorithms that we deal with later, the string representation is not helpful. We could represent the category by a number, e.g. 0 for “Apple”, 1 for “Orange”, 2 for “Lemon” etc. But this system implies an ordering for the categories, which is not the case here.

Instead, what we do is convert the categorical variable into an **indicator variable**, also known as a dummy variable. We create one new column for each category (for example, an “Apple” column, an “Orange” column and a “Lemon” column). We indicate which category the item belongs to by putting a 1 in the corresponding columns, and zeros in the other columns. This form of encoding of categories is also known as **“one-hot” encoding**, since there is always 1 “hot” bit required to indicate the category.

In fact, if we have k categories, we can manage with $k - 1$ columns, by dropping one column, which we regard as the default. For example, if we dropped the “Apple” column, when there were zeros in all the remaining columns, we would assume that the item represented was an “Apple”.

2.3 Working with tabular data

Reading stored data Tabular data can be stored in a variety of formats, and Data Science packages (for example Pandas or R) have functions to read in this data:

- Text file: common formats are comma-separated variable (CSV) and tab separated variable (TSV). There are various standards for the precise formatting of the files, for example if the data enclosed by cells are enclosed by quotes. Sometimes there are a few lines of metadata at the top of the file. Data science package functions to read in text files have many options, for example allowing you to skip lines at the head of the file, or deal with separators other than tabs or commas.
- Binary file: common formats are Excel spreadsheets or Open Document Format spreadsheets. The modern versions of both of these are in fact ZIP files containing an XML file with the spreadsheet information, and any other files (e.g. embedded images).
- Databases: for example MySQL or SQLite. Here there is a database server, and data is extracted by running SQL (Structured Query Language).¹

Text files have the advantage of being human-readable, and also enforce a discipline of encoding information solely in the content of the cells – it is not possible to encode information by colour-coding cells, as it is in spreadsheets. Conversely, spreadsheets have the advantage of being format-able, which can help with readability.

Selecting rows and columns Data science packages have methods for extracting rows and columns from tables. With tidy data, extracting a row selects all the data connected with one observation. Extracting a column selects every instance of one variable.

Filtering data We call finding a subset of the data based on the value of some variable **filtering**. For example, we may wish to filter out all the squirrels longer than 220 mm.

2.4 Data wrangling

Before you can undertake data science analyses, you need to get (or “wrangle”) the data into a suitable form, a process that is called **data wrangling**.

Cleaning data The quality of any analysis of data can only be as good as the data itself – and the data itself may have many types of problems:

- Data entry: for example, if we have collected data in a free-text survey, respondents may not have entered a time in a uniform format (e.g. “16:00”, or “4pm”, or “4” or “16” or “17 (CET)”) may all mean the same thing. Or someone may have typed “08” when they meant “80”.
- Mixing text and numbers: for example, we might have recorded “210g” or “0.21kg” in the weight cell for our squirrels.
- Missing data: perhaps a sensor wasn’t working for a few minutes or hours, so some readings are missing. Does it record “0” in this situation? Or “-1”. The metadata should tell us, but maybe it doesn’t. Missing data can affect inferences from data.

¹You can learn how to process and analyse data using SQL in [Introduction to Databases](#).

Table 2.3: Time taken to complete obstacle course by squirrels

Name	Date	Time (s)
Fiona	2021-04-06	67.5
Fiona	2021-04-10	50.2
Cameron	2021-04-08	55.6
Lily	2022-07-13	45.0

- **Mislabelled data:** A column may have been mislabelled. For example, we might be recording the temperature of heating water entering and leaving the Informatics Forum, and the temperature of heating water entering and leaving Dugald Stewart Building (DSB). To get a measure of how much heat the Forum is using, we subtract the temperature of the water leaving from the water entering, and the same for DSB. But if we subtract the temperature of water leaving DSB from water entering the Forum, we will get nonsense. (This scenario actually happened.)
- **Duplicated data:** perhaps someone has made a copy-and-paste error in a spreadsheet.
- **Faulty data:** perhaps a sensor goes wrong, and starts giving values that are implausible.

In summary, there are many potential problems with data: **the data is out to get you!** One of the most important jobs of a data scientist is to check the data quality, and fix problems. You need to approach the data in the spirit of critical evaluation: Is this value reasonable? Is that pattern strange? Does the data look too good?

We use the term **data cleaning** to describe the process of checking and fixing problems in data. You should start data checking and cleaning after loading the dataset, but problems with data also emerge in the process of visualisation, which we will come to later.

Data cleaning is very time-consuming: it is commonly said that 80% of the time of a data science project is spent on cleaning (Dasu and Johnson, 2003). Some more recent estimates suggest time spent cleaning is less than 30% (Anaconda, 2020), though the same report also estimates loading and visualising data take around 20%. In any event, data cleaning can take a long time, needs to be done carefully, and can be quite fiddly and frustrating.

Representing missing data Modern data science packages have special values – NA or NaN (**Not Applicable** or **Not a Number**) – to describe data that is missing. It’s very helpful to ensure that data that you regard as missing shows as NA or NaN rather than as a default value (e.g. 0 or –1), as this can help with filtering out missing data.

Pandas is somewhat confusing in its handling of missing data. Strictly speaking, NA applies to missing string, categorical or numeric data, whereas NaN applies only to missing numbers. However, Pandas represents missing data as NaN, but functions to check for missing data refer to NA, e.g. `.isna()`.

Once missing data is identified, you have to decide what to do with it, which may depend on the analysis you are undertaking. Sometimes it can be possible to keep observations that contain missing data, but some analyses only work if every variable in every observation has a value. In this case, you may need to drop any rows containing NaN.

2.5 Merging tabular data using database-style joins

Often we wish to combine data from two sources that relates to the same entities. For example, we might have recorded the times that some of our squirrels could undertake an obstacle course. Now, suppose we wanted to compare the times of completion to the characteristics of the squirrel (Weight, Length, Sex and Age). Both Table 2.1 and Table 2.3 share a common variable: “Name”. We can match up the rows of both tables by using the `merge` function in Pandas with “Name” as the **key** that binds the tables. In relational database terminology the equivalent operation is called a **join**.

Notice that not all the squirrels in the Table 2.1 undertook the obstacle course, and there is one squirrel (“Lily”) who is present in Table 2.3 but not in the first table. There are various ways of merging to deal with this mismatch, described in the next paragraphs.

Inner join In an **inner join** (Table 2.4), only the squirrels in both datasets will be present in the joined dataset. Note that the key Fiona in Table 2.1 matches two rows in Table 2.3, so there are two corresponding

Table 2.4: Results of inner join applied to Tables 2.1 and 2.3.

Name	Weight (g)	Length (mm)	Sex	Age	Date	Time (s)
Fiona	342	222.0	Female	1-2 years	2021-05-06	67.5
Fiona	342	222.0	Female	1-2 years	2021-05-10	50.2
Cameron	330	215.0	Male	2+ years	2021-05-08	55.6

rows in the joined table. The data describing Fiona’s characteristics is repeated in these rows, but the unique information about the data and time of the obstacle course run are not repeated. Jakub isn’t present in this table, since Jakub hasn’t had a time recorded on the obstacle course in Table 2.3.

Table 2.5: Results of left join applied to Tables 2.1 and 2.3.

Name	Weight (g)	Length (mm)	Sex	Age	Date	Time (s)
Jakub	320	211.0	Male	Under 1 year	nan	nan
Fiona	342	222.0	Female	1-2 years	2021-05-06	67.5
Fiona	342	222.0	Female	1-2 years	2021-05-10	50.2
Cameron	330	215.0	Male	2+ years	2021-05-08	55.6

Left and right joins In a **left join** (Table 2.5), all squirrels present in the first (left) table passed to the merge function will be retained in the merged table. When the key isn’t present in the second table, we fill in the missing values with NaN. In Table 2.5 Jakub has NaN values for data and time, since Jakub hasn’t had a time recorded on the obstacle course in Table 2.3.

In a **right join** all items in the right-hand table are retained, and the resulting table will have NaN values when a key is present in the right-hand table but not the left-hand table.

Table 2.6: Results of outer join applied to Tables 2.1 and 2.3.

Name	Weight (g)	Length (mm)	Sex	Age	Date	Time (s)
Jakub	320.0	211.0	Male	Under 1 year	nan	nan
Fiona	342.0	222.0	Female	1-2 years	2021-05-06	67.5
Fiona	342.0	222.0	Female	1-2 years	2021-05-10	50.2
Cameron	330.0	215.0	Male	2+ years	2021-05-08	55.6
Lily	nan	nan	nan	nan	2022-07-13	45.0

Outer join In an **outer join** (Table 2.6): All squirrels in both datasets will be present in the joined dataset. When the key isn’t present in the first or second table, we fill in the missing values with NaN. In Table 2.6 Jakub and Lily are present because they are present either in Table 2.1 (Jakub) or Table 2.3 (Lily). Jakub has NaN values for data and time, since Jakub hasn’t had a time recorded on the obstacle course, and Lily’s age, height, weight etc. are missing, since Lily didn’t appear in Table 2.1.

2.6 Split-apply-combine

Split-apply-combine A commonly used operation in data science is **split-apply-combine**. It involves:

- Splitting the data into groups based on some criteria.
- Applying a function to each group independently. This could include computing a summary statistic for each group, such as a mean or count; performing some group-specific computation such as standardising the data within a group; as well as filtering to discard data that has only a few members, for example.
- Combining the results into a data structure.

For example, we might want to find the fastest time for male and female squirrels. In this case, we could split the joined table (Table 2.4) by Sex, apply the `min` function to each subtable, and then combine the results in one table (Table 2.7).

Table 2.7: Results of **splitting** by Sex, **applying** `min()` and then **combining**, applied to the inner join (Table 2.4).

Sex	Time (s)
Female	50.2
Male	55.6

Advantages of split-apply-combine Data science programming environments, such as R or the Pandas package in Python have functions that can achieve split-apply-combine operations in one line. We could write code to achieve the same effect, but it would be less clear, at least to data scientists who are familiar with data science programming idioms.

The split-combine-apply paradigm is also found in big data, where it is referred to as **MapReduce**. Here the groups created by the split operation can be processed in parallel before being combined.

Related Python Lab: Introduction to Jupyter notebooks and Pandas

<https://github.com/Inf2-FDS/FDS-S1-01-introduction>

In this lab you will learn the very basics of the Python library `pandas`, which is used for data management. By the end of the lab you should be able to:

- use a Jupyter notebook
- load different data file types
- display data
- filter your data for specific values, and
- apply basic statistical computations on the data.

We also share the Inf2-IADS Lab 1, which covers:

- The Python interpreter
- Basic data types: Numbers (int and float), strings, booleans, lists (including list comprehension), tuples, sets and dicts.

Related Python Lab: Data wrangling with Pandas

<https://github.com/Inf2-FDS/FDS-S1-02-data-wrangling>

In this lab you will learn to prepare your data for future use. By the end of the lab you should be able to:

- get an overview of your data,
- clean data, and
- combine data from multiple datasets.

Related Python Lab: Data wrangling II – groupby and regular expressions

<https://github.com/Inf2-FDS/FDS-S1-05-data-wrangling-02>

In this lab, we will build on what you have previously learned on preparing data for future use. By the end of the lab you should be able to:

- use **regular expressions** (regexps) to parse textual data

- organise columns into indicator variables when appropriate
- apply group-wise computations to your data – i.e. Split-apply-combine (groupby)

You'll also encounter simple web scraping with `pd.read_html`

Chapter 3

Descriptive statistics

Recommended reading

Modern Mathematical Statistics with Applications, Sections 1.1, 1.3 and 1.4

3.1 Introduction to descriptive statistics

Statistics The German word *Statistik* arose in the 18th century and originally referred to “data about the state” (country). The first use of “statistical” in the English language was in 1791 in the *Statistical Account of Scotland*. Sir John Sinclair, an elder in the Church of Scotland, sent a questionnaire to ministers in every parish (church district) in Scotland. The questionnaire asked many questions about agriculture, industry, economics, employment, poverty and education, as well as “The state of the manners, the morals, and the religious principles of the people”. In fact empires and dynasties have been collecting data about population and trade for much longer than this, going back to the Han dynasty in China and the Roman Empire.

Descriptive statistics It’s not humanly possible to make sense of a large raw dataset. For example, suppose we know the salary of every member of staff in the University of Edinburgh – the list would be very long. To make sense of this data we can try to summarise it or visualise it. **Descriptive statistics** refers to methods of summarising data. This topic introduces the notation we’ll use for the sample and population mean and variance of numeric univariate data. We will also introduce quantiles and skewed distributions.

3.2 Distributions of numeric variables

We will focus here on numeric **univariate data**, i.e. data comprising one numeric variable, for example the weights of a number of squirrels that we have captured, weighed, and released. A quick way to start to understand the data is to visualise how it is distributed using a **histogram** (Figure 3.1). Each bar in a histogram covers an interval of the variable known as a **bin** – for example in Figure 3.1 the bins are 10 g wide. The area of each bar is proportional to the **frequency** of items observed within its bin, i.e. how many items are observed within the bin.

There are four ways of determining the heights of the bars:

- A **frequency histogram** displays the frequency on the y -axis. From this plot we can see that there was 1 squirrel between 290 g and 300 g in weight, 4 squirrels between 300 g and 310 g, and so on.
- The **relative frequency histogram** sets the height of each bar to the number items in its bin expressed as a fraction or percentage of the total number of items. The sum of the heights of all bars is equal to 1.
- In a **frequency density histogram** the y -axis shows the frequency per unit on the x -axis, so that the area of each bar of equal to the frequency of items in the bin. In the example here, because each bin is 10 g wide, the frequency density (measured in Squirrels per gram) is a tenth of the raw frequency. The total area under the histogram is equal to the total number of items.

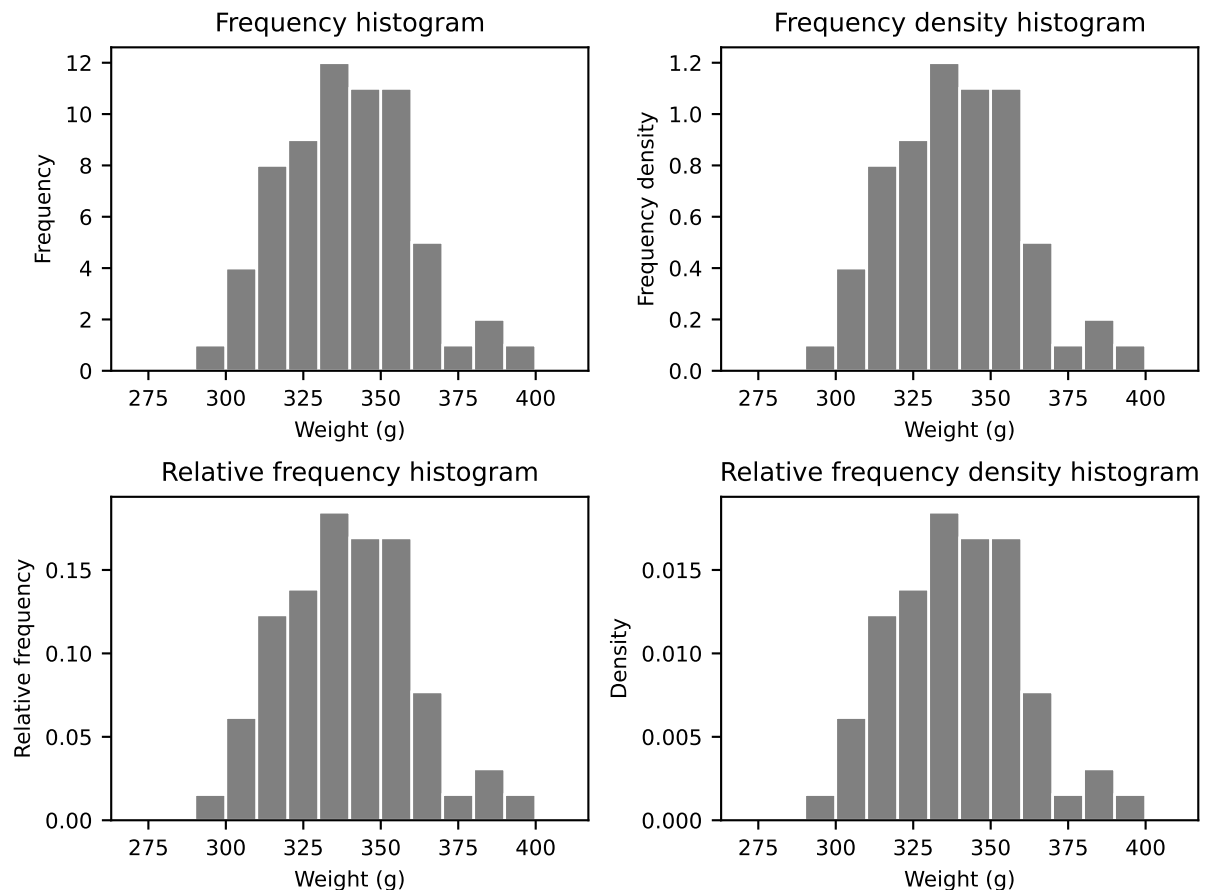


Figure 3.1: Histograms of weight in grams of a sample of squirrels recorded in the winters of 1985 and 1986 in coniferous woods in North Belgium (Wauters and Dhondt, 1989). The bin width in both histograms is 10g.

- In the **relative frequency density histogram**, the total area under the histogram is equal to 1. The height of each bar corresponds to the relative frequency density (sometimes **density** for short) of finding an item in that range. The histogram can be called the **empirical distribution** of the data.

⚠ Histograms with unequal bin sizes

It is possible, and often sensible, to construct bins of different sizes. For example in ranges of value where data is quite sparse, larger bins can help to create a visually smoother histogram. However, with bins of varying sizes, it is no longer true that the area of the bars in frequency and relative frequency histograms is proportional to the frequency. However, the area is still proportional to frequency in the case of frequency density and relative frequency density histograms.

If the histogram has one clear peak we say it is **unimodal**. If it has two peaks it is **bimodal**. Histograms with multiple peaks are **multimodal**. Figure 3.1 appears to be bimodal, but the peak between 380g and 390g is small, and there are not many individuals in this sample, so we should be careful to assume that this distribution is truly representative of all squirrels.

3.3 Sample and population mean

Populations and samples It's important to distinguish between **populations** and **samples**. The population is the set of all the things we are interested in, for example, all 400 Scottish wildcats in the Highlands (Figure 3.2). The sample is a subset of the population that we observe – for example 10 wildcats that we trap, measure and release again into the wild.



Figure 3.2: Scottish wildcats, *Felis silvestris silvestris*, a critically endangered species. Credit: Peter Trimming / CC BY 2.0)

We refer to the size of the population with N and the size of the sample with n . Note that we don't always know N exactly. In the case of the wildcats, $N = 400$ is an estimate – there is no practical way of counting all the wildcats in Scotland. In other cases we do know N exactly, for example if the population were a pile of exam papers.

Definition of sample mean For a numeric variable x with n observations or instances sampled from a population, $x_1, x_2 \dots x_n$, the **sample mean** is defined as:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i \quad (3.1)$$

Sometimes the sample mean is written informally as $1/n \sum x_i$. When reporting the mean of a set of numbers, one convention is to report to one more decimal place than the accuracy of the x_i 's. For example, if the age of 6 cats is 3, 4, 5, 6, 6 and 7 years, the mean would be reported as 5.2 years, not 5.1666 years. Note that the units of the mean should be quoted; i.e. "5.2 years" *not* just "5.2".

The sample mean is a measure of where the centre of the set of instances is. It is guaranteed to lie between the minimum and maximum x_i . It has the same units as the x_i .

Population mean For a numeric variable x with N instances, $x_1, x_2 \dots x_N$, the **population mean** is defined as:

$$\mu = \frac{1}{N} \sum_{i=1}^N x_i \quad (3.2)$$

With bivariate data (two variables) or multivariate data (more than one variable), we can distinguish between the population means of variables x, y, z by using subscripts: e.g. μ_x, μ_y, μ_z .

The population mean is a measure of where the centre of all instances in the population is. It is guaranteed to lie between the minimum and maximum x_i . It has the same units as the x_i .

The sample mean is an estimate of the population mean. We will consider how good an estimate it is when we learn about [Statistical inference](#). For now, it is enough to know that it depends on how the sample is chosen (randomly or by some other method), and the values of n and N .

3.4 Sample and population median

Definition of median The **sample median** $\tilde{x}$ of a variable x is the "middle" value, when the sampled observations x_i are ordered from smallest to largest. To be more precise, if there are n observations, then the

sample median is defined:

$$\tilde{x} = \begin{cases} (\frac{n+1}{2})^{\text{th}} \text{ ordered value, if } n \text{ is odd} \\ \text{mean of } (\frac{n}{2})^{\text{th}} \text{ and } (\frac{n}{2} + 1)^{\text{th}} \text{ ordered values, if } n \text{ is even} \end{cases} \quad (3.3)$$

For example, the median age of 6 cats aged 3, 4, 5, 6, 6 and 7 is 5.5 years. The median age of 5 cats aged 3, 4, 5, 6 and 7 is 5 years. Note that we should quote the units, as for the mean.

By analogy with the population mean, the **population median** $\tilde{\mu}$ of a variable is the median of the entire population.

Median and mean The mean and median of a sample or population are generally not the same. For example, the mean age of 6 cats aged 3, 4, 5, 6, 6 and 7 years is 5.2 years, but the median is 5.5 years.

- If a distribution is **symmetric**, $\bar{x} = \tilde{x}$
- If a distribution is **positively skewed** or **right-skewed**, $\bar{x} > \tilde{x}$
- If a distribution is **negatively skewed** or **left-skewed**, $\bar{x} < \tilde{x}$

Outliers Suppose the age of the cats had been 3, 4, 5, 6, 6 and 18. The mean age would now be 7 years, but the median is unchanged at 5.5. An instance that appears to be far away from most of the other numbers is called an **outlier**. The example shows that the median is less affected by outliers than the mean. For this reason, the median can be seen as a better way of measuring a typical value of a variable.

It is often worth checking outliers, to make sure that they are real data. Depending on how the data has been collected, an outlier might be due to a faulty sensor, or a mistake in data entry or in the logic of an automated programme collecting data. However, outliers may well be real data, and should not just be removed as a matter of course.

3.5 Variance and standard deviation

Definition of sample variance and sample standard deviation For a numeric variable with n observations or instances, $x_1, x_2 \dots x_n$, the **sample variance** is defined as:

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 \quad (3.4)$$

The **sample standard deviation** is defined as:

$$s = \sqrt{s^2} = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2} \quad (3.5)$$

The sample variance and standard deviation give a measure of how spread out the data is. It's an average of a measure of distance of each point from the sample mean (we'll come to why we divide by $n-1$ rather than n later).

One measure of distance we could use is the magnitude (absolute value) of the **deviation from the mean** of each observation: $x_1 - \bar{x}, x_2 - \bar{x}, \dots, x_n - \bar{x}$. We cannot just use the deviations, since they add up to 0 (think about it!). The magnitude of each deviation $|x_i - \bar{x}|$ is guaranteed to be positive. However, the average of magnitudes is not as nicely behaved mathematically as the average of the square of the deviations, as defined in Equation 3.4.

It's important to quote the units of the standard deviation and variance. The standard deviation has the same units as the quantity in question and the variance has those units squared.

Definition of population variance and population standard deviation By analogy with the population mean, for a numeric variable in a population of N instances, $x_1, x_2 \dots x_N$, the **population variance** is defined as:

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2 \quad (3.6)$$

The **population standard deviation** is defined as:

$$\sigma = \sqrt{\sigma^2} = \sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2} \quad (3.7)$$

Why the divisor $n - 1$ in the sample variance? In short, we'd like the sample variance to be a statistically unbiased estimate of the population variance (we define estimator bias formally in the chapter on [Estimation](#)). The squared deviations between the sample mean and the observations $(x_i - \bar{x})^2$ will tend to be smaller than the squared deviations between the population mean and the observations $(x_i - \mu)^2$, so if we divided by n we'd underestimate the sample variance. For the truly interested, we prove that if we divided by n we would have a biased estimate later ([Estimation bias and variance](#)).

A second way of thinking about this is that we know that the sum of the deviations is 0:

$$\sum_{i=1}^n (\bar{x} - x_i) = 0 \quad (3.8)$$

So if we know all but one $(n - 1)$ of the deviations, we can use the above equation to deduce the deviation we don't know. We say that this means there are $n - 1$ **degrees of freedom**, and it turns out that it makes sense to divide by $n - 1$.

In practice, when n is large, the difference doesn't matter much. It's worth being aware that various Python packages have different conventions about dividing by $n - 1$ or n . `pandas.Series.std` divides by $n - 1$ whereas `numpy.std` divides by n . This behaviour can be changed by specifying the `ddof` parameter in either function.

Scaled quantities Sometimes quantities can be scaled, for example if the units change. If a variable $y = cx$, where c is a scaling constant, then the following relationships hold:

$$\begin{aligned} \bar{y} &= c\bar{x} \\ s_y^2 &= c^2 s_x^2 \\ s_y &= cs_x \end{aligned} \quad (3.9)$$

We'll leave that as an exercise for you to prove.

Another way of writing the variance It's sometimes helpful to rearrange the summation over the squared deviations in the sample or population variance:

$$\begin{aligned} \sum (x_i - \bar{x})^2 &= \sum (x_i^2 - x_i \bar{x} - \bar{x} x_i + \bar{x}^2) \quad \text{expanding} \\ &= \sum x_i^2 - \sum x_i \bar{x} - \sum \bar{x} x_i + \sum \bar{x}^2 \quad \text{splitting up the summation} \\ &= \sum x_i^2 - n\bar{x}\bar{x} - n\bar{x}\bar{x} + n\bar{x}^2 \\ &= \sum x_i^2 - n\bar{x}^2 \end{aligned} \quad (3.10)$$

3.6 Standardised variables

Definition A fundamental concept for machine learning and statistics is **standardised variables**. The standardised version of a variable x is, by convention, denoted by z , and is also referred to as a *z-score*. The standardised version of i th instance is defined as:

$$z_i = \frac{x_i - \bar{x}}{s_x} \quad (3.11)$$

💡 Example of standardised variables

The mean length of squirrels shown in Figure 3.1 is $\bar{x} = 216.185$ mm and the standard deviation is $s_x = 5.238$ mm. Therefore, the standardised length of a squirrel that is 210 mm long is

$$z = \frac{x - \bar{x}}{s_x} = \frac{210 - 216.185}{5.238} = -1.181$$

Therefore the standardised length of this squirrel is -1.181 .

Properties of standardised variables The standardised variable z_i has several nice properties:

1. It has zero mean, i.e. $\bar{z} = 0$
2. It has unit variance, i.e. $s_z^2 = 1$
3. It is dimensionless (has no units) – for example, in the example above, the standardised length is not given the units of mm.

🔧 Proving that standardised variables have zero mean and unit variance

As an optional exercise, prove the first two properties above.

3.7 Quantiles

Definition of a percentile The y th percentile of a set of numeric observations $x_1, \dots, x_N$ is the value of x_i that is above $y\%$ of the values. For example, a baby that is on the 95th percentile for weight will weigh more than 95% of other babies.

The median as the 50th percentile By definition 50% of observations are less than the median, so we could also think of the median as the 50th percentile.

Lower and upper quartiles The 25th percentile is called the **lower quartile** (since it encloses the lower quarter of the distribution) and the 75th percentile is called the **upper quartile**. The difference between the upper and lower **quartiles** is called the **interquartile range**. The interquartile range is a measure of the spread of the distribution of values.

Quantiles Percentiles and quartiles are all examples of the general concept of **quantiles**. q -Quantiles are the values that divide the population or sample into q separate nearly equally sized groups. Percentiles are 100-quantiles and quartiles are 4-quantiles. Other common uses are deciles (10-quantiles) and quintiles (5-quantiles).

3.8 The mode

Definition of the mode The **mode** is the most frequent element in a set of data. In contrast to all the summary statistics described so far, the mode can be computed for categorical data as well as numeric data. For example the most popular name given to baby boys in Scotland in 1964 was David – David was therefore the mode of baby boys' names in 1964.

There can be two or more modes in a dataset, when there are two or more elements with the largest frequency. Note that a bimodal distribution does generally have two modes, since one peak may be shorter than the other (see for example Figure 3.1).

Chapter 4

Introduction to data ethics

Essential reading

Parts 1 and 2 of *An Introduction to Data Ethics* by Shannon Vallor

4.1 Introduction to ethics

Data ethics Even if we're not aware of them, ethical issues arise in our daily interactions with each other and with technology. A way into thinking about data ethics is the potential benefits and harms of data science practices in scenarios.

Example: Consequences of incorrect classification of personal data By agreeing to the use of service agreement for technology, we opt in (often unwittingly) to allowing institutions to take control of our personal data and apply algorithms to them. Although this is legal, it compromises our privacy, and can lead to harmful consequences. For example, Google has a system to detect images containing child pornography in photos processed or stored by its systems – a system which could have tremendous benefits in preventing harm to children. However, this system can produce false positives, as in the case of a man who took a picture of his sick son to send to the doctor and was labelled as a potential child pornographer, and had his Google account blocked (Hill, 2022).

Example: AI-assisted tax collection In 2022, the French government started applying a very basic AI algorithm to Google Maps images to find the pools much faster than could be done by humans checking the images. It consequently made several million euros in late tax returns from people who had swimming pools they had not reported (Willsher, 2022). We can consider there were ethical benefits from the tax collected being redistributed through society, and the new system was fairer than the previous one. But suppose the system was extended to detect other types of undetected structures such as verandas or gazebos, and this system produced false positives?

The ethical ramifications of algorithms are a significant issue that is affecting all of us. We will focus on these issues from different aspects.

What is ethics? The word **ethics** derives from the Ancient Greek word *ethos* meaning "character, personal disposition". Ethics can be regarded as the study of moral phenomena, and addresses questions such as "how best do I live?" and "how best do I act in a given situation". Ethical principles have been around since the age of Greek philosophers. They provide us with a kind of checklist of what is right and what is wrong in society. Our job as data scientists is to understand and apply that checklist to some of the stages in the data science life-cycle.

Universal and context-specific ethical guidelines Some ethical guidelines are so universal that they're naturally understood, for example, "Thou shalt not kill". That's a rule of behaviour that has been around for thousands of years, and that is a guide to good behaviour that people in society should be able to follow. Other types of ethical behaviour are contextual and depend on the culture, the person and the way in which they are distributed and applied and because ethical reasoning is so subjective. For example, in

some countries, no-one stands in a queue. But in the UK that would be considered inappropriate behaviour. So culture and context matter. When building a data science program, it's very difficult to understand the context.

The ambiguity of moral decisions – the trolley problem People are often ambiguous about what are good moral decisions, depending on the context. The trolley problem is one of the tools that philosophers use to talk about the moral dilemmas that come about with ethical reasoning. A runaway train is hurtling towards 5 people on a railway track. You look around, but there's no way of warning them. You notice that you're standing next to a lever that you can operate. You can divert the train onto another track and save the people. But there is a problem: there is also a person on the other track. If you hit the lever and divert the train, then they die. If you do nothing then 5 people die, but if you pull the lever then 5 people are spared but 1 person dies. About 50% of the people in different places all over the world were asked how they would behave in this hypothetical scenario are willing to switch the points, for the net benefit of saving four people.

Philosophers consider variations of this problem to understand how people think about different contexts. Now imagine there are no points any more, but there's a man standing on top of a bridge. Now the only way to save these five people is by pushing the man on the tracks, thereby stopping the train with his body. The number of people willing to push the man off the bridge is fewer than those willing to switch the points, even though the net effect is the same (saving 5 people at the expense of one person).

Relevance of the trolley problem This is much research on understanding how should autonomous vehicles should behave in situations when they might kill people. Suppose the brakes have failed in an autonomous vehicle that is approaching a junction that a school bus is crossing. In order to avoid the school bus, the vehicle needs to slam into a tree and kill the driver. If you think about utility, that is the rational outcome¹.

When surveyed, most people agreed the autonomous vehicles should take moral decisions for the benefit of society and the common good. But people in the survey weren't willing to buy such a car. The trolley problem has a real, strong relevance to how computers should be making decisions. This example demonstrates that data ethics is a wide, surprisingly complex area of research that requires philosophers and computer scientists and machine learning people to work together.

Ethical sensitivity and ethical reasoning **Ethical sensitivity** is an awareness that a particular situation may pose an ethical dilemma. **Ethical reasoning** is the process of reasoning about the ethics of a particular situation: what could the potential consequences of a decision be?; who might benefit or be harmed?; how could any harms be avoided or mitigated? Ethical sensitivity and ethical reasoning are relevant to all the stages of the data science life cycle, from when data is collected, how data is used, how data is distributed, and even how it is controlled. We shall touch on several of these aspects. We do not expect you to be expert data ethicists by the end of this course, but we do hope you will have developed ethical sensitivity and had some experience of ethical reasoning.

4.2 Data protection and privacy

Privacy The UK is networked with more cameras on the streets than any other Western country, and those cameras are used by the police and other enforcing law agents to help keep the peace. We as the public are willing to give up some of our privacy in return for order and discipline in our lives.

But how far are we willing to go to maintain order? Would we also allow cameras into our home, for example, if we knew that they could be used by police to help track potential criminals? Many of us may not allow the use of surveillance cameras in our home. Indeed, concerns over privacy, like many social issues, are culture-dependent. This is why there is no "one size fits all" approach to privacy.

GDPR The EU has a very successful list of regulations on data protection and privacy called the **General Data Protection Regulation Rights** or GDPR <https://gdpr-info.eu/>. GDPR lists regulations that concern rights of the data subjects, duties of data controllers or processors, transfers of personal data to third countries and more. Importantly, it also details liability or penalties for breach of rights. GDPR has also been widely applied outside the EU and is considered to be a gold standard in data protection.

¹You can participate in an experiment presenting similar scenarios at <https://www.moralmachine.net/>.

The Data Protection Act 2018 is the UK's implementation of the GDPR. The [UK Information Commissioner's Office](#) upholds information rights enacted under the Data Protection Act, and various other pieces of information-related legislation, including on electronic communications and freedom of information.

One provision of the GDPR allows individuals the right to access and rectify data that concerns them, and even the right to erase data. This is also known as the **right to erasure** or the **right to be forgotten**. This right can be exercised if personal data is no longer necessary for its original purpose, or that it harms someone individual interests and there is no overriding legitimate interest to keep it. For example, Germany's highest court has ruled that a German man convicted of murder in 1982 has the right to have his name removed from online search results ([BBC, 2022](#)).

Influence of algorithms on behaviour As data scientists, we should be aware that data is used by algorithms to intervene in people's lives, making them engage in behaviour that they would not have done otherwise. For example, convenience chain stores in the USA have been known to track consumer behaviour for the purpose of recommending products for them. Ride-sharing apps have been known to use intervention methods to get drivers to stay longer on the road, when they are predicted to quit and go home. Video sharing apps use recommendation algorithms to lure people to consuming more and more content, while this content becomes more radical or extreme in nature. This phenomenon has been coined "[down the rabbit hole](#)".

Throughout the 2010s, personal data belonging to millions of Facebook users was collected without their consent by British consulting firm Cambridge Analytica, predominantly to be used for political advertising.

There are several possible ways of mitigating the influence of AI algorithms on behaviour that we can apply to respect people's privacy and autonomy, such as:

- being transparent about whether AI is being used to recommend or to motivate users
- allowing people to opt out at any stage from receiving AI generated interventions.

4.3 Bias

Everyday and technical meanings of bias In everyday usage, by **bias**, we mean prejudice and discrimination, or more specifically, an inclination or prejudice for or against someone, something, or a group, especially in a way considered to be unfair. Bias has a more technical, statistical meaning to describe the systematic deviation from a true state. Bias can be exhibited not only by people, but also by computers. When it does this, it unfairly favours someone or something over another person or thing.

A particular form of statistical bias will be considered in more detail in the chapter on [Estimation](#). Here we'll focus on biases in selecting and measuring data, and algorithmic bias.

Selection bias **Selection or sampling bias** is when the sample chosen is not representative of the population under investigation. For example a lecturer might ask the set of students who appear at a 9am lecture about some aspect of how the course runs – this set of students is not necessarily representative of the population of students on the course. As discussed in the chapter on [Randomness, sampling and simulation](#), a random sample from a population is unbiased, and non-random samples may be biased.

A related example of selection/sampling bias is data that is missing, but not at random. For example, suppose a temperature sensor failed regularly during the daytime, but not at night. An estimate of the average temperature using only the available readings would be likely to be lower than the true average temperature, because night-time temperatures are likely to be lower than daytime ones. In contrast missing data at random would affect daytime and night-time readings with equal probability, and would therefore not induce bias in the average.

Confirmation bias and cherry-picking **Confirmation bias** is defined as the search for and use of information to support an individual's, beliefs or hypotheses ([Spencer and Heneghan, 2018](#)). For example a psychological experiment ([Lord et al., 1979](#)) showed that US students who were in favour of the death penalty judged the methodology of a fictional study that confirmed their belief in the deterrent effect of the death penalty more favourably than another fictional study that showed that the death penalty had the opposite effect, against their existing beliefs. Students opposed to the death penalty also showed confirmation bias, judging the study supporting their opinion less harshly than the one in conflict with their opinion. Related to confirmation bias is **cherry-picking**, the practice of seeking out aspects of the available data that support one's point of view.

Measurement bias By **measurement bias**, we mean some practical or conceptual factor of measuring data which leads to a systematic error in the data. An example of a practical measurement error is a thermometer that records a temperature 1 degree Celsius above the true temperature; the temperature is systematically higher than it should be. Note that a thermometer could randomly give a temperature that deviates from the true temperature, but if the mean deviation is zero, it is not biased.

An example of a conceptual measurement bias might be the attempt by Lombard in the 19th century to estimate the longevity of various professions by examining the profession at the age of death (Wainer, 1999). He found that students were the profession that died youngest, around the age of 20. Can we conclude from this that being a student is a dangerous profession?

Algorithmic Bias **Algorithmic bias** describes systematic and repeatable errors in a computer system that create “unfair” outcomes, such as “privileging” one category over another in ways different from the intended function of the algorithm. Algorithmic bias can emerge from many factors, including but not limited to the design of the algorithm or the unintended or unanticipated use or decisions relating to the way data is coded, collected, selected or used to train the algorithm.

Machine learning algorithms can introduce bias at different stages in the data science life cycle; they may provide biased recommendations and decisions. Noteworthy examples of machine bias that have been widely reported in the media include:

- The COMPAS (Correctional Offender Management Profiling for Alternative Sanctions) algorithm used in US court systems to predict the likelihood that a defendant would become a re-offender. This model predicted twice as many false positives for black offenders than for white offenders.
- In 2015, Amazon realised that their algorithm used for hiring employees was found to be biased against women.
- In 2019, Facebook was found to be in violation of the U.S. constitution, by allowing its advertisers to deliberately target adverts according to gender, race, and religion, all of which are protected classes under the country’s legal system.

How can we mitigate algorithmic bias? Unfortunately, there are no quick answers to this question. Humans are the ultimate source of bias, as they are the ones generating the data and writing the algorithms that use the data to learn and make decisions. There is a lot of work (outside the scope of this course) on detecting fairness and bias in machine learning. It is also clear that regulatory practices should be developed that limit algorithm discrimination and also provide guarantees such as the right to explanation. We will discuss several of these issues in the data ethics workshop.

Related Workshop: Data ethics discussion

In this workshop we will discuss two ethical case studies:

- OK Cupid web scraping
- Facebook emotional contagion experiment

Chapter 5

Exploratory data analysis, data communication and visualisation

Essential reading

The Big Book of Dashboards ([Wexler et al., 2017](#)), Chapter 1. Access via the [Resource List](#).

5.1 The importance of visualisation for exploring and communicating data

Having acquired data (see [Data collection and statistical relationships](#) and [Introduction to data ethics](#)), the next step in a data science project is to explore the data (remember the data science process, Figure 1.1). The first step in exploration is to find any problems in the data, like the ones mentioned in the section on [Data wrangling](#). We can do some basic exploration just by using packages like Pandas to show summary statistics (min, max, mean, etc) and find values categorical variables can take (as in the Lab on [Data wrangling with Pandas](#)). However, another tool is crucial: visualisation.

We define **visualisation** as the process of conveying information through graphical representations of data. The influential statistician John Tukey explained the power of visualisation thus:

The simple graph has brought more information to the data analyst's mind than any other device. It specializes in providing indications of unexpected phenomena. ([Tukey, 1962](#))

and

The greatest value of a picture is when it forces us to notice what we never expected to see. ([Tukey, 1977](#))

It is the property of forcing “us to notice what we never expected to see” that makes visualisation so powerful, both in the exploration phase of the data science process and in the communication phase. In the next two sections we will look at data exploration, and the role visualisation plays in it, and the principles of data communication, which are wider than visualisation. However, due to its central role in both exploration and communication, the largest part of this chapter will be about visualisation.

5.2 Visualisation for Exploratory Data Analysis

Exploratory Data Analysis Once we have got a dataset, the next step in a data science project is to explore the data. The objectives of **Exploratory Data Analysis** (EDA) are to:

- enable unexpected discoveries in the data, for example:
 - identifying any possible anomalies in the data, e.g. missing data or data points that appear unlikely
 - trends and patterns

- suggest hypotheses about the causes of observed phenomena
- assess assumptions on which statistical inference will be based
- support the selection of appropriate statistical tools and techniques
- provide a basis for further data collection through surveys or experiments

In exploratory data analysis, we should try to avoid any preconceived ideas about relationships in the data, but we should use our pre-existing knowledge to assess whether data make sense.

💡 Example of EDA using descriptive stats: Informatics Forum electricity consumption

Suppose we are given a dataset of electricity consumption from the Informatics Forum and asked to describe it briefly. We're told that the two columns "Time" and "Electricity consumption (kW)" should tell us the mean electrical power consumption used in the Forum within each 10 minute interval in the time period covering the years 2013 to 2024. To check for problems in the data, we first compute some descriptive statistics using the pandas `.describe()` function:

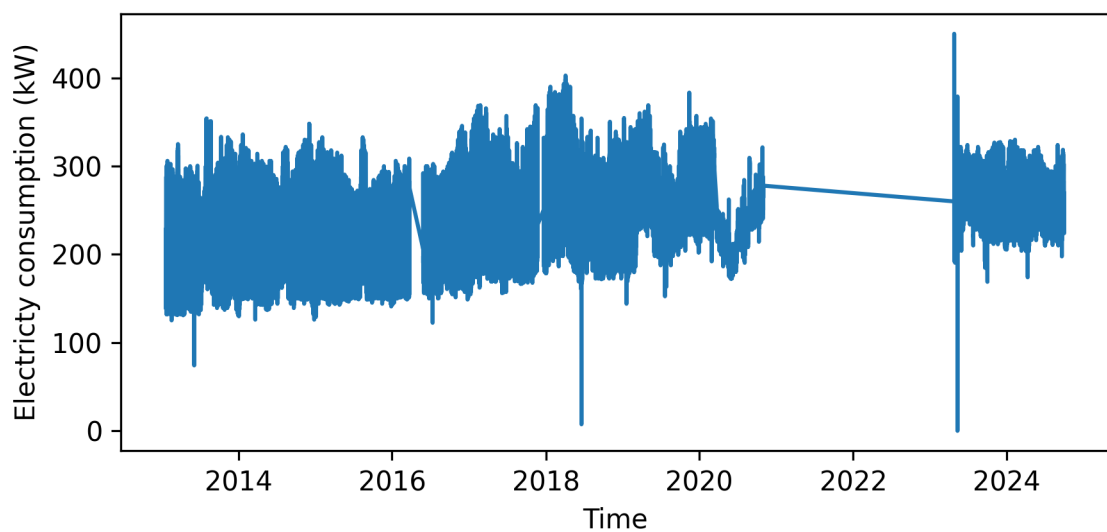
	Time	Electricity consumption (kW)
count	458771	458770.000000
mean	2018-01-26 07:20:21.862532864+00:00	228.967876
min	2013-01-21 17:40:01+00:00	0.000000
25%	2015-03-31 08:36:01.500000+00:00	192.000000
50%	2017-09-05 11:40:15+00:00	226.490066
75%	2020-01-13 11:25:03.500000+00:00	261.818182
max	2024-09-29 19:30:01+00:00	450.000000
std	NaN	45.640890

We notice that:

- The minimum and maximum values of the Time (in date format) look reasonable.
- A back-of-the-envelope calculation suggests the electricity consumption descriptive statistics are OK: We know that there are about 500 people who have offices in the Forum, so the mean electricity consumption of 228 kW (kilowatts) equates to about 0.5 kW per person. That's about 12 kWh (kilowatt-hours) in one day (24 hours). Given that we can tell from our electricity bills that we use about 3 kWh a day at home, this number of 12 kWh seems a bit high, but on the other hand, we know there are lots of servers running in the Forum, and suppose that they may be using a lot of power.
- The mean and the median electricity consumption are close, so the distribution is not skewed.
- The maximum electricity consumption (450 kW) is of the same order of magnitude as the mean – we might suspect a measurement error if it was much bigger.
- The minimum value electricity consumption is 0 kW, which could have happened in a power cut.

💡 The power of visualisation in EDA: Informatics Forum electricity consumption

We could now give a brief description of the data, but we decide to do a quick visualisation of the electricity consumption over time:



Instantly we can see that something looks odd. For most of the time the electricity consumption is clearly fluctuating, but there's a period around 2022 when it's suspiciously flat. We need to investigate what's causing the anomaly, which we didn't spot from the descriptive stats.

In exploratory data analysis, we should aim for plots that are informative and insightful, but we don't need to spend too much time on the quality of the presentation of the plots. We (or our close colleagues) are the audience for these plots, and that audience should be familiar with the data. The visualisation in the Informatics Forum electricity example does the job of showing us a problem – our task now is to investigate the problem rather than make the visualisation look better.

5.3 Visualisation for data communication

Aims of communicating data The aims of communicating data are to:

- present data and analysis honestly, accurately and clearly in a way that is appropriate for the audience
- interpret the data and analysis as fairly and objectively as possible

Data communication may provide evidence and support for a course of action and influence and persuade.

Data communication as story-telling Good data communication is like good story-telling or journalism:

- engaging
- perhaps surprising
- truthful
- understandable

We make choices about the “story” we want to tell, but we must try not to hide “what the data say”. As in a story, we need to introduce the data (the characters): where did the data come from?; who collected the data?; why? We then need to describe the data (often using descriptive stats), and describe any important relationships and trends. We don't need to give every detail – just as when describing a film plot, it can be helpful to miss out subplots that aren't crucial for the main plot.

Exercise: choices in data communication

We often have choices in data communication. Some ways of communication may be truthful, but are not necessarily helpful, as illustrated by this example.

94% of camels have one hump, with the remaining 6% having two humps. There are at least three mathematically correct ways of communicating this fact:

1. The mean number of humps camels have is 1.06
2. The median number of humps camels have is 1, with a maximum of 2
3. The mean number of humps camels have is 1.06, with a standard deviation of 0.62

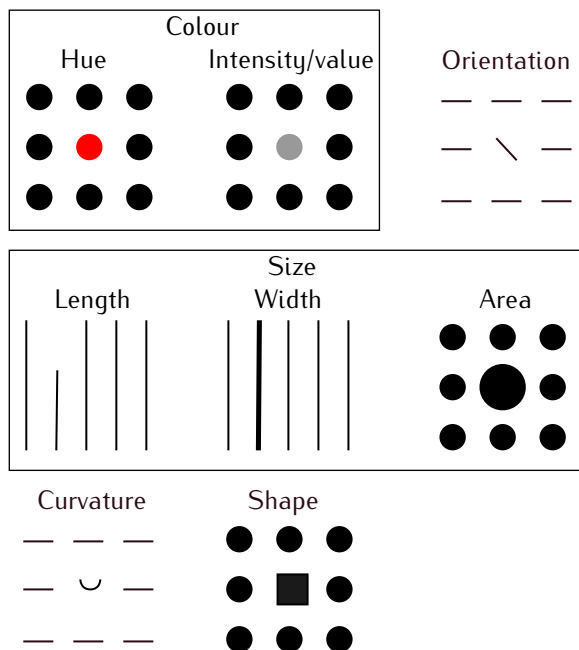
Which option would you choose? Or do you prefer the original statement of the fact?

Visualisation for data communication Visualisation is a key part of data communication. This may sound obvious, but it is surprising how a changing the presentation of a dataset can make features of the dataset become more apparent.

In contrast to using visualisations for data exploration, when communicating data, we need to focus on the presentation of the visualisation, to make it as easy as possible for our intended audience to understand. For example, the small axis labels that you could read clearly enough when you were sitting at your laptop might be illegible when displayed on slides and viewed from the back of a lecture theatre. If the viewer of a visualisation can't read the labels, the visualisation loses much of its meaning!

i Why does visualisation work?

The human visual system is very good at identifying what are called **preattentive attributes** in scenes. Below are illustrated examples of preattentive attributes. Each example demonstrates how an item with a preattentive attribute that differs from its surroundings can cause that item to “pop out”, i.e. direct our attention to it.



Different authors list different sets of preattentive attributes; for example [Wexler et al. \(2017\)](#) list 12. The illustration above contains attributes that are relevant for visualisation, and which have considerable evidence for them as fundamental features ([Wolfe and Horowitz, 2017, 2004](#))

By careful use of these features we can construct visualisations in which aspects of data “pop-out” to us. However, we can also create visualisations that mislead or confuse, for example by using too many colours or shapes. Creating visualisations that represent the data truthfully and informatively is an art that takes practice to acquire.

5.4 How we create visualisations

Before the computer age, visualisation was a painstaking process of drawing lines accurately on paper. The earliest work that we would recognise as modern visualisations were in the late 18th century by the Scottish economist [William Playfair](#) (not the architect), who created visualisations to illustrate various economic data.

Nowadays, visualisations can be created by computer packages such as Excel, or via interactive cloud-based packages such as Power BI and Tableau.

In this course we will focus (in the labs) on creating static visualisations programatically using Python's Matplotlib and Seaborn packages (though we could also have used R). Creating plots using computer code can seem more time-consuming than doing so using a graphical package. However, an advantage is that we can reproduce the steps taken to analyse data and produce a visualisation. We also have fine control over the design of our visualisations.

We focus on static rather than interactive visualisations partly because this is an introductory course; there are whole courses on visualisation, and we can only scratch the surface of what is possible. Also, static visualisations are still widely used (for example in reports, and undergraduate dissertations).

i Tidy data and the grammar of graphics

As covered in the chapter on [Data](#), in a tidy dataset, each column corresponds to a variable (or attribute) and each row corresponds to an instance or observation possessing each of the variables. Tidy data makes data exploration easy when paired with the **grammar of graphics** approach ([Wilkinson, 2005](#)).

The grammar of graphics is a general framework of producing visualisations from data, in which after various processing steps such as transformation and summarisation, variables are mapped onto attributes (also referred to as aesthetic elements) found in a visualisation. Describing the full grammar of graphics approach is beyond the scope of course, but the concept of mapping is very useful. For example, we might want to map one numeric variable to the x -axis (Position), another numeric variable to the y -axis (Position) and a third categorical variable to a colour (Hue or Value) for each category.

The python [Seaborn library](#) provides functions to create plots using this mapping principle quickly. The full grammar of graphics approach is implemented in libraries such as ggplot2 in R and python, and [Seaborn objects](#) in python.

5.5 Principles of visualisation

When learning, practising and assessing visualisation, how do you (and we) define “good”? Several sets of guidelines about good visualisation practice exist in the literature and online, though each set focuses on different aspects of visualisation and their level ranges from very general to very specific. We have constructed five principles and associated guidance that is:

- appropriate for an entry-level undergraduate data science course where students produce static visualisations;
- actionable, meaning students and markers can assess visualisations against the guidance; and
- (iii) concise enough to fit on one page, provided on a [resource on the FDS OpenCourse website](#).

We will use these principles throughout FDS, including for coursework.

The principles are mainly based on those of [Tufte \(1982, 2001, 2006\)](#) and [Schwabish \(2021\)](#), with guidance on colour from [Wexler et al. \(2017\)](#). We describe them below, expanding slightly on the handout version.

! Principle 1: Show the data

Show as much of the data as possible without making a confusing visualisation. There are multiple ways of representing the same dataset, and no “right” answer.

The following guidance plot should help show as much of the data as possible:

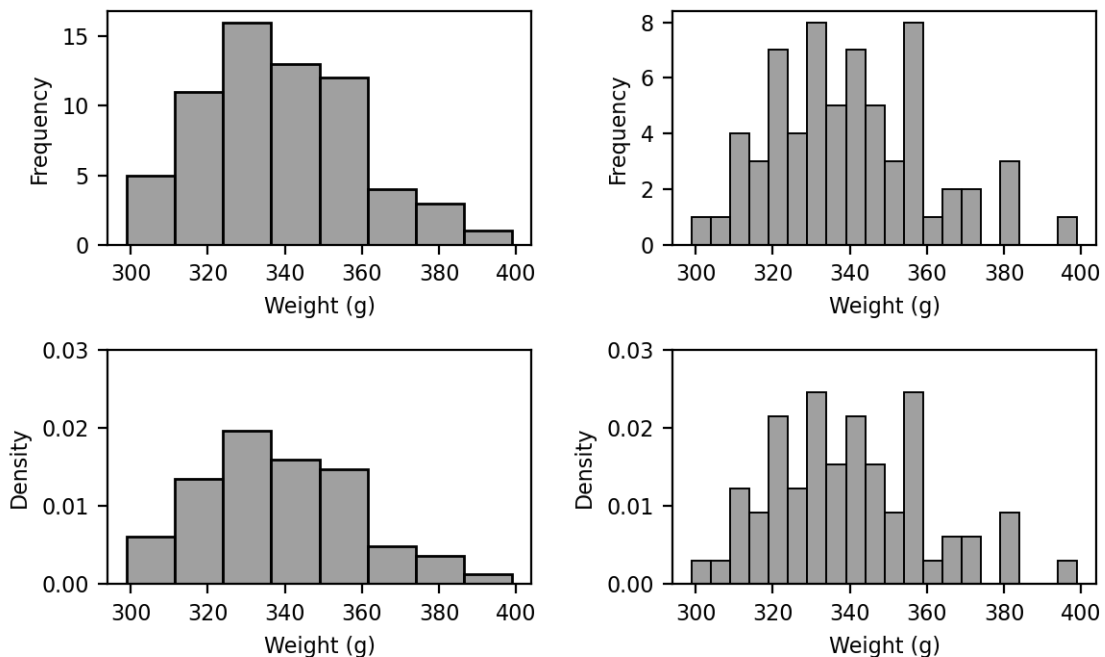
- **Identify the variables and their type.** For tabular data in a tidy (long) format, each column corresponds to a numeric, categorical, ordinal or unstructured variable.
- **Choose an appropriate plot type to show one or two variables.** Use the guide at <https://www.data-to-viz.com/>. E.g.:
 - One numeric variable → histogram or density plot shows distribution
 - One categorical variable → bar plot shows counts
 - One categorical variable and one numeric variable → bar plot can show mean of numeric variable for each category; box plot or violin plot show distribution.

- Two unordered numeric variables → scatter plot shows relationship
- One ordered and one ordered numeric variable (e.g. time series) → line plot
- **Consider showing extra variables by using length, shape, size and colour.**
 - E.g. In a scatter plot (two numeric variables), the colour and shape of each marker can represent two categorical variables, thus displaying four variables. Size can represent ordinal variables.
 - But assess whether the plot is too complex to read. adding information using marker properties can detract plot. one
- **Consider using a table.** Data patterns are generally clearer in graphics than tables. However, tables are a form of visualisation, and good for conveying raw data or dealing with large numbers of variables.
- **Use colour effectively.** (Wexler et al., 2017, pp. 14–18)
 - Choose an appropriate colour scale, depending on if the data is sequential (numeric), diverging (numeric with a zero point in the scale) or categorical.
 - Colour can also be used to highlight features in the visualisation, e.g. the largest two bars in a bar plot or the largest values in each column of a table.
- **Encourage the eye to compare several pieces of data.**
 - E.g. Use multiple plots with the same scale (“small multiples”), which can work better than using large numbers of symbols or colours on a single plot.
- **Present many numbers in a small space.**
 - E.g. A box plot of a numeric variable uses as much space as a bar plot, but conveys more information.
- **Consider how much to aggregate.** Histogram bin size or merging categories changes the patterns revealed.
- **Choose appropriate transforms.**
 - E.g. For a positive variable that varies over many orders of magnitude changes when the variable is both small and large.

💡 Examples of plot types

Histograms: An example of a histogram is the weights of squirrels (see [Distributions of numeric variables](#)).

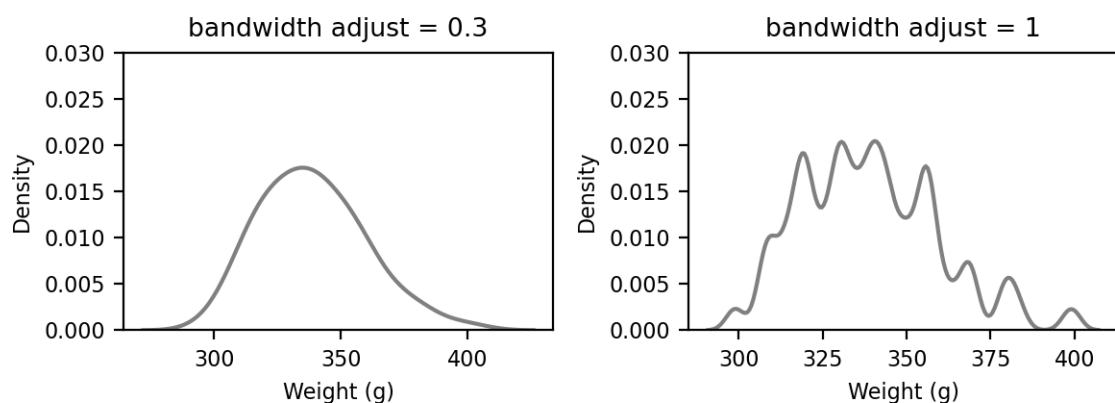
Histograms of squirrel weight showing count and density and different numbers of bins



As well as the choice of displaying frequency or density, the size of bins is your choice: up to a point smaller bin sizes show more detail, but a very small bin size will obscure the distribution. Bins can be of different sizes – as long as the height of the bars is the frequency density or the relative frequency density, the area of each bar will be proportional frequency, so the bars won't over-represent frequencies in large bins.

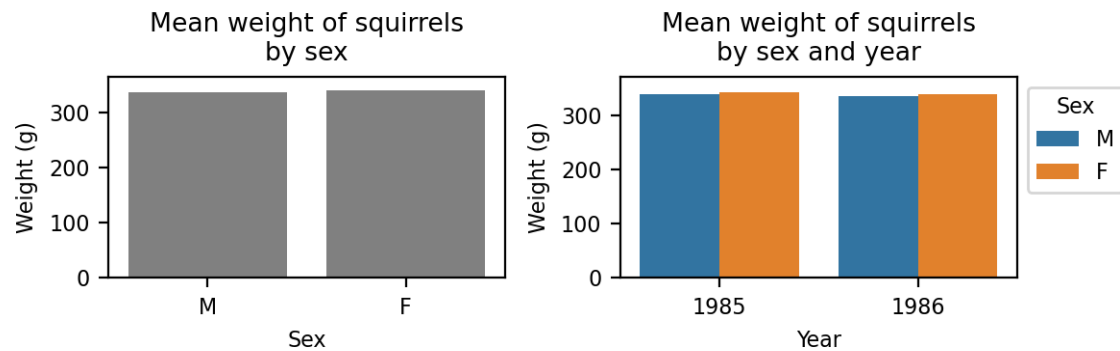
Density plots also show the distribution of a numeric variable and can be seen as a smoothed histogram.

Density plots of squirrel weight with large and small bandwidth

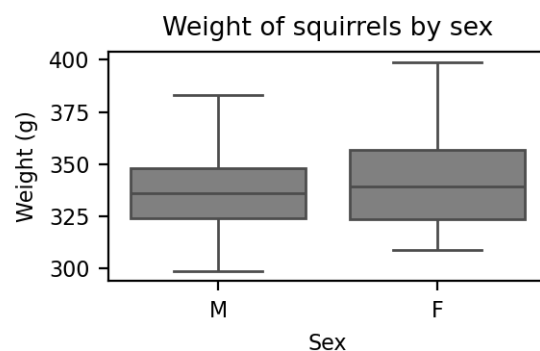


The density is an estimate of the probability density function underlying the distribution, and can be generated using kernel density estimation: the estimated density at each point is computed by placing normal distributions of a particular width (called the bandwidth) around each point. Analogous to the bin width in histograms, narrower bandwidths lead to a more bumpy appearance, and wider ones a smoother appearance. Packages such as Seaborn set the bandwidth automatically, but it's possible to adjust it.

Bar plots: e.g. the mean weight (numeric variable) of male and female (categorical variable) squirrels. We can use hue and the position of the bars to show more than one variable.



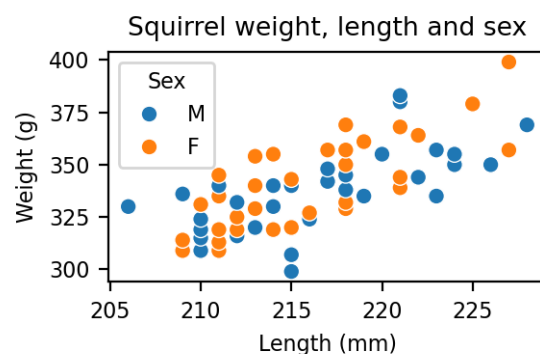
Boxplots represent the distribution of a numeric variable for multiple categories, e.g. the weights of male and female squirrels.



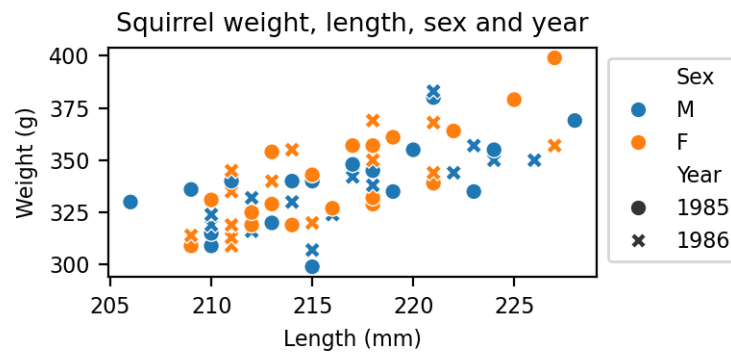
Boxplots (also known as box and whisker plots) are one-dimensional representations of a distribution in which the box extends from the lower to the upper quartiles of the data, while the line across the box represents the (see [Quantiles](#)). The 'whiskers', the lines extending from the box, can represent different things, as described in the [Wikipedia article on boxplots](#). By default, Matplotlib defines the end of the upper whisker as the value of the largest data point that lies within 1.5 times the interquartile range from the upper quartile, and the lower whisker as the value of the smallest data point that lies within 1.5 times the interquartile range from the lower quartile. Data points that lie outwith the whiskers are regarded as potential outliers, and are represented by dots or circles. Note, however, that these "outliers" may arise from a distribution with a long tail, and not be isolated from the other number as implied by the definition in [Descriptive statistics](#). Since the whiskers can represent multiple statistics, ideally their meaning should be indicated in the plot caption.

💡 Examples of showing multiple variables by using length, shape, size and colour

Using colour For example, in a scatterplot of squirrel weight versus length, we can indicate sex using colour, thus displaying 3 variables.



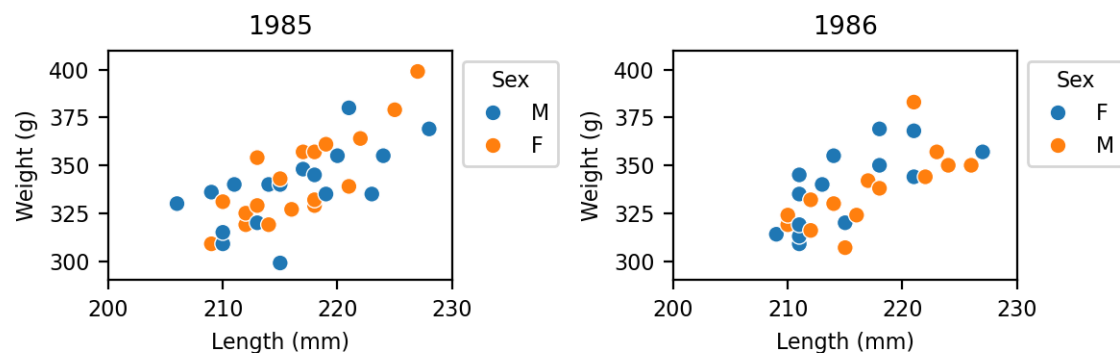
Using colour and shape: We could also indicate age categories (ordinal) by changing the size or the shape of the markers (4 variables).



However, we must be careful that adding information using marker properties does make the plot too complex to read.

💡 Encourage the eye to compare several pieces of data

This example shows the four variables (weight, length, sex and year) but may be easier to read than the plot with shapes indicating sex.



❗ Principle 2: Make the meaning of the data clear

A visualisation is meaningless if it's not labelled.

Every plot should have:

- An informative title or caption
- All variables labelled with *meaningful* text (e.g. "Price (US dollars)" *not* "x94" or "price_usd")
- Any units given (e.g. "Length (mm)" *not just* "Length")
- Scales (tick marks, colour bars) for numeric variables
- Graphical and textual annotation where appropriate – e.g. highlight a time series with known events
- Description of any error bars, e.g. 95% confidence interval, standard deviation, or standard error

❗ Principle 3: Avoid distorting what the data have to say

Design choices can mean the instant impression given by preattentive processing of the visualisation is quite different to the numbers in the dataset.

- **Use appropriate scales and baselines.**
 - In a bar chart, a non-zero baseline (the lowest point on the y -axis) makes small differences look large.
- **Be aware of limitations of human perception.**
 - Humans are better at comparing lengths than areas $\rightarrow$ consider whether area represents a given numeric variable well
 - Humans are better at comparing lengths than angles $\rightarrow$ consider alternatives to pie charts, especially with many categories

Tufte (1982) measures the level of distortion in a visualisation by the “Lie factor”:

$$\text{Lie factor} = \frac{\text{size of effect shown in graphic}}{\text{size of effect in data}} = \frac{\frac{\text{Larger visual area} - \text{Smaller visual area}}{\text{Smaller visual area}}}{\frac{\text{Maximum data} - \text{Lower data}}{\text{Lower data}}}$$

Here by “size of effect” we mean the relative change from the “control” condition. For example, if there is an increase of 50% over a control condition, we say the **effect size** is 50%. A lie factor of 1 means no distortion – the lie factor indicates how much the visualisation has exaggerated the effect size in the data.

💡 Example of lie factor calculation

Consider the left-hand panel of Figure 1 in [Kramer et al. \(2014\)](#), which shows the percentage of words written by Facebook users in a week that were classified as emotionally positive for users in a control condition, or in an experimental condition in which 10% of users’ friends’ posts containing negative words were omitted from their feeds. Suppose that the quantity measured in the control condition is y_C and the quantity measured in the experimental condition is y_E . Then

$$\text{size of effect in data} = \frac{y_E - y_C}{y_C}$$

In the “Negativity reduced” condition (left), $y_C = 5.24\%$ words, and $y_E = 5.30\%$ words. Therefore, the effect size in the data is $(5.30 - 5.24)/5.24 = 1.132\%$. However, if we look at the size of the bars in the data, the control bar 0.24 high, and the experimental bar is 0.30 high. Thus, the effect size in the graphic is $(0.30 - 0.24)/0.24 = 25\%$. Therefore, Tufte’s lie factor is $25/1.132 = 21.8$.

i Baselines and the lie factor

Suppose the baseline used in the graph is y_0 . Then

$$\text{size of effect in graphic} = \frac{(y_E - y_0) - (y_C - y_0)}{y_C - y_0} = \frac{y_E - y_C}{y_C - y_0}$$

We substitute the two equations into the equation for the lie factor to give:

$$\text{Lie factor} = \frac{y_C}{y_C - y_0}$$

If $y_0 = 0$, i.e. when the baseline of the variable we are measuring is at zero”, we can see that the lie factor is 1.

For positive y_C with a baseline above 0, the lie factor will be above 1, i.e. the size of the effect will be exaggerated. We can make the lie factor less than one (i.e. understating the size of the effect) by having a negative baseline – but this rarely happens.

! Principle 4: Make the data accessible

Make visualisations accessible so that they are meaningful for everyone.

- **Make sure text is legible**, i.e. font size of minimum 8 points in a PDF, or about 20 points in a presentation. (Surprisingly often in talks, it's impossible to read plot labels, even from the front row.)
- **Use colours that work for people with colour-vision deficiency.** [Wexler et al. \(2017\)](#), Chapter 1 has an excellent introduction to using colour in visualisations.

⚠ If you shrink or expand a figure, the font size changes

Suppose you are writing a report in which the width of the page is 8 inches, with 1 inch left and right margins. The width of the body text is therefore 6 inches. You generate a figure that has a `figsize` of (6, 4), and a `fontsize` of 8 points, save the figure to file, and include it in your document. Since `figsize=(6, 4)` means "6 inches wide and 4 inches tall", when you include the figure, the font size remains 8 points. All is well.

Suppose now you are having difficulties fitting all the numbers on the plot in. You decide to set `figsize=(12, 8)` (12 inches wide and 8 inches high) and keep the font size at 8 point. Bingo! Your numbers now fit. You put the image in the document so that it fills the width of the page. *But now the actual font size is 4 points, because you have shrunk the image by a factor of two to fit it into the page. Your plot will not be accessible to some people. This is a very common mistake.*

! Principle 5: Focus on the content

Minimise distractions for the viewer's brain.

- Avoid chartjunk – e.g. colours that don't have any meaning or 3D bar charts
- Reduce clutter – e.g. more tick labels than needed or vertical grid lines with a categorical x-axis
- Use consistent colours for plots in the same study
- Use an appropriate number of decimal places
- Check spelling is correct

5.6 Effective and ineffective visualisations

A good way to improve your visualisation skills is to learn from effective and ineffective visualisations. In his book *The visual display of quantitative information*, originally published in 1982, [Tufte \(2001\)](#) highlighted some outstanding examples of visualisation from the time of Playfair onwards, criticised some trends in visualisation that he identified as being prevalent by 1980, and proposed some principles for graphical integrity and good visualisation. His book is inspirational, but we should also bear in mind that some of the visualisations he most admires would be impossible to create in Matplotlib or Seaborn, since they lack some of the flexibility of hand-drawing or using a graphics package. Also, some of his principles were reacting against problems that now occur less often than they did then, due to the standardised nature of visualisation software.

You can find modern examples of graphical excellence on some websites – for example the plots on the [Financial Times](#) and [Guardian](#) websites. [WTF Visualizations](#) has a wide range of ineffective visualisations. You could try assessing both good and bad visualisations against the principles.

Related Python Lab: Visualisation for Exploratory Data Analysis - using Tidy Data and Seaborn

<https://github.com/Inf2-FDS/FDS-S1-03-visualisation-exploratory-data-analysis>

In this lab you will learn how to produce visualisations for exploratory data analysis. By the end of the lab you should be able to:

- explore relationships between variables with scatter and line plots
- use aesthetic elements like colour to represent variables
- explore how distributions differ across variables with histograms and box plots

Related Python Lab: Visualisation for Data Communication - Fine-tuning using Matplotlib

<https://github.com/Inf2-FDS/FDS-S1-04-visualisation-data-communication>

In this lab you will learn different plotting methods using Matplotlib and some of the techniques you can use to improve a visualisation for data communication. By the end of the lab you should be able to:

- understand how Matplotlib works under the hood
- draw line plots, bar charts and scatter plots
- fine tune plots produced in either Matplotlib or seaborn to improve data communication

Related Workshop: Visualisation

<https://opencourse.inf.ed.ac.uk/inf2-fds/course-materials/semester-1/week-4/task>

In this workshop, you'll try to apply the visualise principles and guidance introduced in this chapter to examples of Coursework 1 from 2020-21.

Chapter 6

Data collection and statistical relationships

Further reading (not examinable)

- The first few chapters of the online book *Causal Reasoning for the Brave and True* by Matheus Facure Alves is a very readable introduction to causal inference, for those interested.
- Interested students could look at *The Book of Why* (Pearl and Mackenzie, 2018) for a comprehensive and general introduction to the science of causal reasoning.

The chapter on [Data](#) covers the nature of data, under the assumption it has already been collected. This chapter considers how data is collected the first place, and how the way in which the data was collected affects the conclusions we can draw from the data. Often we collect data about multiple variables, and would like to understand if there are causal relationships between them. We'll introduce the correlation coefficient, which gives an indication if variables are associated. However, we must be mindful of the well-known saying "correlation is not causation", and will introduce the notions of confounding variables and causal inference.

6.1 Collecting data

Where does data come from? Data comes from many sources, for example:

- Official records (e.g. census records)
- Surveys, e.g. opinion polls, the weekly class rep survey
- Management systems, e.g. student and staff records, the timetable
- Healthcare records
- Wearable devices
- Computer logs
- Monitoring systems (e.g. electricity meters)
- Scientific observations
- Drugs trials
- Sports results
- Computer games stats

The nature of these sources varies: student and staff records are necessary for running a University and healthcare records are needed to run a hospital; they should be complete record, and need to be instantly available.

Census records help to plan future services, but the designers of the census questions know that they will not succeed in getting the whole population to respond, and it takes a number of years to quality assure

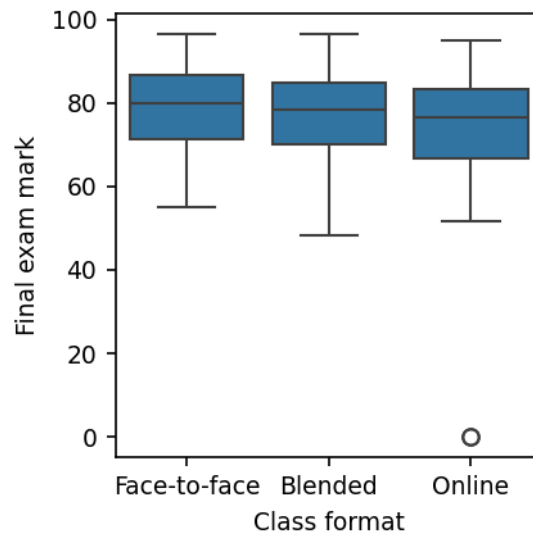


Figure 6.1: Outcomes of a randomised controlled trial in which 323 students were assigned to learn a class either face-to-face (Control), or online or blended (Treatments) (Alpert et al., 2016).

Class format	gpa	gender	sat_math_NEW	sat_verbal_NEW	falseexam
Blended	3.19	0.55	627.05	586.86	77.09
Face-to-face	3.19	0.63	633.17	589.13	78.55
Online	3.09	0.54	632.90	581.59	73.64

Table 6.1: Means of some variables in the three randomly assigned groups in which students were assigned to (Alpert et al., 2016). The outcome variable is “falseexam”; all the others are variables that are being controlled for by the random design.

the data and conclusions drawn from it before it is released. A survey such as an opinion poll asks questions of far fewer people than a census, and the result is wanted quickly.

Computer logs are needed to diagnose problems, but it is probably not so important to generate stats from them. Sports results and computer games stats are there mostly for pleasure, but sports and games enthusiasts will care deeply that they are correct.

Experiments In a scientific experiment the aim is to have a high degree of control (for example, we heat the solution of 80 millimolar sodium chloride solution and 5g of iron filings to a temperature of 80 celsius and then keep it there for 5 minutes). This means that we can repeat experiments and reproduce the results; we have confidence in our findings. In a subject like psychology, it is typically harder to have as much control in experiments as in chemistry or physics – humans are not identical like iron filings and sodium chloride solution – but with the aid of inferential statistics it is still possible to design and interpret experiments.

Studies In other disciplines, controlled experiments over a short time period in a lab setting are not possible. For example in medicine, it may take months or years for the effect of a new drug to become clear, during which time each patient will be subjected to a different set of influences. However, what are referred to as **studies** can be designed to minimise the effect of unavoidable variability.

The gold standard of study is the **randomised control trial**. Individuals participating in the trial are assigned at random to a **treatment group** and **control group**. In a medical context, the control group is given an existing drug (or a placebo) and the treatment group is given the drug to be tested. The **outcome** (e.g. “Did the participant get better?” or “How many headaches did the participant have?”) in both groups is measured. If the groups are big enough, there will be a similar distribution of characteristics such as age and sex in both groups, and so any differences in the outcome should be due to the treatment.

In other contexts it may be hard to devise a placebo, and the control group may be regarded as the

Table 6.2: Examples of study types and features of studies that make them prospective or retrospective. Based on <https://www.statsdirect.co.uk/help/basics/prospective.htm>.

Type	Definition	Example(s)
Cohort study, longitudinal study	A group (cohort) of subjects is followed over time and the outcome is measured during or at the end of study period.	Born in Bradford birth cohort study
Case-control study	Subjects with positive outcomes are identified and controls without the outcome are found, and then features of the subjects with the positive and negative outcomes are compared.	Investigation of breast cancer (Lane-Clayton, 1926 reanalysed by Press, 2010)
Randomised control trial	Individuals are assigned at random to a control group and a treatment group, and the outcomes of the two groups are compared.	Drugs trial, A/B testing
Feature	Definition	Example(s)
Prospective study	Outcomes are not known at the start of the study period, but features of the subjects of study are.	Randomised control trials. Longitudinal or cohort studies are often prospective.
Retrospective study	Outcomes are known at the start of the study, and data is gathered on features of the subjects with the different outcomes.	John Snow's investigation of cholera around the Broad St tap. Case-control studies are often retrospective.

status quo. For example, in a randomised trial of online learning (Alpert et al., 2016), students were assigned at random to Face-to-face, Online and Blended groups at the start of a course. The control condition is Face-to-face, and Online and Blended are treatments. The outcome was the mark on the final exam, shown in Figure 6.1. The boxplot shows that the median of the final mark is lower in the Online and Blended formats than in the Face-to-face format – we sometimes say that there is an **association** between class format (a categorical variable) and final exam mark (in this case a numeric outcome).

In Table 6.1 the means of some of the variables of the participants in the group are shown along with the outcome (“falseexam”). We can see the means of these non-treatment variables (such as GPA, the graduate point average, a measure of pre-university performance) are fairly similar between the groups, whereas the outcome is quite different, suggesting that the treatment is not just associated with the outcome, but potentially affecting (or causing) the outcome.

At the moment, we don't know if this effect is really significant, or might still be down to the chance of the composition of the groups. We'll return to this chapter in the Chapter on A/B testing. (In computer science, when testing the effect of two situations on a user, randomised control trials are called A/B tests.)

Table 6.2 lists various types of study, including randomised control trials. Studies may have prospective (“forward-looking”) features or retrospective (“backward-looking”) features. Prospective studies, such as randomised control trials, tend pose the question (e.g. “Is the drug effective?”) *before* deciding how to collect the data required to address the question. In retrospective studies, to answer the question (e.g. “Does smoking cause cancer?”), it may be necessary to find data that was not collected with the express purpose of addressing the question.

Observational data In contrast to experimental data and study data, **observational data** is data in which we have partial or no control over the individuals or things we are collecting data from. It is the data we have to hand, not the data we generated or collected to answer a question.

💡 Data story: John Snow and the Broad Street pump – the power of observational data

A classic example of the power of observational data is the case study of the 1854 Broad Street cholera outbreak. John Snow was a doctor, who hypothesised that the cholera was spread via the water supply, rather than through the air. In 1854, 127 people died in the Broad Street area of London over 3 days. The locations of homes of the people died was observational data. The homes were mostly clustered around the water pump in Broad Street where local people got drinking water from. This was strong

enough evidence to persuade the authorities to disable the pump by taking the handle off it. John Snow's map, indicating the homes of the people who died, is often regarded as a classic visualisation, showing how the people who died in the 1854 outbreak mostly lived close to the pump.

Beware of sample bias in big data

We need to be aware of how the data collection method might affect sample bias (see [Selection bias](#)), the bias where the sample differs from the population we are trying to make inferences about. As Tim Harford discusses in his blog, a common mistake is to assume that big data is representative of the population that we are trying to learn about ([Harford, 2014](#)). A classic example of this problem is the story of the 1936 US General election (see box), and Tim Harford also tells the story of how the big data approach of Google flu trends failed after a few years.

Data story: Gallup and the *Literary Review* in the 1936 US General Election

In 1936, the US *Literary Review* magazine sent out 10 million postcards to owners of cars and telephones, asking them if they would vote for Roosevelt (Democrat) or Landon (Republican) in the US presidential election. Of the 2.4 million of postcards returned 55% favoured Landon and 41% Roosevelt, with the remainder for a third candidate. In the same year the pollster George Gallup asked 3,000 people how they intended to vote in the election, but forecast that Roosevelt would win. In the event Roosevelt did win: the small poll had given the right answer. The reason was that George Gallup had tried to make the sample of voters he interviewed an unbiased sample of the population, including people who didn't own cars and telephones.

We need to be wary of making inferences from observational data, where there is much less control of how we found the data than in randomised control trials or surveys. However, if we know the attributes of the items sampled (e.g. sex, age and ethnicity), careful use of [Multiple regression](#) or [Logistic regression](#) can isolate the effects of each attribute on the outcome of interest.

6.2 Obtaining data already hosted online

Why obtain data already online? Although the ideal way of answering a data-scientific question is generally to collect the appropriate data, existing data may be relevant. There is now a wealth of data available online, much of it high-quality, and available under open data licenses. An advantage of obtaining data online is that it is quick – someone else has done the hard work of collecting the data. Disadvantages of obtaining data online are that the descriptions of the data and how it was collected may be unclear.

Ethical and legal considerations in online collection It is worth remembering that using online data involves a relationship between you and the data creator, and possibly the people who funded the data collection, so we need to think a little about law and ethics.

Licences for downloaded file If you're downloading data files, it's likely they will have been issued with a data license, which governs what you are allowed to do with the data, and how you can publish any work that you produce using that data. A few common types of license are:

- **Creative Commons:** These general-purpose licenses can come with conditions, such as that the creator must be credited (BY) or that only non-commercial uses are allowed (NC). The CC0 license is very permissive, meaning that the owner has no control of the data, not even requiring data users to attribute the data creator – but it is always good practice to attribute sources.
- **Open Data Commons:** Similar to the Creative Commons licences, but designed for data.
- **Open Government Licence:** A licence created by public bodies in the UK.

You should check the license or copyright statement, if there is one. [This section of *The Turing Way* \(The Turing Way Community, 2022\)](#) has more information about data licences.

Methods of obtaining online data There are a number of ways of obtaining data hosted online.

- Downloaded as structured files (e.g. tabular data, or data in JSON format) from websites, e.g. scientific, government or charity data repositories
- Via an API or database
- **web scraping** – the process of automating the process of extracting data from a web page.

These are listed in order of difficulty. It's generally better to start with the easier options first – there's no need to scrape data if it's already available in a structured format.

The first rule of web scraping is don't do it unless you have to. Data is increasingly available in structured formats such as CSV and JSON. Web scraping seems cool, but it's also time-consuming and fiddly. It pays to spend some time searching for files in structured formats, in which the data may well be in a cleaner format than on a web page. However, if you can't find the data you need in a structured format, but it is available on a web page, then web scraping can be helpful.

The FDS OpenCourse website has a [list of sources of data](#).

API conditions Sites with an API (e.g. [API-Football](#)) have conditions under which you can take the data – you should respect those conditions if they are not already enforced by the API.

Ethical web-scraping As with data files, the owner of the website you wish to scrape may have invested considerable time and money in creating their site. It's possible that by web scraping you could have an adverse effect on either their intellectual property or their web server.

The law around web scraping is not always clear ([Davies, 2020](#)) but you should always check the terms and conditions of the website before starting scraping. For example, the [Copyright Policy of the Financial Times website](#) says you cannot “Frame, harvest or scrape FT content or otherwise access FT content for similar purposes.” In contrast, the [Time Out website terms and conditions](#) are more permissive.

Beyond the legal restrictions for a specific website, we should also consider general ethical principles of web scraping. [Densmore \(2017\)](#) suggests a number of rules for those undertaking web scraping, including:

- using an API if available
- requesting data at a reasonable rate (not hoarding bandwidth)
- respect content rights, and not passing data off as one's own
- giving back to the data owner when possible, including attributing the data owner in any publication
- scrape only to create value from data, not to duplicate it

6.3 Correlation

In the topic on [Descriptive statistics](#), we considered summary statistics of one numeric variable, for example the weight of a squirrel. In data science and statistics we're often interested in statistical relationships between two or more variables. In this section we'll discuss the most basic statistical relationships: covariance and correlation.

Covariance We take some observational data that links the wages of individuals with the years of education they have experienced (Figure 6.2). We'll call the years of education of the i th individual x_i and the natural log wage y_i (the log transform is helpful here to make the distribution of wages less skewed). The **sample covariance** is defined:

$$s_{xy} = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) \quad (6.1)$$

We expect that individuals that have more education than average will also earn more than average, i.e. both $x_i - \bar{x}$ and $y_i - \bar{y}$ are positive. Thus, for these individuals $(x_i - \bar{x})(y_i - \bar{y})$ will be positive. We also expect that individuals that have fewer years of education than average will have lower wages than average. For these individuals both contributions $x_i - \bar{x}$ and $y_i - \bar{y}$ will be negative, but the product $(x_i - \bar{x})(y_i - \bar{y})$ will be positive. We should thus expect that the covariance of years of education and log wage is positive. If one quantity gets smaller as the other one gets bigger, then, by similar reasoning, the covariance is negative.

There are a few points to note about the covariance:

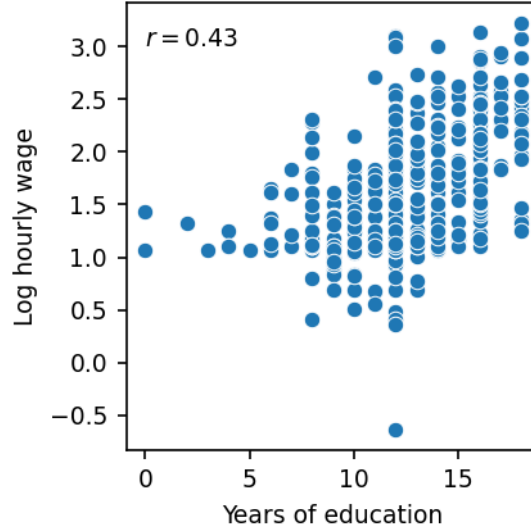


Figure 6.2: Log hourly wage in dollars versus years of education from the 1976 Current Population Survey collected and made available by [Wooldridge \(2020\)](#) and [\(Shea, 2024\)](#). The correlation coefficient of the log hourly wage and years of education is $r = 0.43$.

1. Its units are the product of the units of the two variables; in this example the units would be years (log quantities don't have units).
2. It depends on the scaling of the variables; suppose we measured the individual's education in months (instead of years) and we took log to the base 10 of their wage (which is 2.30 times smaller than the wage to the natural log). The covariance would be $12/2.30$ times larger, even though the relationship between the variables remains the same.

We don't want a measure of how strongly quantities are related to depend on the units they are measured in, so the second point is a problem.

Correlation coefficient The **sample correlation coefficient** (also known as Pearson's correlation coefficient¹) addresses this problem by dividing the covariance by the product of the standard deviations of the two quantities:

$$r = \frac{s_{xy}}{s_x s_y} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}} \quad (6.2)$$

It has a number of properties:

1. It ranges between -1 and 1 .
2. $r = 1$ corresponds to all points lying on straight line sloping upwards and $r = -1$ corresponds to all points lying on a straight line sloping downwards.
3. It has no units, i.e. it is dimensionless
4. It is independent of the units in which x and y are measured
5. r remains the same if we swap the labels of x and y
6. The correlation coefficient of two variables is equal to the covariance of standardised versions of those variables.

To convince yourself of the first property, imagine what would happen if all $y_i = cx_i$.

What's hiding in a correlation coefficient? With multivariate numeric data, a very helpful technique can be to look at the correlation coefficient of every pair of variables. These can be plotted as a heat map, quickly showing where there might be interesting relationships (i.e. correlation coefficients close to 1 or -1).

¹Karl Pearson (1857–1936) was a remarkable polymath – and a eugenicist. Although he invented many statistical concepts, the formula of the correlation coefficient was originally published before he was born.

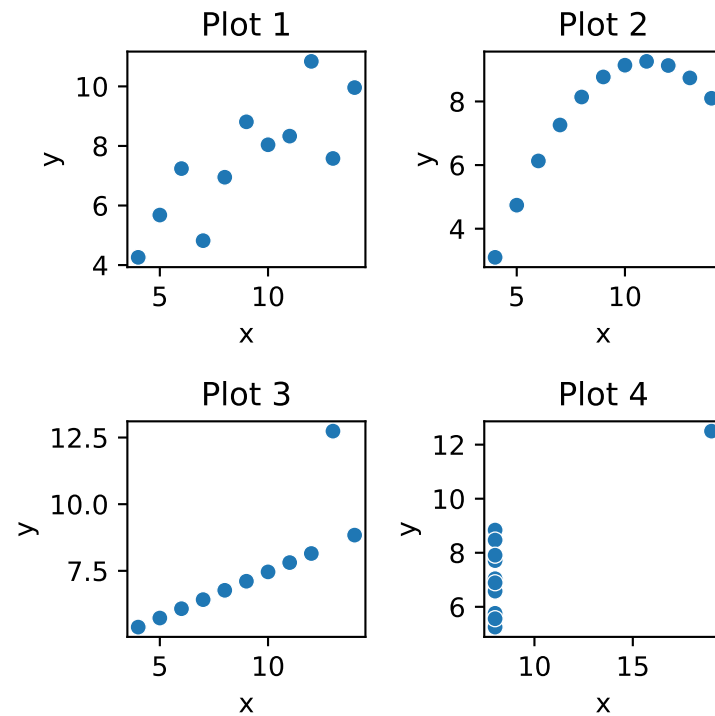


Figure 6.3: What are the correlation coefficients in these plots? “Anscombe’s quartet” ([Anscombe, 1973](#))

 Can you guess the correlation coefficient?

But before using correlation heat maps, we should take a look at Figure 6.3. Can you guess roughly what the correlation coefficient is in each plot? Before going to the next the page, try to think about how you would describe the data in each plot in words.

The answer is they all have the same correlation coefficient, $r = 0.82$. But the visualisation indicates that something quite different is happening in each one:

1. Basically a linear relation, with some noise around it
2. It looks like there is a very precise nonlinear (perhaps quadratic) relationship between the variables here
3. It looks like there would be a linear relationship with $r = 1$ if we took out the outlier point
4. Again there is something suspicious about the outlier point

i The correlation coefficient of two standardised variables

To see why the correlation coefficient of two variables is equal to the covariance of standardised versions of those variables, suppose that z is the standardised version of x and u is the standardised version of y , which we substitute into Equation 6.1 for covariance:

$$\begin{aligned}
 s_{uz} &= \frac{1}{n} \sum_{i=1}^n \frac{z_i - \bar{z}}{s_z} \frac{u_i - \bar{u}}{s_u} \\
 &= \frac{1}{n} \sum_{i=1}^n z_i u_i \quad \text{by properties 1 and 2} \\
 &= \frac{1}{n} \sum_{i=1}^n \left(\frac{x_i - \bar{x}}{s_x} \right) \left(\frac{y_i - \bar{y}}{s_y} \right) \tag{6.3} \\
 &= \frac{1}{s_x s_y} \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) \\
 &= \frac{s_{xy}}{s_x s_y} \quad \text{which is the definition of the correlation coefficient}
 \end{aligned}$$

6.4 Limitations of statistical relationships

Correlation and causation You've probably heard the phrase "correlation is not causation". Sometimes this is phrased as "association is not causation". Let's define what we mean by correlation and causation:

Two variables are **correlated** or **associated** when knowing the value of one gives you information about the value of the other variable. For example:

1. Wearing seat belts is correlated with fewer injuries in car accidents.
2. Older people are more likely to have credit approved by a bank.
3. Ice cream sales are correlated with incidence of sunburn.
4. Taking an online class is associated with lower final exam performance in a randomised controlled trial.
5. Wages are correlated with years in education.

Two variables are **causally related** when changing the value of one of the variables affects the value of the other variable. To return to the examples:

1. Multiple lines of evidence show that wearing seat belts causes fewer injuries in car accidents
2. We might ask if it is really age that causes a higher chance of credit approval, or some other factor associated with age, for example higher income.
3. It is fairly clear that eating ice cream does not cause sunburn, and vice versa. Sunshine is a common cause of both.
4. Since the data comes from a randomised controlled trial, there is fairly strong evidence that the form of online study in the trial helped students less than the form of face-to-face study.
5. We might ask if more years in education is causing wages to change, or if some other variable associated with education (for example family background) is also contributing to higher wages.

We cannot determine whether a relation is causal just from the statistics – we also need to bring in reasoning about how the data originated.

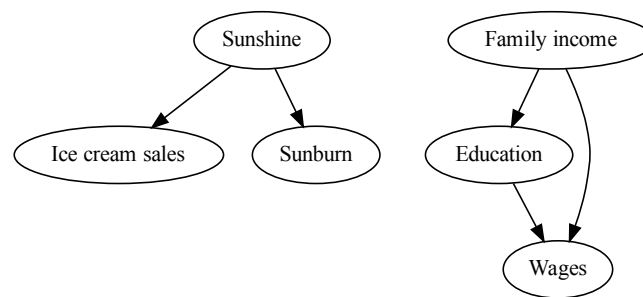


Figure 6.4: Examples of causal graphs.

Graphical causal diagrams A very helpful way of reasoning about causality are causal diagrams. Examples of causal graphs corresponding to some of the examples in the previous lists are shown in Figure 6.4. Here, each node on in the graph is a variable, and the arrows represent our beliefs about which variables cause other variables to change. In the ice cream sales and sunburn example, we can see that sunshine may cause both sunburn and ice cream consumption. Because sunshine is driving both variables, there will thus be a correlation between ice cream sales and sunburn, but the diagram makes clear that we do not believe there is a causal connection between the two.

Confounding variables A **confounding variable** is a variable that influences a treatment variable and an outcome variable, meaning that it is possible for the treatment and outcome to be correlated, even if there may be no causal relationship between them. In the wages example, family income could be a confounding variable, since it could affect education (which could affect wages) or it could affect wages directly (perhaps because higher family income helps with making connections, and hence getting a better paid job).

Spurious correlations Sometimes there appear to be remarkably high correlations between two apparently unrelated variables: for example according to the [Spurious correlations website](#), the annual number of Disney movies released and annual number of motor vehicle thefts each year in the US has a correlation of $r = 0.864$. In this situation it's difficult to think of a common cause as in the case of the ice cream sales and sunburn, and it may well be chance that the time courses of these two statistics align. We give the name **spurious correlations** to correlations that arise by chance between two unrelated variables.

6.5 Causal Reasoning in Graphs

In this section we will show how we can represent and reason about causal relationships using graphs. It turns out that the graph structure makes it easy to identify and to reason about complex effects of variables on each other. The past twenty years have seen a multitude of research on causal reasoning. Interested students could look at *The Book of Why* (Pearl and Mackenzie, 2018) for a comprehensive and general introduction to the science of causal reasoning.

Suppose we wish to determine the effects of a new drug on survival rates in a hospital. We can use a graph shown in Figure 6.5 to represent the experiment design. A node in the graph represents a boolean variable. The node *Medicine* equals True if a given patient has received the drug, and False if they did not receive it. Similarly, the node *Severeness* is True if the patient is already severely ill and False otherwise; the node *Survival* is True if the patient recovers and False otherwise. An edge in the graph between two nodes represents a **causal link** – the source node can affect the value of the destination node. In our example, the edge *Medicine* → *Survived* states a positive relationship, that receiving the drug increases the likelihood of recovery. Similarly, the edge *Severeness* → *Medicine* states that being severely ill also increases the likelihood of receiving the drug. The edge *Severeness* → *Survived* states a negative relationship, that being severely ill decreases the likelihood of recovery.

We note that the relationships encoded in the graphs are probabilistic. In our example, this means that receiving the drug will *probably* lead to recovery, but there is also a small chance that the patient won't

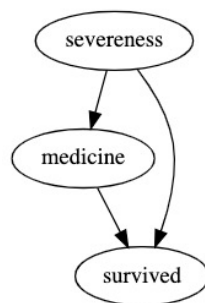


Figure 6.5: Causal graph for drug experiment. From *Causal Reasoning for the Brave and True*

recover. The degree of uncertainty expressed in the graph structure can be controlled by placing probability distributions in the nodes, but this is beyond the scope of this week's lecture, so we won't go into it.

Paths can involve intermediate nodes: the path $Severity \rightarrow Medicine \rightarrow Survived$ means that a patient that was severely ill was more likely to receive the drug, which in turn makes it more likely that the patient survived. Note that reasoning paths in the graph can be non directed. The path from $Survived \rightarrow Medicine$ demonstrates evidential reasoning: a patient who survived was more likely to be given the drug.

How can causal graphs help us with experiment design? We can use the graph to directly observe the relationships between treatments, outcomes and confounding variables. In our example, the node *Medicine* represents the treatment and the node *Severity* represents the outcome. The node *Severity* in the graph is a possible cause of both the treatment and the outcome.

There exists an active path in the graph from the treatment (*Medicine*) to the outcome (*Survived*), which is what we want in an experiment. However, there is also a (non-directed) path $Medicine \rightarrow Severity \rightarrow Survived$ that passes through *Severity*. In this case, *Severity* is a confounding variable, because it creates an active 'backdoor' path from *Medicine* to *Survived* that represents alternative explanations. Indeed, if only highly sick patients get the drug, it might even look like giving the medication decreases the patients' health. That is because the effect of the severity is getting mixed up with the drug's effect.

How can we close backdoor paths? By **conditioning** on the confounding variables in the path. In our example, we would analyse the effect of the drug on survival separately for severe and non-severe cases. Because we are able to observe the value of *Severity*, we can isolate it in the analysis and essentially block the back door path.

In this simple example there was a single confounding variables, but there could be potentially other confounding variables. For example *Allergic* can be a parent node of *Medicine*, because allergic patients cannot get the drug. Even though there does not exist an edge $Allergic \rightarrow Survived$, this node can still affect both the treatment and outcome and would be part of a backdoor path (or paths) that would need to be conditioned on.

Note that in some cases, backdoor paths contain variables that cannot be directly observed. For example, suppose we wish to measure the effect of Education on Wages. We know that Education is a confounding variable because it affects both education and wages, but we can't directly observe it. However, we *can* observe students' college preparation scores (their SAT scores), which is also confounding variable and is a noisy indicator of intelligence. In this case, we can condition on the value of the SAT score to block the backdoor path $Education \rightarrow SAT \rightarrow Intelligence \rightarrow Wages$. Effectively this minimises the effect of the confounding variable.

Part II

Linear models

Chapter 7

Linear Regression

7.1 Regression as prediction

Regression is related to the correlation coefficient, but is subtly different. In regression, we are still dealing with two numeric variables x and y , but we are trying to predict what value of y will be found at a particular x , not just give a measure of how they are correlated. Often we have good reason to believe that there is a causal relation between the variables x and y such that x causes y or y depends on x .

The variable y can be referred to as the **response variable** (or “response” for short) and the variable x can be referred to as the **predictor variable** (or **predictor** for short). We will use this terminology¹ but it is worth noting that the response variable is also referred to as a **dependent variable**, **target variable**, the **outcome** or **endogenous variable** and the predictor variables are also referred to as **features**, **explanatory variables**, **independent variables**, **covariates**, **regressors**, **exogenous variables**. The terminology used tend to vary on the discipline. In particular, in Machine Learning, the pair “target” and “features” or “covariate” is used; we have used “feature” when introducing Machine Learning in [Supervised learning: Classification with Nearest neighbours](#).

A standard dataset used to illustrate regression involves another mammal, *Homo sapiens*. The data was collected by the inventor of regression, Francis Galton², who surveyed the heights of grown-up children and their parents. He expected that parents who were taller than average would have children who were taller than average. Galton’s results are shown in Figure 7.1. The data is simplified. Galton recorded four variables: the height of the mother, the height of the father, the gender of the child and the height of the child. We have reduced the number of variables to three by taking the mean height of each child’s parents to give a variable we call the midparental height.

Suppose we want to predict the height a daughter will grow to given that we know the heights of her parents, and thus her midparental height. We have a reasonable amount of data, so we could just take the mean height of all daughters with a similar midparental height. To be more specific, if we knew the midparental height were x , we could predict the future height of the daughter to be the mean height of all daughters with midparental heights in the range $[x - 0.5, x + 0.5]$.

Figure 7.2 shows the results of this method of prediction. It’s generally true that the taller the parents, the taller the daughters will be, though the relationship is not quite monotonic – it fluctuates a bit at lower midparental heights. It seems reasonable to suppose that, if there were more data, the relationship would be linear. In the next section we’ll see what happens if we assume the relationship is linear.

7.2 Linear regression

The linear regression model We’ll now assume that y depends linearly on the predictor variable x :

$$y = \beta_0 + \beta_1 x \quad (7.1)$$

The variables β_0 and β_1 are called **parameters** or **regression coefficients** and are the intercept and slope of the line respectively. Think of them as dials that we can turn to produce any straight line we want to. Our

¹Up until 2022–23, we used the terms “dependent variable” and “independent variable” in FDS. However, the term “independent variable” is confusing, since in multiple regression, the predictor variables are often not statistically independent.

²Francis Galton (1822–1911) was a remarkable polymath – and a eugenicist. Adam Rutherford’s book *Control: the Dark History and Troubling Present of Eugenics* (Rutherford, 2022) gives a detailed account of Galton’s views and those of his academic offspring and fellow statisticians and eugenicists Karl Pearson and Ronald Fisher.

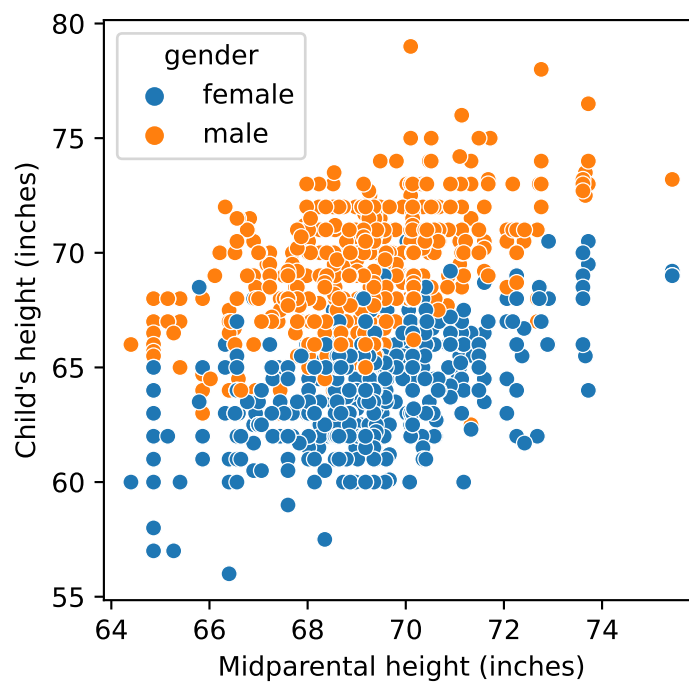


Figure 7.1: Height of daughters and sons plotted against the mean height of their two parents.

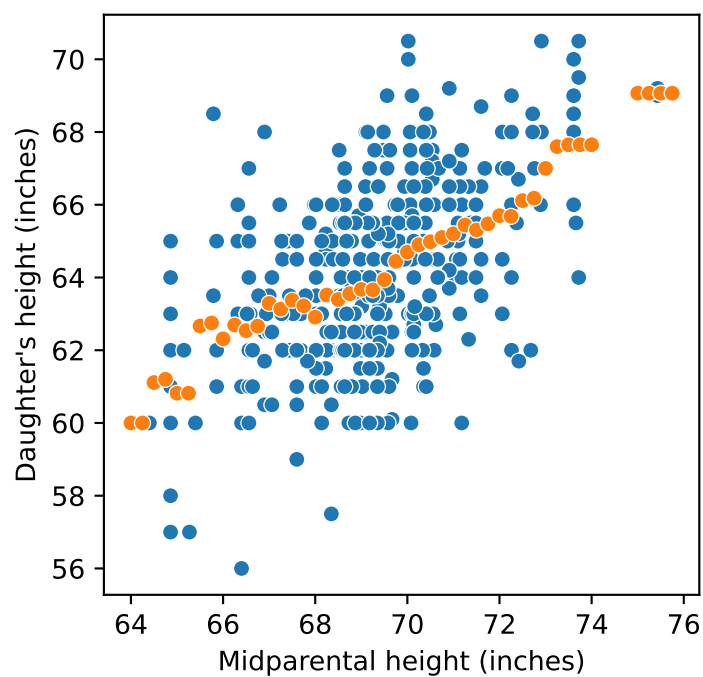


Figure 7.2: Heights of daughters plotted against their midparental heights (blue dots) and predictions (orange dots). The method described in the text doesn't work where there are gaps in the x values of greater than 1 inch.

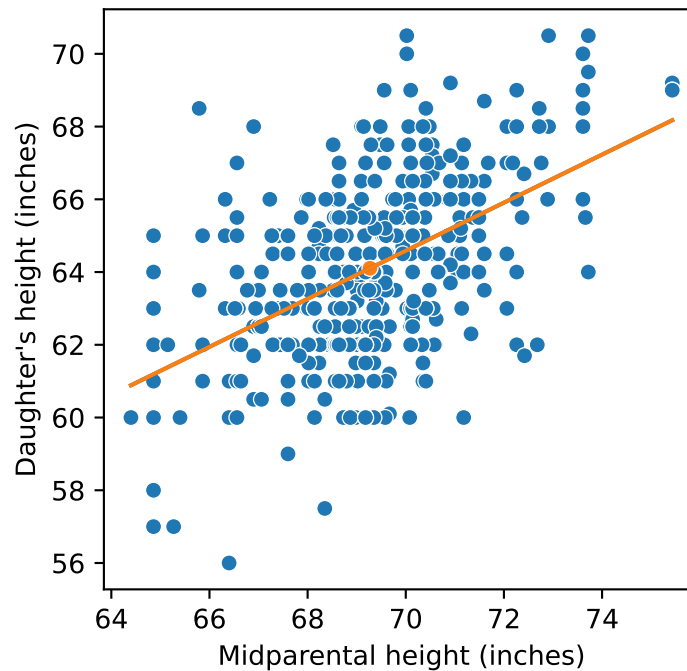


Figure 7.3: Linear regression of daughter's height on their midparental height. The linear regression line is shown in orange and the location of the means of the daughter's height and midparental height is indicated by the orange dot.

aim is to tune the values of β_0 and β_1 so that the line lies close to the data, and will therefore be a good predictor. We refer to Equation 7.1 as the **linear regression model**.

Models The word “model” recurs throughout data science and statistics, and indeed, science in general. A general definition is that a **model** is an abstraction of reality that captures aspects that are important for a given task and omits the rest.

For example, a LEGO® model of a tractor looks like a tractor but doesn't include all the details of a real tractor (Figure 7.4). Similarly, the linear regression model looks like the data, but doesn't contain all the details. Figure 7.3 shows the linear regression model of the data that we are working towards in the next few pages.

Note that linear regression is called a **parametric model**, since it contains explicit parameters (the regression coefficients). In contrast, the method shown in Figure 7.2 is a **non-parametric method**, since the



Figure 7.4: Left: LEGO Expert Builder set 952. (David Sterratt, CC-BY-4.0) Right: Massey Ferguson 135 Tractor (Lyle Buist, reproduced with permission).

prediction does not depend explicitly on any parameters.

i Types of model

The LEGO® tractor is an example of a physical or mechanical model. If we're using equations to connect how gravity and Newton's laws lead to the motion of a pendulum, we'd be using a mathematical model; if we were simulating a city in a computer simulation game, we'd be using a computational model. Biologists refer to animal models, i.e. investigate biological processes such as diseases in an animal, with the aim of understanding these biological process in humans.

The linear regression model is an example of a **statistical model**. We use statistical models to investigate relationships in data and to make predictions. Statistical models can help us to explain the data, but they do not provide a mechanistic explanation of the data. For example, fitting a sine wave to the motion of a pendulum would be a statistical model; in a mathematical model of the pendulum the sine wave would emerge from solving equations describing gravity, Newton's laws, and the weight and length of the pendulum. In these lecture notes the word "model" will generally mean "statistical model", though when we are referring specifically to probabilities (e.g. in the chapter on [Randomness, sampling and simulation](#)), we refer to "probabilistic models".

The principle of least squares To tune the parameters, we need to have a measure of how close the line is to the data. One way of defining "close" is how the mathematicians Gauss³ and Legendre did around 1800: the sum of the squared deviations between the predicted and actual values of y for every value of x . The closeness to the line is a function of β_0 and β_1 :

$$f(\beta_0, \beta_1) = \sum_{i=1}^n (y_i - (\beta_0 + \beta_1 x_i))^2 \quad (7.2)$$

We call this function a **error function** or **loss function**. We are searching for the values of β_0 and β_1 that minimise the error function. We denote these values $\hat{\beta}_0$ and $\hat{\beta}_1$, pronounced "beta 0 hat" and "beta 1 hat" respectively. In statistics, the hat denotes the best estimate of a quantity. Minimising this error function is to the principle of least squares, which states that the best fit is obtained by minimising the sum of the squared deviations between the predicted and actual values of y .

Applying the principle of least squares We can imagine visualising the error function in 3D, by plotting the β_0 and β_1 on the horizontal axes, and $f(\beta_0, \beta_1)$ on the vertical axis. This function will look like a landscape, a valley with a lowest point, i.e. the point we're trying to find.

This is an example of an **optimisation** problem: how to find the arguments of a function that maximise (or minimise) that function. We could try to find the values of $\hat{\beta}_0$ and $\hat{\beta}_1$ that minimise the error function by using a numerical optimisation function in a library, such as the [SciPy optimize function](#). This function allows the user to choose one of a number of optimisation algorithms, which it is beyond the scope of this course describe.

Numerical methods should work (with more or less difficulty) on all optimisation problems. However, in the case of the linear regression model, we can find the optimal values $\hat{\beta}_0$ and $\hat{\beta}_1$ analytically, as shown in the box.

i Analytical derivation of $\hat{\beta}_0$ and $\hat{\beta}_1$

The derivation in this box is not examinable. However, it may be of interest if you want to understand where the expressions for the regression coefficients come from. In order to understand the derivation fully you need to know what partial derivatives are. In short, if we have a function of more than one variable, for example $f(\beta_0, \beta_1)$, when we find the partial derivative with respect to β_0 , we're just imagining that β_0 is the only *variable*, and that β_1 is a constant. Finding the partial derivative with respect to β_0 is then just like finding a derivative with respect to β_0 and pretending that β_1 is a constant. We then do the same thing, but with β_1 being the variable, and β_0 being the constant. We denote the partial derivatives using the symbol '∂' instead of 'd', as in Equation 7.3 below.

We can find values of $\hat{\beta}_0$ and $\hat{\beta}_1$ by setting the partial derivatives of the error function f with

³Gauss was a polymath – and not a eugenicist.

respect to β_0 and β_1 equal to 0:

$$\begin{aligned}\frac{\partial f}{\partial \beta_0} &= \sum_{i=1}^n (-2)(y_i - \beta_0 - \beta_1 x_i) = 0 \\ \frac{\partial f}{\partial \beta_1} &= \sum_{i=1}^n (-2x_i)(y_i - \beta_0 - \beta_1 x_i) = 0\end{aligned}\tag{7.3}$$

We can rearrange these formulae so that it is obvious that they are a pair of simultaneous equations in β_0 and β_1 :

$$\begin{aligned}n\beta_0 + \left(\sum x_i\right)\beta_1 &= \sum y_i \\ \left(\sum x_i\right)\beta_0 + \left(\sum x_i^2\right)\beta_1 &= \sum x_i y_i\end{aligned}\tag{7.4}$$

These equations are called the **normal equations**. Notice that we've used the abbreviated notation for summation here.

The sums over x_i and y_i are related to their sample means: $\sum x_i = n\bar{x}$ and $\sum y_i = n\bar{y}$. These relationships allow us to simplify the normal equations:

$$\begin{aligned}n\beta_0 + n\bar{x}\beta_1 &= n\bar{y} \\ n\bar{x}\beta_0 + \left(\sum x_i^2\right)\beta_1 &= \sum x_i y_i\end{aligned}\tag{7.5}$$

We can then eliminate β_0 to give:

$$\hat{\beta}_1 = \frac{\sum x_i y_i - n\bar{x}\bar{y}}{\sum x_i^2 - n\bar{x}^2} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2}\tag{7.6}$$

The second part of this equation follows from the identity $\sum_{i=1}^n (x_i - \bar{x})^2 = \sum_{i=1}^n x_i^2 - n\bar{x}^2$ (see Equation 3.10 in the [Descriptive statistics](#) topic for proof of the identity). It follows from the first of Equation 7.5 that:

$$\hat{\beta}_0 = \bar{y} - \hat{\beta}_1 \bar{x}\tag{7.7}$$

Properties of the linear regression line Don't worry if you don't follow the derivation above – let's stand back and look at some properties of the least squares linear regression line.

1. The line passes through the point $(\bar{x}, \bar{y})$. We can see this by substituting the value of $\hat{\beta}_0$ into Equation 7.1, which leads to:

$$y = \hat{\beta}_0 + \hat{\beta}_1 x = \bar{y} - \hat{\beta}_1 \bar{x} + \hat{\beta}_1 x = \bar{y} + \hat{\beta}_1 (x - \bar{x})\tag{7.8}$$

When $x = \bar{x}$ it follows that $y = \bar{y}$. In Figure 7.3 you can indeed see the regression line passing through the mean of both variables.

2. Equation 7.6 for the gradient of the line $\hat{\beta}_1$ looks similar to the equation for the correlation coefficient r (Equation 6.2 in the [Data collection and statistical relationships](#) topic) but is different in one respect. The numerator is the same, but the denominator contains only the sum of squared deviations of x rather than the product of the square roots of the sum of squared deviations of x and y . We can in fact relate $\hat{\beta}_1$ and r via the standard deviations of x and y :

$$\hat{\beta}_1 = \frac{s_y}{s_x} r\tag{7.9}$$

3. If we plug this expression for $\hat{\beta}_1$ into the fitted model (Equation 7.8), we get:

$$y = \bar{y} + \frac{s_y}{s_x} r (x - \bar{x}) = \bar{y} + s_y r \left(\frac{x - \bar{x}}{s_x} \right)\tag{7.10}$$

which rearranges to

$$\frac{y - \bar{y}}{s_y} = r \left(\frac{x - \bar{x}}{s_x} \right)\tag{7.11}$$

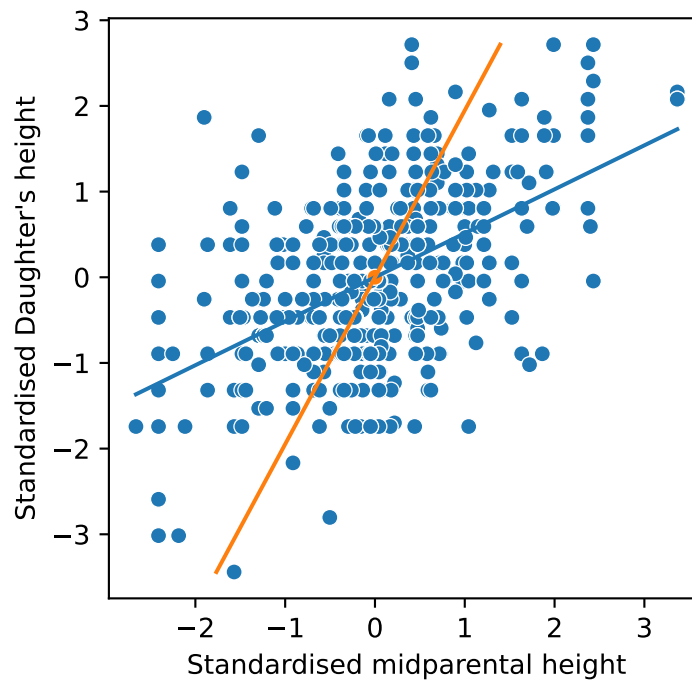


Figure 7.5: Linear regressions of: (i) standardised daughter height on standardised midparental height (orange line) and (ii) standardised midparental height on daughter height (green). For both regression lines $\hat{\beta}_1 = r = 0.51$. Note that children of parents whose height is 3 standard deviations higher than the mean are predicted to be less than 2 standard deviations above the mean – this is what the famous phrase “regression to the mean” is referring to.

Here the fractional terms on the left- and right-hand sides are the standardised versions of the variables x and y , which we also refer to as their z -scores (see section on [Standardised variables](#)). This shows that we can think of making predictions from regression in four steps:

- (a) compute the z -score for x (i.e. we standardise x)
- (b) multiply the z -score by the correlation coefficient
- (c) multiply the resulting value by the scale of y (via s_y), and
- (d) add the mean of y .

This is a simple example of moving into a “normalised space” (the z -score), doing the prediction there, and then going to the target space by multiplying-in the y -scale (i.e. the standard deviation) and adding-in the y -location (i.e. the mean). This view also shows that “learning”, i.e. estimating r , happens in the normalised space, since we correlate the z -score of x with the z -score of y .

Insights from regression with standardised variables We can gain some insights into regression by standardising x and y (Figure 7.5):

1. Since the mean of standardised variables is zero, the regression line passes through $(0, 0)$.
2. The gradient of the regression line of y on x is $\hat{\beta}_1 = r$, from Equation 7.9. The intercept of the regression line $\hat{\beta}_0 = 0$, from Equation 7.7. As the relationship between the variables gets weaker (r gets smaller), the gradient of the regression line decreases.
3. The predicted standardised y is always closer to the mean than the standardised x . In other words, parents who are much taller than average are likely to have children who are taller than average, but not by as much as the parents: the height of these children is likely to be closer to the mean height of children than the tall parents’ height was to the mean parental height. The same is true of very short parents and their children. Galton characterised this observation as “regression to mediocrity”, or regression to the mean.

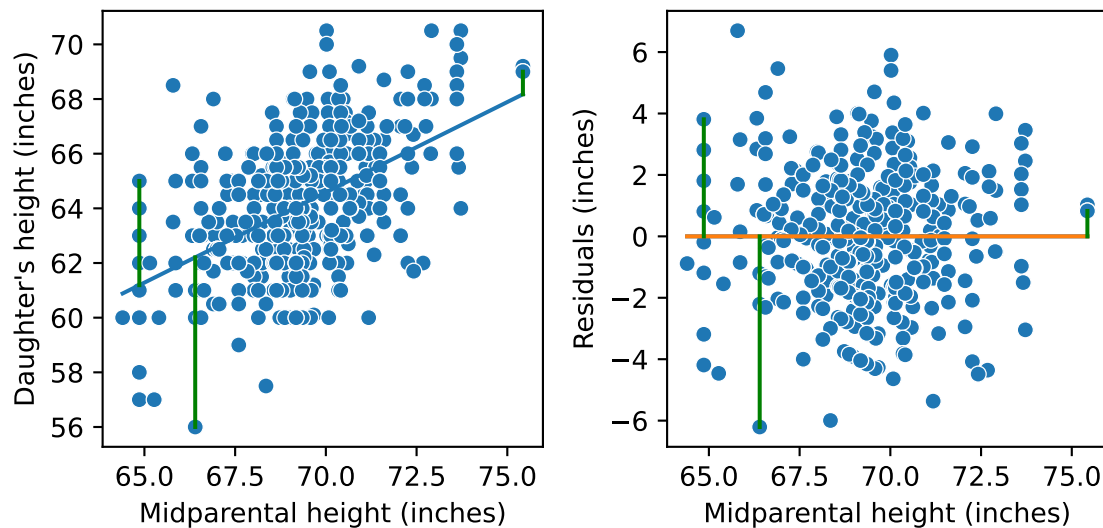


Figure 7.6: Data points and regression line (left) and residual plot (right) for regression of daughter's height on midparental height. The residuals for the same three data points are shown in both plots.

4. If we try to predict x from y , i.e. we flip the variables, we find the same gradient $\hat{\beta}_1$. However, by plotting on the same set of axes, we see that when $-1 < r < 1$ the regression lines of y on x (orange) and x on y (green) are different. If there is perfect correlation ($r = 1$ or $r = -1$), the two lines overlap.

7.3 Visual diagnostics and transformations

Residuals Once we have determined the parameters $\hat{\beta}_0$ and $\hat{\beta}_1$ from the variables x and y , we can compute the **predicted values** $\hat{y}_1, \hat{y}_2, \dots, \hat{y}_n$ of y for each of the x values by substituting x into the estimated regression line:

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_i \quad (7.12)$$

The differences $y_i - \hat{y}_i$ between each pair of predicted and actual values of the response variable are called **residuals**. They can be visualised as the vertical deviations of each data point from the regression line (Figure 7.6, left). Plotting the residuals on their own (Figure 7.6, right) can indicate whether the linear regression model is an appropriate one for the dataset in question.

It is worth noting that residuals from a linear regression always have:

1. Zero mean
2. Zero correlation with predictor variable x
3. Zero correlation with the predicted values $\hat{y}$

These are all true no matter what the data looks like, just like the mean of deviations from the mean is zero.

Nonlinearity The linear regression model is appropriate when the underlying data is linear, but deviations from linearity might not be very apparent when the regression line is plotted with the data. The residual plot makes this more obvious. An example which demonstrates this very clearly is a plot of world population versus year, for the years 1940–2000 (Figure 7.7). The regression line fit looks alright, but when we look at the residuals, we see more clearly that there seems to be something systematically nonlinear going on: the estimate is too low at the start and end of the sequence and the residuals fall and rise again smoothly.

Transforming the data Just because it looks as though data is nonlinear doesn't mean we have to give up on linear regression. We might suppose that the world population grows exponentially (if we looked at a longer time series we would definitely get that idea). To turn an exponential curve back into a linear one, we can transform the data by taking log to the base 10 of the population (Figure 7.8). The residuals now appear to have a less systematic relationship with the predictor variable, and if we turn the residuals of the

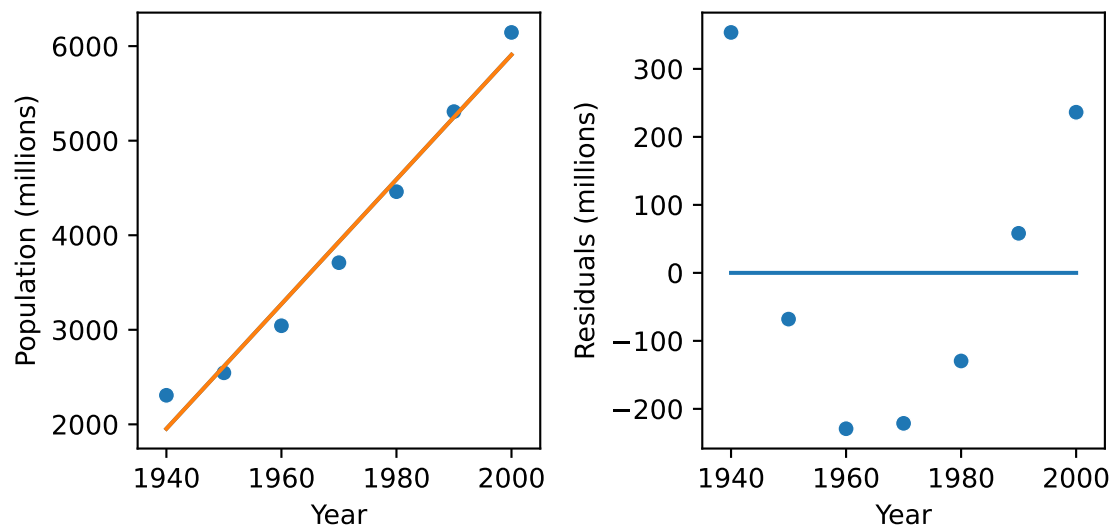


Figure 7.7: World population 1940–2000 (Klein Goldewijk et al., 2017)
Source: HYDE 3.2 database <https://dataportaal.pbl.nl/downloads/HYDE>

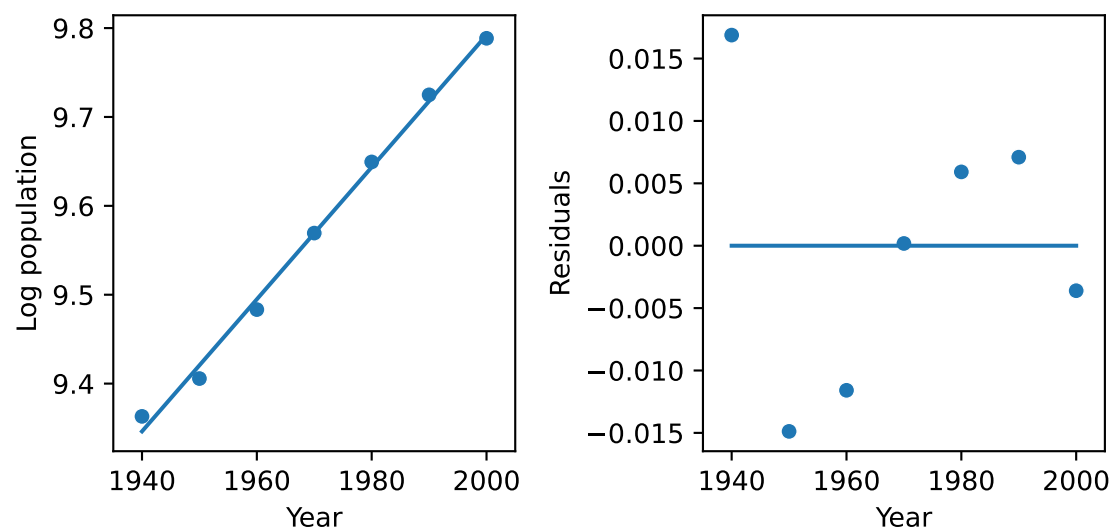


Figure 7.8: Log world population 1940–2000.

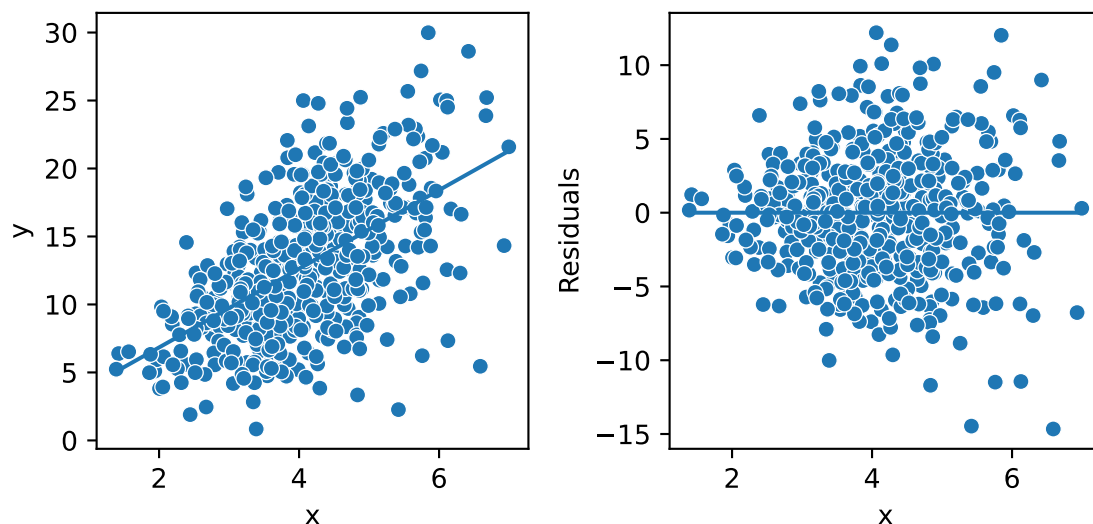


Figure 7.9: A synthetic dataset that exhibits heteroscedasticity.

log population back into raw numbers of people, we find the largest of these residuals is of the order of 70 million, rather than 300 million for the linear fit (Figure 7.7).

We can also transform the predictor variable or both the predictor and response variables. Whether it is appropriate to transform depends on our understanding of the data and the potential underlying processes.

For example, we might expect that the weight of a squirrel is proportional to the cube of its length (assuming that its height and width are proportional to length). This would suggest regressing the cube root of the weight on the length.

Heteroscedasticity A quick look at the variance of the residuals (Figure 7.6) suggests that the variance of residuals doesn't change as a function of the predictor variable. But suppose we had some data that looks like Figure 7.9. The variance here clearly increases as the predictor variable increases, and we say that the residuals exhibit **heteroscedasticity**. The word heteroscedasticity comes from the ancient Greek “hetero” (different) and “scedastic” (spread) – in other words the variance of the residuals changes as we go along the x -axis.

In the chapter on [Statistical inference and regression](#), we'll look at linear regression from a probabilistic perspective, and see that datasets exhibiting heteroscedasticity violate the assumptions that we're using in least-squares regression, namely that the variance of the residuals is independent of x , which we call **homoscedasticity**. (The Greek word *homos* means “same”.)

7.4 Numerical diagnostics

There are a number of ways of measuring how good a regression fit is.

(Root) Mean Squared Error We have already seen the sum of the squared deviations; this is what we minimised with respect to the parameters in order to fit the regression line. However, this quantity scales with the number of points, and we would prefer something that gives an indication of the typical size of a residual. The **mean squared error** (often written MSE) is defined by:

$$\text{MSE} = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (7.13)$$

and the **root mean squared error** (RMSE) is just the square root of this. The RMSE gives an indication of how far data points deviate from predictions in absolute terms.

Coefficient of determination Another way of looking at the performance of the regression is to consider how much of the variance in y we can *explain* with x . If we have a perfectly linear relationship (correlation coefficient $r = 1$ or -1), then if we know x , we can predict y precisely. Thus, we have *explained* all the variance s_y^2 of y . When the relationship has a correlation of magnitude lower than 1, if we know x , we can predict y better than if we didn't know x , but we can also expect that the actual value of y will be some way from our prediction, and we say that there is unexplained variance.

The **coefficient of determination** metric quantifies how much variance is explained. It is defined in terms of two quantities:

1. The **total sum of squares** (SST), which we define as the sum of squared deviations from the mean of y :

$$\text{SST} = S_{yy} = \sum (y_i - \bar{y})^2 = (n - 1)s_y^2 \quad (7.14)$$

It is a measure of the total variance in y before we know anything about x .

2. The **sum of squared errors** (SSE):

$$\text{SSE} = \sum (y_i - \hat{y}_i)^2 \quad (7.15)$$

The SSE is n times the MSE defined in Equation 7.13.

The coefficient of determination is then defined:

$$R^2 = \frac{\text{SST} - \text{SSE}}{\text{SST}} = 1 - \frac{\text{SSE}}{\text{SST}} \quad (7.16)$$

R^2 is measure of “goodness of fit”. $R^2 = 1$ indicates that the model predicts the data perfectly, whereas a value of 0 indicates no predictive value. In physical and biological sciences we might expect R^2 to be of the order of 0.8, but this can be lower in other disciplines such as social sciences. We can also think of R^2 as 1 minus the MSE normalised by the variance of y . If the mean squared error is a high fraction of the variance, the R^2 will be low.

The notation suggests the coefficient of determination is related to the correlation coefficient r . For a linear regression line we can indeed prove that $R^2 = r^2$. The disadvantage of using R^2 for linear regression is that it doesn't indicate if the correlation is positive or negative. However, R^2 is more versatile: it can measure how well a nonlinear model fits the data. For a nonlinear model, R^2 is not generally equal to the squared correlation coefficient. The R^2 definition also generalises to multiple linear regression, which we will come to in the next section.

Related Python Lab: Linear models

<https://github.com/Inf2-FDS/FDS-S1-06-linear-models>

In this lab you will learn to apply linear regression to datasets and evaluate your model. By the end of the lab you will be able to

- Apply linear regression using the statsmodels package
- Interpret some of the visual and numerical diagnostics from the package
- Transform variables to produce a better model of the data
- Apply multiple regression to some data

Chapter 8

Multiple regression

8.1 The principle of multiple regression

Multiple predictor variables Often we want to investigate how a variable depends on more than one predictor variable. For example, we might expect a student's grade in a calculus course could be predicted by their grades in four previous assessments. In this case the response variable is the grade, and we have four predictor variables. The process of predicting a response variable from multiple predictor variables is called **multiple regression**.

Dealing with categorical variables The predictor variables may also be categorical. In the example of predicting a child's height from the heights of their parents, introduced in the chapter on [Linear Regression](#), there were three variables in the data: midparental height, height of child and gender of child. Although gender is a categorical variable, we can convert it to a numeric variable by encoding "Daughter" as 0 and "Son" as 1. As described in the section on [Tabular data and variables](#), this new numeric variable is called a dummy variable or indicator variable, and categorical variables with more than two values (for example colours) can also be encoded using multiple indicator variables. We can treat indicator variables mathematically in the same way that we would treat variables that are naturally numeric.

The multiple regression model To predict the child's height given their midparental height and their gender, we could try to use a non-parametric method, as we did when thinking about regression as prediction. However, we choose here to focus on extending the linear regression model introduced in the last chapter, i.e. a parametric model, in which the parameters were the coefficients β_0 and β_1 . Suppose we have two predictor variables, $x^{(1)}$ and $x^{(2)}$ and one predictor variable y . (We're using this rather cumbersome notation for the predictor variables so that we don't get confused with the notation for instances of variables.) The multiple regression model is then expressed:

$$y = \beta_0 + \beta_1 x^{(1)} + \beta_2 x^{(2)} \quad (8.1)$$

Geometrically, this is a 2D plane in 3D space, with β_1 being the gradient of a cut through the plane with constant $x^{(2)}$ and β_2 being the gradient of a cut through the plane with constant $x^{(1)}$.

Principle of fitting the multiple regression model The principle of fitting the multiple regression model is exactly the same as the principle of fitting the linear regression model, just with more variables. We use the least squares principle, but this time we have to adjust three coefficients, β_0 , β_1 and β_2 , to manoeuvre the plane around so that we minimise the sum of the squared distances between the predicted and actual values of y over all the data points.

In maths, we'll modify the function f (as defined in the chapter on [Linear Regression](#)) that we're minimising to be:

$$f(\beta_0, \beta_1, \beta_2) = \sum_{i=1}^n (y_i - (\beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2}))^2 \quad (8.2)$$

Note that x_{ij} means the i th instance of the j th variable. We could have used the more cumbersome notation $x_i^{(j)}$, but x_{ij} is simpler, and also makes more sense when we come to minimise the function analytically.

Table 8.1: MSE, RMSE and coefficient of determination for various simple and multiple regression models applied to the heights data.

Model	MSE	RMSE	R^2
Child's height on midparental height and gender code (mult. regression)	4.69	2.16	0.63
Daughter's height on midparental height (regression)	4.08	2.02	0.26
Son's height on midparental height (regression)	5.27	2.30	0.23
Child's height on midparental height (regression)	11.48	3.39	0.10

As with linear regression, once we have gone through the maths, we denote the values of the coefficients that minimise f to be $\hat{\beta}_0$, $\hat{\beta}_1$ and $\hat{\beta}_2$, and we can compute the predicted value of the response variable for any values of predictor variables as:

$$\hat{y} = \hat{\beta}_0 + \hat{\beta}_1 x^{(1)} + \hat{\beta}_2 x^{(2)} \quad (8.3)$$

For now, we will skip the derivation of how to find the regression coefficients – it is sufficient to know that it can be done using an extension to the derivation shown in the Linear Regression slides. If you're interested, we give the derivation in the chapter on [Mathematics of multiple regression](#).

8.2 Interpreting multiple regression coefficients and metrics

Interpretation of the coefficients We'll return to Galton's height data. We'll name the variables as follows:

- $x^{(1)}$: midparental height in inches
- $x^{(2)}$: gender (0 for "Daughter" and 1 for "Son")

After fitting the regression function (Equation 8.1) we find the coefficients:

- $\hat{\beta}_0 = 16.41$ inches: the intercept – nominally the height of a child born to parents with zero height!
- $\hat{\beta}_1 = 0.69$: for every inch of midparental height, we expect the child to be 0.687 inches taller.
- $\hat{\beta}_2 = 5.21$ inches: we expect sons to be 5.21 inches taller than daughters.

(Root) mean squared error Since the MSE and RMSE (see the chapter on [Linear Regression](#) just depend on y_i and $\hat{y}_i$, we can compute them using the same formulae. We find the MSE and RMSE in Table 8.1, in which we've also included the MSE and RMSE from:

- the regression of daughter's height on midparental height shown in the [Linear Regression](#) chapter
- the regression of son's height on midparental height (equivalent to the regression of the daughter's height)
- the regression of child's height on midparental height, i.e. as if we did not know the gender of the child.

The RMSE and MSE for the multiple model is mid-way between the values for the two single regression models where we know the gender. This is consistent with the picture of what looks like two regression lines (Figure 8.1). However, if we don't know the gender (imagine we make the prediction before the child is born and without the benefit of a prenatal scan), then the MSE and RMSE are much bigger. This is because the regression line goes midway between the lines that we find for the single-sex regressions or the multiple regressions, so there is more spread around the line.

Interpretation of coefficient of determination Since the coefficient of determination (see the [Linear Regression](#) chapter) just depends on y_i , $\hat{y}_i$ and $\bar{y}$, we can compute it using the same formula, and here we find that $R^2 = 0.633$. This is a lot more than the $R^2 = 0.263$ we found for regressing daughter's height on midparental height. Does this indicate that the multiple regression model is a better fit? That doesn't seem to make sense, as we've just seen the MSE and RMSE of the multiple regression model is about the same as for the single-sex regression models.

The reason is in that the multiple regression, the mean height that is in the total sum of squares is the mean of *all* children, making the total sum of squares (SST) term in the coefficient of determination bigger. As we have more-or-less the same sum of squared errors (SSE) after fitting, there's a much bigger improvement SST – SSE, and therefore a higher value for R^2 .

This isn't the whole picture about the coefficient of determination – we'll return to it later.

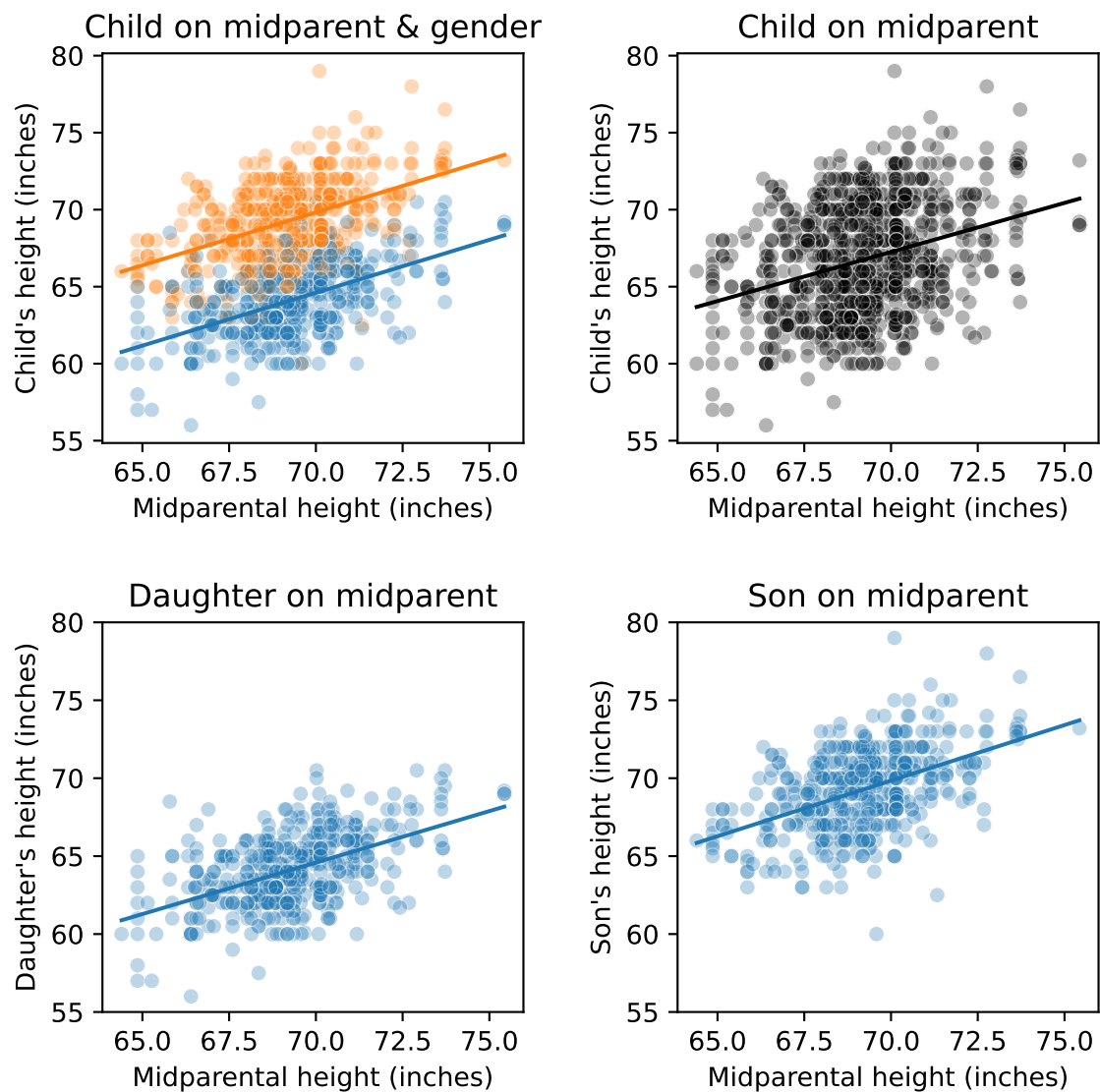


Figure 8.1: Multiple and single regression models of the height data. **Child on midparent & gender:** Multiple regression of child's height on midparent height and child's gender. **Child on midparent:** Single linear regression of child's height on midparent's height. Gender is not a variable, hence the black and grey colour coding. **Daughter on midparent** and **Son on midparent:** Single linear regression of daughter's height on daughter's midparent height (as in last chapter) or son's height on son's midparent height.

Lurking variables Suppose that in the chapter on [Linear Regression](#), we had ignored the gender, and tried to predict the height of sons and daughters just based on their midparental height. In this case we'd have had much higher MSE and RMSE (see table above) and the coefficient of determination would have been 0.10, since there is more variance measured around the mean height of daughters and sons. If we had not considered gender in the simple linear regression, gender would have been a **lurking variable**. In the simple linear regression, we controlled for the lurking variable of gender by only considering the midparental heights of daughters.

Using multiple regression to control for lurking variables By including gender in the multiple regression analysis, we have also controlled for it, and quantified the size of the effect of gender – a double win. However, one side-effect of this model is that the gradients of both regression lines have to be the same, and multiple regression will find a value of $\hat{\beta}_1$ that is not quite optimal for either sons or daughters.

In terms of accurate prediction, we could do better by having two linear regression models: one for the daughters and one for the sons.

8.3 Interaction terms and nonlinear fits

Interaction terms to allow more flexibility There's an alternative to having two models: introduce **interaction terms**, with an extra coefficient, β_3 :

$$y = \beta_0 + \beta_1 x^{(1)} + \beta_2 x^{(2)} + \beta_3 x^{(1)} x^{(2)} \quad (8.4)$$

In the fourth term, we've effectively introduced a new variable $x^{(3)} = x^{(1)} x^{(2)}$, and we can just feed $x^{(3)}$ into the machinery for computing the coefficients. This new variable allows us to have different gradients for the daughters and sons. It will be zero for daughters (since $x^{(2)} = 0$) but it will be $\beta_3 x^{(1)}$ for sons (since $x^{(2)} = 1$), so the effective height gradient for sons will be $\beta_1 + \beta_3$ rather than just β_1 for daughters.

It turns out that the fits we get are exactly the same as the fit we obtain from two separate regression analyses. Why bother then with this interaction term? By using interaction terms, we can avoid having to create separate datasets for each analysis. Also, had $x^{(2)}$ been a continuous variable, we couldn't create a linear regression model for each of its values.

Nonlinear fits The same idea allows us to use linear regression to fit nonlinear curves. Let's go back to our simple linear regression $y = \beta_0 + \beta_1 x$, and suppose that the residual plot shows a quadratic pattern. We can create a new variable x^2 and treat it as a variable in a multiple regression:

$$y = \beta_0 + \beta_1 x + \beta_2 x^2 \quad (8.5)$$

Effectively, we've started to think of y being a function of x and x^2 .

We don't just have to be confined to polynomial functions: we can try other functions too; the problem is choosing the one that is appropriate for the situation. There is also a problem with **over-fitting**: the more flexibility we add to our function the better we can fit – but is the fit really showing something that's generally true about the population data, or is picking up features that have occurred by chance because of the limited size of the sample?

8.4 Interpreting and refining multiple regressions on many variables

We can extend linear regression to many dimensions, but here we'll just increase the number of predictor variables to four. Let's take a look at an example we mentioned at the start: the prediction of student grades. Figure 8.2 shows the data from 80 students. There look to be correlations between all 4 predictor variables and the grade.

Adjusted R^2 We can use the Python `statsmodels` package to run linear regression on the plot. When we do this, we can retrieve a lot of information about the fit (Figure 8.3), but we will just focus on two measures:

- **R-squared** = 0.289: This is the coefficient of determination R^2 , as defined in the [Linear Regression](#) chapter.
- **Adj. R-squared** = 0.251. This is the **adjusted coefficient of determination**, as will be defined below.

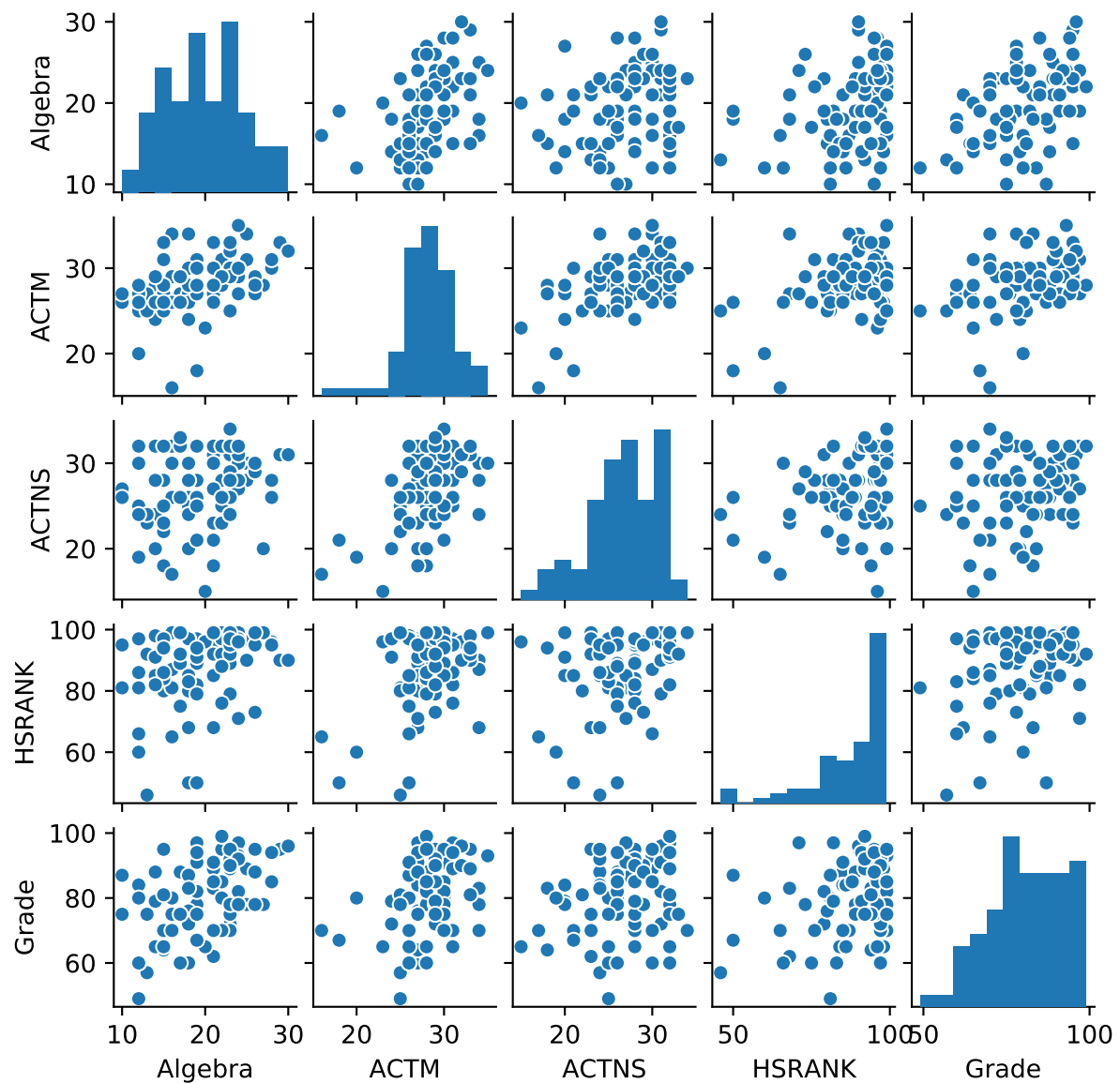


Figure 8.2: Paired scatter plot showing grades of 80 students on a calculus course (**Grade**) and their performance on three earlier tests (**Algebra**, **ACTM** and **ACTNS**) and their high school rank (**HSRANK**). Data from [Edge and Friedberg \(1984\)](#), downloaded from website accompanying [Devore and Berk \(2012\)](#), where it is Example 12.25 <https://extras.springer.com/?query=978-1-4614-0391-3>.

	coef	std err	t	P> t	[0.025	0.975]
Intercept	36.1215	10.752	3.360	0.001	14.703	57.540
Algebra	0.9610	0.264	3.640	0.000	0.435	1.487
ACTM	0.2718	0.454	0.599	0.551	-0.632	1.175
ACTNS	0.2161	0.313	0.690	0.492	-0.408	0.840
HSRANK	0.1353	0.104	1.306	0.196	-0.071	0.342

Dep. Variable:	Grade	R-squared:	0.289
Model:	OLS	Adj. R-squared:	0.251
Method:	Least Squares	F-statistic:	7.622
Date:	Tue, 20 Oct 2020	Prob (F-statistic):	3.30e-05
Time:	18:46:00	Log-Likelihood:	-294.31
No. Observations:	80	AIC:	598.6
Df Residuals:	75	BIC:	610.5
Df Model:	4		
Covariance Type:	nonrobust		

Figure 8.3: Output from the `statmodels.formula.api.ols` routine applied fitting a multiple regression of Algebra, ACTM, ACTNS and HSRANK on Grade (Figure 8.2). At this point, we will focus on R-squared, Adj. R-squared and the `coef` column.

If the number of variables is k and the number of instances is n , the adjusted coefficient of determination is given¹:

$$R_a^2 = 1 - \frac{n-1}{n-(k+1)} \frac{\text{SSE}}{\text{SST}} \quad (8.6)$$

This is very similar to the coefficient of determination, but is, by definition, lower than the coefficient of determination – the term $(n-1)/(n-(k+1))$ is bound to be greater than 1, which means that $R_a^2 < 1 - \text{SSE}/\text{SST} = R^2$.

The reason for adjusting the coefficient of determination is that as we add more and more variables, it becomes easier to fit the data, and therefore the goodness of fit may increase just because of the increase in variables. By adjusting the coefficient of determination, we counteract this tendency.

Meaning of the coefficients Looking at the second table of output in Figure 8.3 we see that:

- Intercept $\hat{\beta}_0 = 36.1215$
- Algebra $\hat{\beta}_1 = 0.9610$
- ACTM $\hat{\beta}_2 = 0.2718$
- ACTNS $\hat{\beta}_3 = 0.2161$
- HSRANK $\hat{\beta}_4 = 0.1353$

The interpretation is that an increase of 1 point in the Algebra test predicts an increase of 0.961 points in the final grade, whereas an increase of 1 in the HSRANK predicts only an increase of 0.135 in the final grade.

Correlated predictor variables The scatter plots (Figure 8.2) show considerable correlation between the predictor variables. We may imagine that if we did a single linear regression with the *Algebra* scores as the predictor variable, we could explain a considerable amount of the variance in the *Grade*. We can check this by re-running the model with just the *Algebra* as a predictor, in which case we find $R_a^2 = 0.231$, not much less than the R_a^2 for the model with 4 predictor variables. Those extra variables don't seem to have explained much more.

We can also investigate what would happen if we tried to fit using only one of the other variables. When we regress *Grade* on *ACTM*, we find $R_a^2 = 0.124$, which is about half the size of R_a^2 for regressing on *Algebra*. Note that the $R_a^2 = 0.247$ when we regress on both *Algebra* and *ACTM* is less than the sum of the two adjusted coefficients of determination ($0.231 + 0.124 = 0.355$). The reason for this is the correlation between *Algebra* and *ACTM*: knowing *Algebra* tells us a lot about *ACTM*, so *ACTM* doesn't add very much new information. The lesson to learn here is that we need to think carefully about which variables to include as predictor variables in a multiple regression analysis. It may make sense to take out a variable that adds little to the adjusted coefficient of determination. In fact, some of the output in Figure 8.3 that we haven't discussed gives us clues about which variables to drop, but we will discuss that later in the course.

Collinear variables The extreme case of correlation is when a pair of predictor variables are perfectly correlated, for example: $x^{(1)} = cx^{(2)}$, with c being a constant. Here the correlation coefficient between $x^{(1)}$ and $x^{(2)}$ is 1 or -1 . In this case a linear regression function in a stats package like `statsmodels` in Python will complain about a singular matrix. The full mathematical explanation of why this problem arises is in the section on [Interpreting the equation for the coefficients](#). Basically, the problem is that there are an infinite number of solutions of the values of $\hat{\beta}_1$ and $\hat{\beta}_2$.

Collinear variables can occur surprisingly frequently, for example, there could be two columns with a height in centimetres and the same height in inches. To fix the problem of collinear variables, we should take out one of the correlated variables – it is adding no information anyway.

Highly correlated variables When the magnitude of the correlation between two predictor variables is almost one, there can still be numerical stability issues that cause numerical routines in stats packages to fail. Also, small differences in the correlation of $x^{(1)}$ and $x^{(2)}$ with y can lead to very different estimates of the coefficients. The interpretation of the coefficients therefore needs particular care, or one of the highly correlated predictor variables should be removed.

¹It would be nice to use the letter D to denote the number of dimensions of the predictor variables, but this is not the convention in the literature on multiple regression. Normally it's either k (as used here and by Devore and Berk (2012)) or p

 Related Workshop: Interpretation of correlation and linear regression

<https://opencourse.inf.ed.ac.uk/inf2-fds/course-materials/semester-1/week-9/workshop>

The goals of the task and workshop are:

- to develop your ability to interpret multiple regression analyses
- to develop your skills in reading and critiquing data-driven methods and claims from case studies, in order to identify and discuss the extent to which stated conclusions are warranted given evidence provided

Chapter 9

Mathematics of multiple regression

9.1 Derivation of coefficients in multiple regression

Derivation of multiple regression coefficients, part 1 We start by repeating Equation 8.1 in the [Multiple regression](#) chapter for the regression line with two predictor variables:

$$y = \beta_0 + \beta_1 x^{(1)} + \beta_2 x^{(2)} \quad (9.1)$$

We could work generally with k variables, but we will stick with 2 variables for now, as we feel it gives a better intuition about what's going on. To find values of $\hat{\beta}_0$, $\hat{\beta}_1$ and $\hat{\beta}_2$ that minimise:

$$f(\beta_0, \beta_1, \beta_2) = \sum_{i=1}^n (y_i - (\beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2}))^2 \quad (9.2)$$

we will modify the method we used for simple linear regression ([Linear regression](#) chapter). We could set the partial derivatives of the error function (or loss) f with respect to β_0 , β_1 and β_2 equal to 0, and work from the resulting three equations. However, we get some nice insights by proceeding in two steps. First we'll set the partial derivative of f with respect to β_0 to be 0:

$$\frac{\partial f}{\partial \beta_0} = \sum_{i=1}^n (-2)(y_i - \beta_0 - \beta_1 x_{i1} - \beta_2 x_{i2}) = 0 \quad (9.3)$$

All the x_{ij} and y_i are constants, so this is an equation in β_0 , β_1 and β_2 . If we divide through by $-2n$ and define the mean of the j th predictor variable $\bar{x}^{(j)} = 1/n \sum_i x_{ij}$ we get:

$$\begin{aligned} \bar{y} - \beta_0 - \beta_1 \bar{x}^{(1)} - \beta_2 \bar{x}^{(2)} &= 0 \\ \Rightarrow \beta_0 &= \bar{y} - \beta_1 \bar{x}^{(1)} - \beta_2 \bar{x}^{(2)} \end{aligned} \quad (9.4)$$

We now substitute this expression for β_0 into Equation (9.1) for the regression line:

$$\begin{aligned} y &= \bar{y} - \beta_1 \bar{x}^{(1)} - \beta_2 \bar{x}^{(2)} + \beta_1 x^{(1)} + \beta_2 x^{(2)} \\ y - \bar{y} &= \beta_1 (x^{(1)} - \bar{x}^{(1)}) + \beta_2 (x^{(2)} - \bar{x}^{(2)}) \end{aligned} \quad (9.5)$$

This equation shows that the regression plane will pass through the point $(\bar{x}^{(1)}, \bar{x}^{(2)}, \bar{y})$, which is what happens with the linear regression line in simple linear regression. It is as though the plane is pinned to this location, and there are only two coefficients left to adjust: β_1 and β_2 .

It will simplify the following analysis to define new versions of the variables in which the mean has been subtracted:

$$y_i^* = y_i - \bar{y} \quad , \quad x_{ij}^* = x_{ij} - \bar{x}^{(j)} \quad (9.6)$$

Here we've used the stars to denote that this is a version of the variable with the mean subtracted – the star does not mean “complex conjugate”!

Derivation of multiple regression coefficients, part 2 We can now start again from the least squares function f^* that is a function of β_1 and β_2 , and which contains our mean-subtracted variables:

$$f^*(\beta_1, \beta_2) = \sum_{i=1}^n (y_i^* - (\beta_1 x_{i1}^* + \beta_2 x_{i2}^*))^2 \quad (9.7)$$

We can now partially differentiate with respect to β_1 and β_2 to give the normal equations:

$$\begin{aligned} \frac{\partial f^*}{\partial \beta_1} &= \sum_{i=1}^n (-2x_{i1}^*)(y_i^* - \beta_1 x_{i1}^* - \beta_2 x_{i2}^*) = 0 \\ \frac{\partial f^*}{\partial \beta_2} &= \sum_{i=1}^n (-2x_{i2}^*)(y_i^* - \beta_1 x_{i1}^* - \beta_2 x_{i2}^*) = 0 \end{aligned} \quad (9.8)$$

We can divide both sides of the normal equations by -2 and then rewrite them as one matrix equation:

$$\begin{aligned} &\begin{pmatrix} \sum_{i=1}^n x_{i1}^*(y_i^* - \beta_1 x_{i1}^* - \beta_2 x_{i2}^*) \\ \sum_{i=1}^n x_{i2}^*(y_i^* - \beta_1 x_{i1}^* - \beta_2 x_{i2}^*) \end{pmatrix} = 0 \\ \Rightarrow &\begin{pmatrix} \sum_{i=1}^n x_{i1}^*(\beta_1 x_{i1}^* + \beta_2 x_{i2}^*) \\ \sum_{i=1}^n x_{i2}^*(\beta_1 x_{i1}^* + \beta_2 x_{i2}^*) \end{pmatrix} = \begin{pmatrix} \sum_{i=1}^n x_{i1}^* y_i^* \\ \sum_{i=1}^n x_{i2}^* y_i^* \end{pmatrix} \end{aligned} \quad (9.9)$$

Now suppose that we define the matrix $\mathbf{X}$ to be an $n \times 2$ matrix in which the first column contains the n values of x_{i1}^* and the second column contains the n values of x_{i2}^* . We call this matrix the **design matrix** or **regressor matrix**. We'll define $\mathbf{y}$ to be the vector containing the n values of y_i^* . By the definition of matrix multiplication, you should be able to verify that the right-hand side of the equation is equal to the matrix product $\mathbf{X}^T \mathbf{y}$:

$$\mathbf{X}^T \mathbf{y} = \begin{pmatrix} \sum_{i=1}^n x_{i1}^* y_i^* \\ \sum_{i=1}^n x_{i2}^* y_i^* \end{pmatrix} \quad (9.10)$$

($\mathbf{X}^T \mathbf{y}$ is referred to as the **moment matrix**.)

We can also rewrite the left-hand side in terms of a matrix multiplication of a 2×2 matrix and a vector of coefficients:

$$\begin{pmatrix} \sum_{i=1}^n x_{i1}^*(\beta_1 x_{i1}^* + \beta_2 x_{i2}^*) \\ \sum_{i=1}^n x_{i2}^*(\beta_1 x_{i1}^* + \beta_2 x_{i2}^*) \end{pmatrix} = \begin{pmatrix} \sum_{i=1}^n x_{i1}^* x_{i1}^* & \sum_{i=1}^n x_{i1}^* x_{i2}^* \\ \sum_{i=1}^n x_{i2}^* x_{i1}^* & \sum_{i=1}^n x_{i2}^* x_{i2}^* \end{pmatrix} \begin{pmatrix} \beta_1 \\ \beta_2 \end{pmatrix} \quad (9.11)$$

There are two things to note about the 2×2 matrix:

1. If you look back to the definitions of sample variance and sample covariance (Equation 6.1 in the chapter on [Linear Regression](#)), you will see that its diagonal elements are $(n-1)$ times the variances of $x^{(1)}$ and $x^{(2)}$ and its off-diagonal elements are $(n-1)$ times the covariance of $x^{(1)}$ and $x^{(2)}$.
2. The matrix can be written as $\mathbf{X}^T \mathbf{X}$, which can be referred to as the **normal matrix**.
3. We define the **covariance matrix** to be $\mathbf{S} = \frac{1}{n-1} \mathbf{X}^T \mathbf{X}$

We can now rewrite Equation 9.9 as a matrix equation:

$$\mathbf{X}^T \mathbf{X} \boldsymbol{\beta} = \mathbf{X}^T \mathbf{y} \quad (9.12)$$

We can then solve it for the vector of coefficients:

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \quad (9.13)$$

The vector on the left contains the values of $\hat{\beta}_1$ and $\hat{\beta}_2$, which we can substitute back into Equation 9.4 to get $\hat{\beta}_0$.

Notes on our derivation The derivation can be extended to k predictor variables – Equation 9.13 looks the same, but $\hat{\boldsymbol{\beta}}$ is a $k \times 1$ vector, and the design matrix $\mathbf{X}$ is an $n \times k$ matrix. The normal matrix (and covariance matrix) end up being $k \times k$ matrices.

It's worth noting that there are other ways of deriving the coefficients – many treatments in textbooks add a column of 1s to the design matrix and include β_0 in the vector of coefficients. We've chosen not to do this, as it makes the connection with the covariance matrix clearer.

It's also worth noting that although you can program this equation yourself, real-world multiple regression routines use other matrix formulations for reasons of efficiency and numerical stability.

9.2 Interpreting the equation for the coefficients

Equation 9.13 is elegant but also a bit abstract, so we'll try to interpret what the terms in it mean. To do so we'll consider 2 cases, sticking with 2 predictor variables for simplicity.

Interpretation of the derivation 1: no covariance We'll now suppose that the s_1^2 , the variance of $x^{(1)}$, and s_2^2 , the variance of $x^{(2)}$, are non-zero, but the covariances are 0 (Figure 9.1, left), so the covariance matrix is:

$$\mathbf{S} = \begin{pmatrix} s_1^2 & 0 \\ 0 & s_2^2 \end{pmatrix} = \frac{1}{n-1} \mathbf{X}^T \mathbf{X}$$

We therefore have

$$(\mathbf{X}^T \mathbf{X})^{-1} = \frac{1}{n-1} \mathbf{S}^{-1} = \frac{1}{n-1} \begin{pmatrix} 1/s_1^2 & 0 \\ 0 & 1/s_2^2 \end{pmatrix}$$

and so now our estimates of $\hat{\beta}_1$ and $\hat{\beta}_2$ are

$$\begin{pmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{pmatrix} = \begin{pmatrix} \frac{1}{(n-1)s_1^2} \sum_{i=1}^n x_{i1}^* y_i^* \\ \frac{1}{(n-1)s_2^2} \sum_{i=1}^n x_{i2}^* y_i^* \end{pmatrix} = \begin{pmatrix} s_{1y}/s_1^2 \\ s_{2y}/s_2^2 \end{pmatrix} \quad (9.14)$$

Here we've used the notation s_{1y} as shorthand for the covariance of $x^{(1)}$ and y .

Now we'll denote the correlation of $x^{(1)}$ and y by r_{1y} , and remember that the correlation coefficient between two variables is defined as covariance divided by the product of the standard deviations, in this case: $r_{1y} = s_{1y}/(s_1 s_y)$. Therefore, we can write the covariance of $x^{(1)}$ and y in terms of the correlation $s_{1y} = r_{1y} s_1 s_y$. We have similar definitions for r_{2y} . Substituting in Equation 9.14 gives us:

$$\begin{pmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{pmatrix} = \begin{pmatrix} r_{1y} s_y / s_1 \\ r_{2y} s_y / s_2 \end{pmatrix} \quad (9.15)$$

These are exactly the coefficients we found by doing a single linear regression of y on $x^{(1)}$ and $x^{(2)}$ separately. How much a predictor variable that is not correlated with another predictor variable influences our estimate of y depends purely on its correlation with y .

Interpretation of the derivation 2: the general case Now we'll allow the covariances to be non-zero. We simplify calculations by denoting the correlation between $x^{(1)}$ and $x^{(2)}$ as r_{12} , so $r_{12} = s_{12}/s_1 s_2$. Thus, the covariance $s_{12} = r_{12} s_1 s_2$, and the covariance matrix can be written as:

$$\mathbf{S} = \begin{pmatrix} s_1^2 & r_{12} s_1 s_2 \\ r_{12} s_1 s_2 & s_2^2 \end{pmatrix} = \frac{1}{n-1} \mathbf{X}^T \mathbf{X}$$

The inverse of the normal matrix is:

$$(\mathbf{X}^T \mathbf{X})^{-1} = \frac{1}{n-1} \mathbf{S}^{-1} = \frac{1}{n-1} \frac{1}{1-r_{12}^2} \begin{pmatrix} 1/s_1^2 & -r_{12}/(s_1 s_2) \\ -r_{12}/(s_1 s_2) & 1/s_2^2 \end{pmatrix} \quad (9.16)$$

When we multiply by

$$\mathbf{X}^T \mathbf{y} = (n-1) \begin{pmatrix} r_{1y} s_1 s_y \\ r_{2y} s_2 s_y \end{pmatrix}$$

we end up with

$$\begin{pmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{pmatrix} = \frac{1}{1-r_{12}^2} \begin{pmatrix} r_{1y} s_y / s_1 - r_{12} r_{2y} s_y / s_1 \\ r_{2y} s_y / s_2 - r_{12} r_{1y} s_y / s_2 \end{pmatrix} = \frac{1}{1-r_{12}^2} \begin{pmatrix} (r_{1y} - r_{12} r_{2y}) s_y / s_1 \\ (r_{2y} - r_{12} r_{1y}) s_y / s_2 \end{pmatrix} \quad (9.17)$$

The coefficient values found in the no-covariance case (Equation 9.15) are still there (when $r_{12} = 0$), but we see that when there is a non-zero correlation between the predictor variables, the coefficients are altered. The no-covariance estimate for $\hat{\beta}_1$ (i.e. r_{1y}) is adjusted by subtracting a fraction r_{12} of the correlation of the response variable with the other predictor variable, r_{2y} . This makes sense, since if didn't make this correction, the contribution $\hat{\beta}_2 x^{(2)}$ would "pollute" our estimate of y . Likewise, $\hat{\beta}_2$ is adjusted by subtracting $r_{12} r_{1y}$.

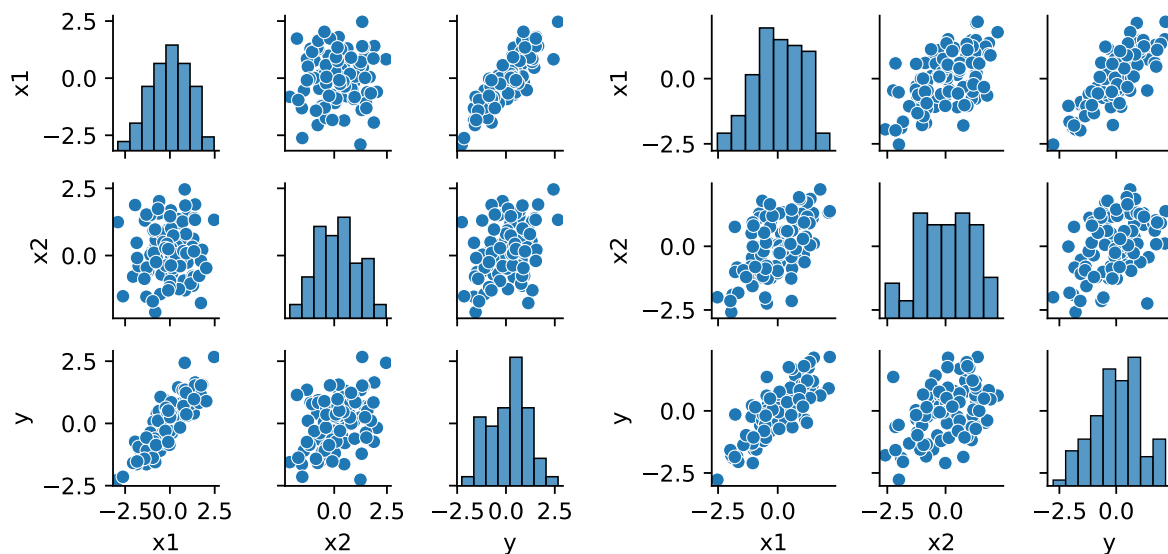


Figure 9.1: Correlated data, generated for the purposes of demonstration. In both plots the correlations between the predictor variables and the response variables are $r_{1y} = 0.8$ and $r_{2y} = 0.4$, and all variables have unit variance. In the left-hand plot the predictor variables are uncorrelated $r_{12} = 0$. The expected regression coefficients are $\hat{\beta}_1 = r_{1y} = 0.8$ and $\hat{\beta}_2 = r_{2y} = 0.4$. In the right-hand plot, the predictor variables have a correlation $r = 0.5$, but due to the correlation between the predictor variables, the regression coefficients change to $\hat{\beta}_1 = 0.8$ and $\hat{\beta}_2 = 0$. See text for details.

For example, suppose we have unit variance variables ($s_1 = s_2 = 1$) and $r_{12} = 0.5$, $r_{1y} = 0.8$, $r_{2y} = 0.4$ (Figure 9.1, right). Then we have

$$\begin{pmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{pmatrix} = \frac{1}{1 - 0.5^2} \begin{pmatrix} 0.8 - 0.5 \times 0.4 \\ 0.4 - 0.5 \times 0.8 \end{pmatrix} = \begin{pmatrix} 0.8 \\ 0 \end{pmatrix}$$

Essentially the predictor variable with the stronger correlation with the response variable reduces the influence of the predictor variable with the weaker correlation. The more informative predictor variable “wins” the competition to explain the response variable.

There is a special case when r_{12} is non-zero, but $r_{1y} = r_{2y}$. The coefficients become:

$$\begin{pmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{pmatrix} = \frac{1 - r_{12}}{1 - r_{12}^2} \begin{pmatrix} r_{1y}s_y/s_1 \\ r_{2y}s_y/s_2 \end{pmatrix} = \frac{1}{1 + r_{12}} \begin{pmatrix} r_{1y}s_y/s_1 \\ r_{2y}s_y/s_2 \end{pmatrix}$$

Here $x^{(1)}$ and $x^{(2)}$ have equal correlation with y , so are equally informative about y . As the correlation between the predictor variables gets stronger the coefficients are scaled down; when r_{12} approaches 1, the coefficients are half what they would be in the case of uncorrelated predictor variables. If two people are singing the same song, you can halve the volume of both singers, and still hear the same information.

Collinearity The extreme case of correlation is $r_{12} = 1$, when $x^{(1)} = cx^{(2)}$, with c being a constant. In this case the denominator of Equation 9.17 is zero. This reflects the fact that the determinant of the covariance matrix $\mathbf{S}$ is zero – the two rows of the matrix $\mathbf{X}^T\mathbf{X}$ in Equation 9.11 are multiples of each other. There is therefore no solution to Equation 9.13, and a linear regression function in a stats package like `statsmodels` in Python will complain about a singular matrix.

When r_{12} is large there are still problems, since small differences in the correlation of $x^{(1)}$ and $x^{(2)}$ with y can lead to very different estimates of the coefficients. As stated in the Chapter on [Multiple regression](#), the interpretation of the coefficients therefore needs particular care.

Chapter 10

Dealing with high dimensions – PCA

10.1 The principle of Principal Components Analysis (PCA)

The challenges of high dimensions In the multiple regression topic, in the student grade prediction example, we were beginning to see two challenges of dealing with more than one predictor variable:

The challenge of visualisation We can see a lot in the paired correlation plots. With 4 predictor variables, the visualisation works, but what about if we had 26 variables? The Scottish Index of Multiple Deprivation (SIMD, Table 10.1) records 26 variables for each of 6527 data zones in Scotland. A 26×26 grid of scatter plots is going to be difficult to read.

The challenge of interpretation In the grades example, the test grades (predictor variables) were correlated, which made the interpretation of the regression coefficients challenging – and this was with only 4 predictor variables. In the SIMD example, we might expect many of the 26 variables to be correlated, e.g. the time it takes to drive to the nearest primary school and the time it takes to drive to the nearest secondary school.

There is also another problem with high-dimensional data, called the **curse of dimensionality**: essentially a large number of dimensions makes it harder for distance-based methods such as clustering and nearest neighbours to work effectively – we'll come back to the curse of dimensionality in the following lectures on clustering and nearest-neighbour methods.

In **dimensionality reduction** methods these challenges are addressed by reducing the number of dimensions in the data while retaining as much useful information as possible. There are a number of dimensionality reduction methods which differ in what aspects of the data they preserve.

Principal components analysis We are going to discuss one method of dimensionality reduction called **principal components analysis** (PCA).

PCA can be applied to a set of D numeric variables with n datapoints. In contrast to linear regression, all variables are treated equally: there is no response variable that we are trying to predict, just a set of variables whose structure we're trying to understand better. The result of PCA is a set of up to D new variables (with n datapoints). We can keep $k \leq D$ of the most informative new variables.

In PCA the objectives are:

Table 10.1: Scottish Index of Multiple Deprivation, 2016 edition ([Scottish Government, 2016](https://simd.scot)). <https://simd.scot>. It has $n = 6527$ data points (postcode zones), each associated with $D = 26$ variables.

Location	Employ- ment	Illness	Attain- ment	Drive Primary	Drive Secondary	Crime	...
Macduff	10	95	5.3	1.5	6.6	249	...
Kemnay	3	40	5.3	2.4	2.4	168	...
Hilton	0	10	6.3	2.2	3.0	144	...
Ruchill	8	130	4.9	1.7	5.6	318	...
Belmont	2	50	6.1	3.1	3.2	129	...
...	...	...	...	...	...	...	...

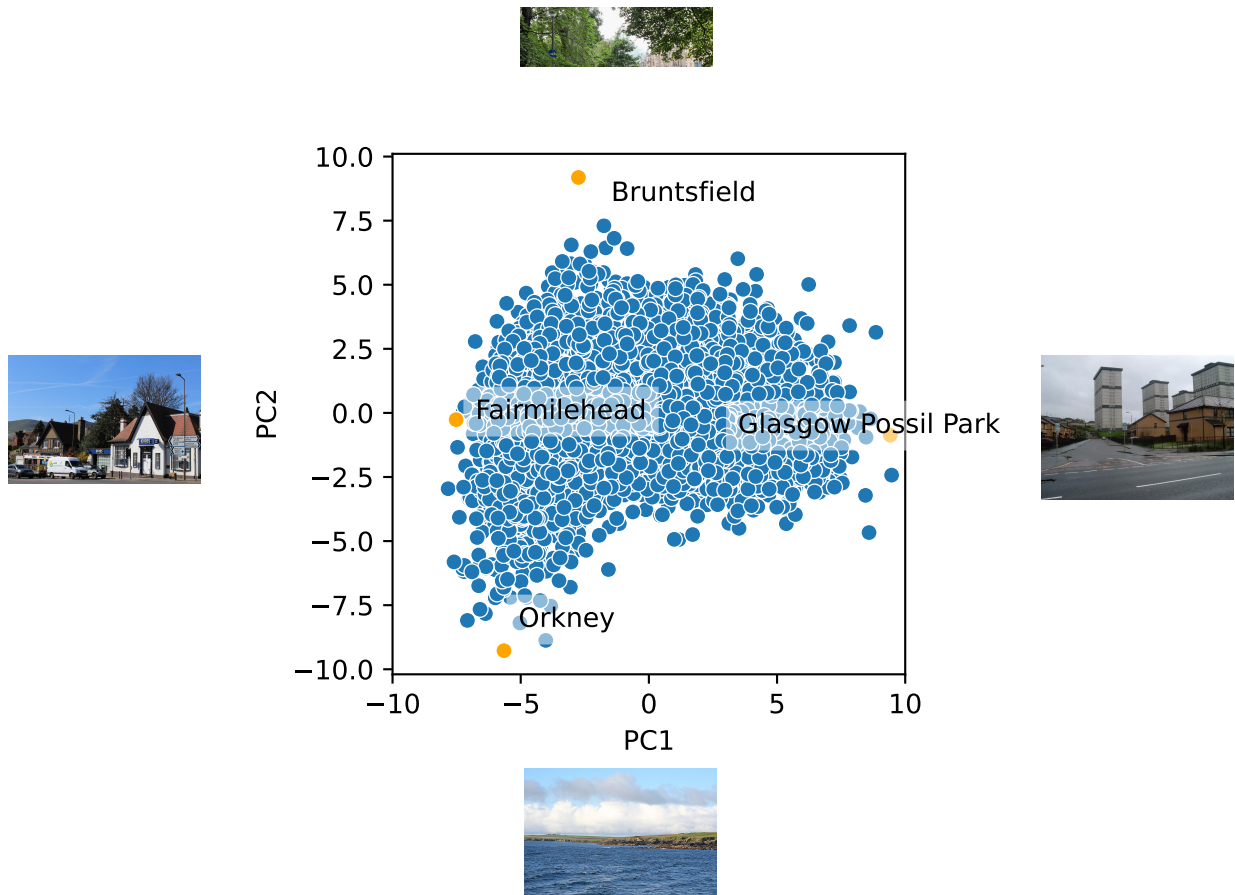


Figure 10.1: Scatter plot of first and second principal component scores (PC1 and PC2) of 6527 data points in the SIMD dataset (blue dots). Locations of 4 data zones are indicated in orange dots next to an image from that data zone. All photos released under CC licence from geograph.co.uk. Credits: Orkney © Des Colhoun; Possil Park © Stephen Sweeney; Bruntsfield © Leslie Barrie; Fairmilehead © Jim Barton.

1. change the angle we view the data from to see things clearly
2. ignore small details in the data that don't affect the big picture.

We'll specify these objectives more precisely and explain how PCA works later. First, we will show the results when PCA is applied to the SIMD example (Table 10.1).

Example of PCA We can use PCA to reduce the number of variables D in the SIMD data from $D = 26$ to $k = 2$, allowing us to visualise all $n = 6527$ data points (Figure 10.1). In this plot, the i th datapoint has coordinates (t_{i1}, t_{i2}) in which each coordinate is a linear combination of the standardised data z_{ij} shown in Table 10.1:

$$\begin{aligned} t_{i1} &= p_{11}z_{i1} + p_{21}z_{i2} + \cdots + p_{D1}z_{iD} \\ t_{i2} &= p_{12}z_{i1} + p_{22}z_{i2} + \cdots + p_{D2}z_{iD} \end{aligned} \tag{10.1}$$

The weights $p_{11}, p_{21}, \dots, p_{D1}$ are elements of the **first principal component** and t_{i1} is the first principal **component score** of the i th datapoint; we will explain how to find them later. Likewise, $p_{12}, p_{22}, \dots, p_{D2}$ form the second principal component and t_{i2} is the second principal **component score** of datapoint i . The weights in the principal component indicate how much influence each original variable has over each principal component score – sometimes they are referred to as **loadings** or **weights**. The axes in Figure 10.1 are labelled PC1 (first principal component – “PC” stands for principal component) and PC2 (second principal component).

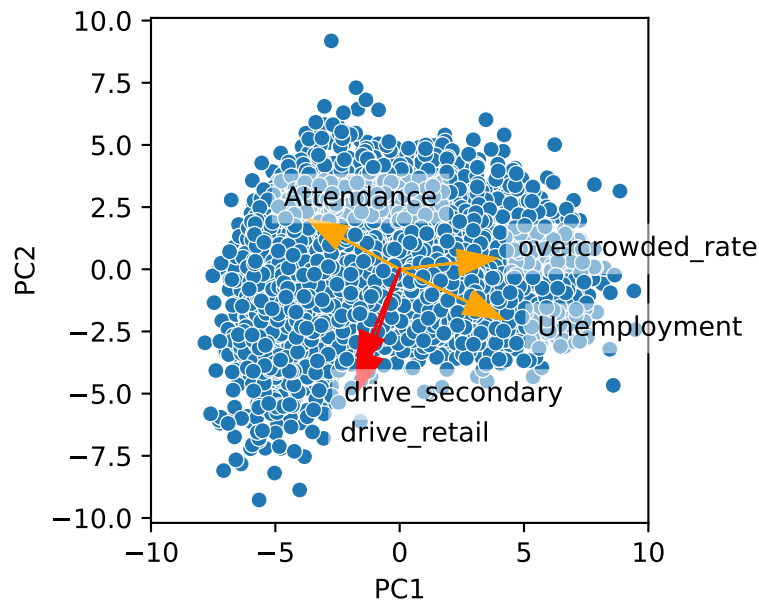


Figure 10.2: Scatter plots of first and second principal component scores of SIMD data zones (blue dots). The projection of three original variables related to deprivation are shown as orange arrows emanating from the origin. High unemployment and overcrowded rate are found in areas with higher deprivation, whereas high school attendance is found in areas with low deprivation. These vectors are more closely aligned with the first principal component (PC1), which we therefore interpret as “Deprivation”. Red arrows indicate the projections of the time take to drive to the nearest secondary school or retail outlet. As these are aligned with PC2, we therefore interpret PC2 as being related to distance to services, or “Remoteness”.

💡 Standardised variables notation for multivariate data

In the section on [Standardised variables](#) defined the standardised version of a single variable x with instances x_i . Here, as in the chapter on [Multiple regression](#) we have multiple variables, denote the j th variable $x^{(j)}$, and denote the i th instance of the j th variable x_{ij} . We standardise the j th variable $x^{(j)}$ by subtracting its mean $\bar{x}^{(j)}$ and dividing by its standard deviation s_j , so that

$$z_{ij} = (x_{ij} - \bar{x}^{(j)})/s_j$$

As a general rule, we should always standardise variables before applying PCA to them.

To see the influence of each original variable on PC1 and PC2 scores, we can project the j th original variable onto the plot by setting z_{ij} to 1 and all the other z 's to 0 in Equation 10.1. In this case, the coordinates we'll be plotting are (p_{j1}, p_{j2}) . The orange arrows in Figure 10.2 show the projections of the variables for unemployment, overcrowding (in housing) and school attendance. Unemployment and overcrowding have high PC1 scores. In contrast, school attendance has a low PC1 score. This all makes sense if we identify the first component score with “Deprivation”. We can rephrase the previous sentences as “unemployment and overcrowding are found in areas of high deprivation and high school attendance is found in areas of low deprivation”.

The red arrows in Figure 10.2 show the projections of the time to drive to the nearest retail outlet and time to drive to the nearest secondary school. These vectors have higher magnitude PC2 scores than PC1 scores. We therefore identify PC2 as being to do with “remoteness” – low values of PC2 indicate the zone is more remote.

Note that the correlation between the PC1 and PC2 scores is zero. It is a general property of PCA there are no correlations between the scores of different principal components.

In this particular example, the visualisation shows a unimodal distribution of data with little obvious structure. Later on in the course we will see examples where PCA reveals clusters of data – though still

with zero correlation.

Even if no structure is apparent, reducing the dimensionality of the data can be useful for further analysis. For example, suppose we have data on cancer screening rates in each data SIMD zone, we could then do multiple regression of the cancer screening rate on the new deprivation and remoteness variables. This is probably going to give us coefficients that are a lot more interpretable than regressing on all 26 variables.

Projecting principal component scores back into the data space Suppose we have identified the first two principal component scores t_{i1} and t_{i2} of area i . We might wish to project them back into the data space, to see what the original variables looked like. To do this we can use the following equations to give approximations (indicated by the tilde) to the original standardised variables:

$$\begin{aligned}\tilde{z}_{i1} &= p_{11}t_{i1} + p_{12}t_{i2} \\ \tilde{z}_{i2} &= p_{21}t_{i1} + p_{22}t_{i2} \\ &\vdots \\ \tilde{z}_{iD} &= p_{D1}t_{i1} + p_{D2}t_{i2}\end{aligned}\tag{10.2}$$

We can include more terms for higher PCs, right up to the D th PC. In general, the j th component of the i data point is given:

$$\begin{aligned}z_{ij} &= p_{j1}t_{i1} + p_{j2}t_{i2} + \dots + p_{jD}t_{iD} \\ &= \sum_{k=1}^D p_{jk}t_{ik}\end{aligned}\tag{10.3}$$

Once we've got the standardised variables, we can convert back to the original variables using the formula $x_{ij} = z_{ij}s_j + \bar{x}_j$.

Principal component equations in vector notation The equations used so far may make more sense when expressed as vectors. The j th principal component is actually a vector in the original data space:

$$\mathbf{p}_j = (p_{1j}, p_{2j}, \dots, p_{Dj})^T\tag{10.4}$$

All the principal component vectors are orthogonal to each other. With this notation we can write Equation 10.2 as a linear combination of the principal component vectors, weighted by the principal component scores:

$$\mathbf{z}_i = t_{i1}\mathbf{p}_1 + t_{i2}\mathbf{p}_2 + \dots\tag{10.5}$$

The dots indicate that we could go up to $t_{iD}\mathbf{p}_D$.

We can rewrite Equation 10.1, in which we computed the scores, as the scalar product of the i th standardised data point and the j th principal component:

$$t_{ij} = \mathbf{z}_i \cdot \mathbf{p}_j\tag{10.6}$$

We'll extend this notation to matrix notation in the derivation.

10.2 Principle of finding principal components

A 2D example We'll now discuss the principle of how to determine the principal components with an imaginary 2D example. Suppose we ask if there are different types of Informatics students, perhaps based on their preferences for programming languages and for drinks. We ask students if they prefer, on a scale of 1–9, Haskell (1) to Java (9), and if they prefer Tea (1) to Coffee (9), and find the data in Table 10.2.

Plotting the data (Figure 10.3 left) shows that students' preferences for drinks and programming languages are correlated. It seems that we could characterise every Informatics student by one number that is low if they like Haskell and tea, and high if they like Java and coffee. If we could rotate the axes (Figure 10.3 right), the new x -axis would give us this number.

Once we've done the rotation (changed the angle), we end up with the data plotted against a new set of axes, which are the principal components (Figure 10.4, top). The new x -axis, which tells us a lot about the students' preferences for Java and coffee or tea and Haskell, is the first principal component (PC1). The new y axis is the second principal component (PC2). It is worth noting two points:

Table 10.2: Imaginary data about Informatics students' preferences for programming languages and drinks.

Student ID	Language	Drink
1	9	8
2	3	1
3	8	7
4	2	2
5	3	3
6	8	6
7	2	3
8	8	8
9	1	2
10	6	7

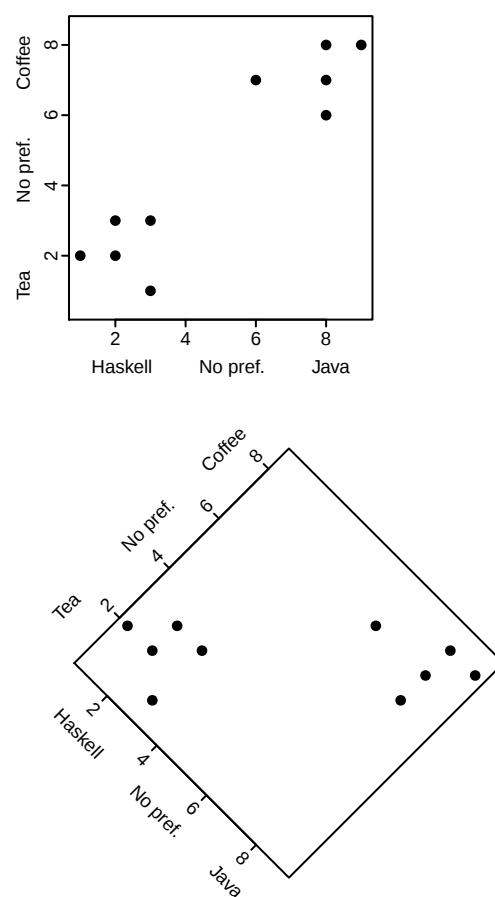
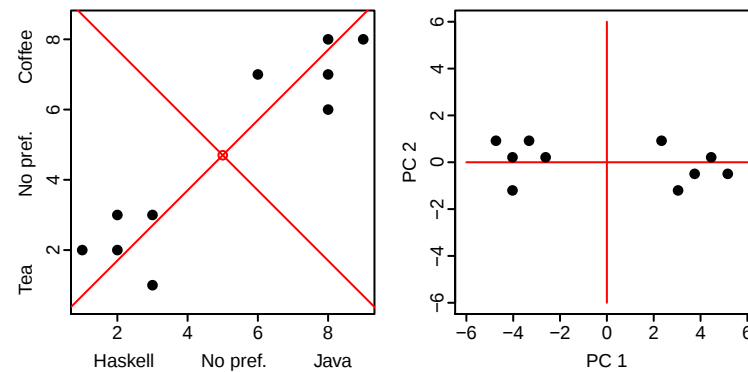


Figure 10.3: Informatics students' preferences for drinks and programming languages, as plotted initially (left), and rotated (right).

1. Change the angle we view the data from to see things clearly



2. Ignore small details in the data that don't affect the big picture

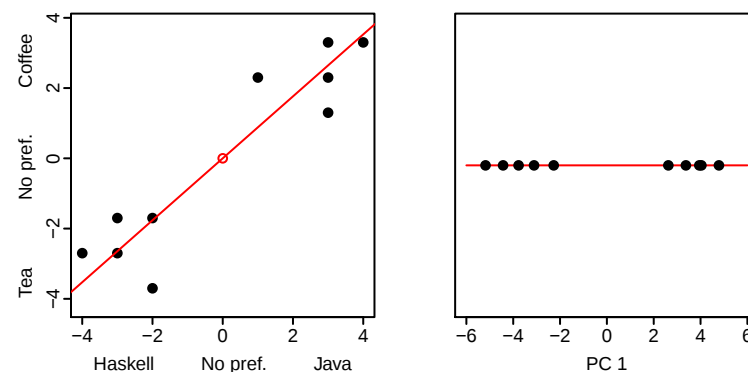


Figure 10.4: Visualisation of how PCA achieves the two objectives in the text.

- The correlation between the new PC1 and PC2 scores is zero. It is a general property of PCA that correlations between scores is zero.
- We have not lost any information about the data; we can reconstruct the original data by reversing the rotation. It is a general property of PCA that it is possible to reconstruct the data if scores of all D principal components are retained.

The second principal component doesn't seem so informative, so we could just ignore it altogether (Figure 10.4, bottom). Thus, we have ignored small details in the data that don't affect the big picture. We have performed dimensionality reduction by reducing the number of values describing each data point from two to one.

Objective of rotation There are two questions that we haven't answered so far:

1. How do we choose how much to rotate the axes?
2. What counts as “informative”?

The answer to both questions is “variance”. In Figure 10.4 (top), the variance of the data in the PC1 direction is much greater than the variance of the data in the PC2 direction. The high variance PC1 is telling us a lot about the informatics students, whereas the low variance PC2 tells us little. Therefore, in order to choose how to rotate the axes, we use the variance as an objective. In fact there are two ways of formulating PCA:

1. Maximum variance formulation: find an axis that maximises the variance of the data projected onto it
2. Minimum variance formulation: find an axis that minimises the variance of the data projected onto it

It doesn't matter which formulation we use; the answer is the same either way.

Explained variance The variance in the original x (Programming language) and y (Drink) directions was 9.7 and 7.7. The sum of these two variances is the total variance, i.e. 17.4. It turns out that the sum of the variance along the principal components is exactly the same. However, the variance of the PC1 scores is 16.5, i.e. 96% of the total variance. We therefore say the PC1 explains 96% of the variance.

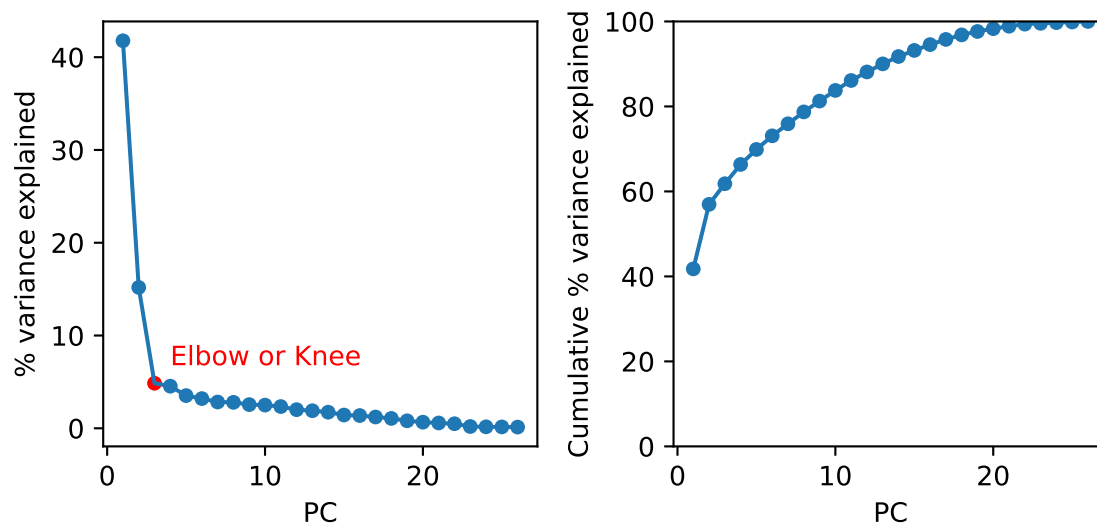


Figure 10.5: Scree plot for PCA applied to SIMD example (left). The elbow (or knee) is indicated in red. The Cumulative scree plot (right).

More than 2D In general, we can find D principal components in D dimensions. The principal components are all orthogonal to each other, and each principal component explains a certain fraction of the variance. We order the principal components from the one that explains most variance to the one that explains least.

In the SIMD example, the first principal component explains 41.7% of the data and the second explains a further 15.2%. Thus, the first two principal components together explain 56.9% of the variance. We can visualise how much each principal component explains in a **scree plot** or **cumulative scree plot** (Figure 10.5).

How many components to choose? Obviously if we are visualising data, we can only look straightforwardly at up to 3 dimensions. The scree plot helps us to choose how many components to include if we are using PCA as a preprocessing step. A rule of thumb is to use the components to the left of the **elbow** or **knee** of the scree plot, i.e. the point where the gradient changes sharply. In Figure 10.5 this point is indicated in red, and the rule of thumb would suggest that we use PC1 and PC2. There are more principled ways of choosing, which we won't cover at this point, and it may also be that successful application of PCA requires more components.

In the next section, we'll look at the maths of how to find the directions of the principal components and the associated variances. However, you should already know enough to skip to the section after, which is about applying PCA to help with a regression problem.

10.3 PCA and regression

PCA as preprocessing PCA is often used as a preprocessing step before another method, e.g. linear regression or K-means. Here we'll see how it can help simplify the grades example from the linear regression lecture. Figure 10.6 shows the results of applying PCA to the predictor variables in this example. Note the correlations between the PC scores are all zero; the general property of PCA already mentioned. However, the correlations between the PC scores and the Grade are non-zero.

We can regress the Grade y on the principal component scores $t^{(1)}, t^{(2)} \dots$:

$$y = \hat{\beta}_0 + \hat{\beta}_1 t^{(1)} + \hat{\beta}_2 t^{(2)} + \dots \quad (10.7)$$

When we regress on all 4 PC scores, we get exactly the same predictions and coefficient of determination as we do for regressing on all variables (Table 10.3). This makes sense, since by keeping all 4 components we have not lost any information about the data. It is more surprising that the coefficient of determination with if we regress on only the first two PC scores is almost as high. Furthermore, the adjusted coefficient of determination is actually higher when regression on the first two principal components, due to there being fewer variables. There is no combination of any two of the original variables that gives as high a coefficient of determination.

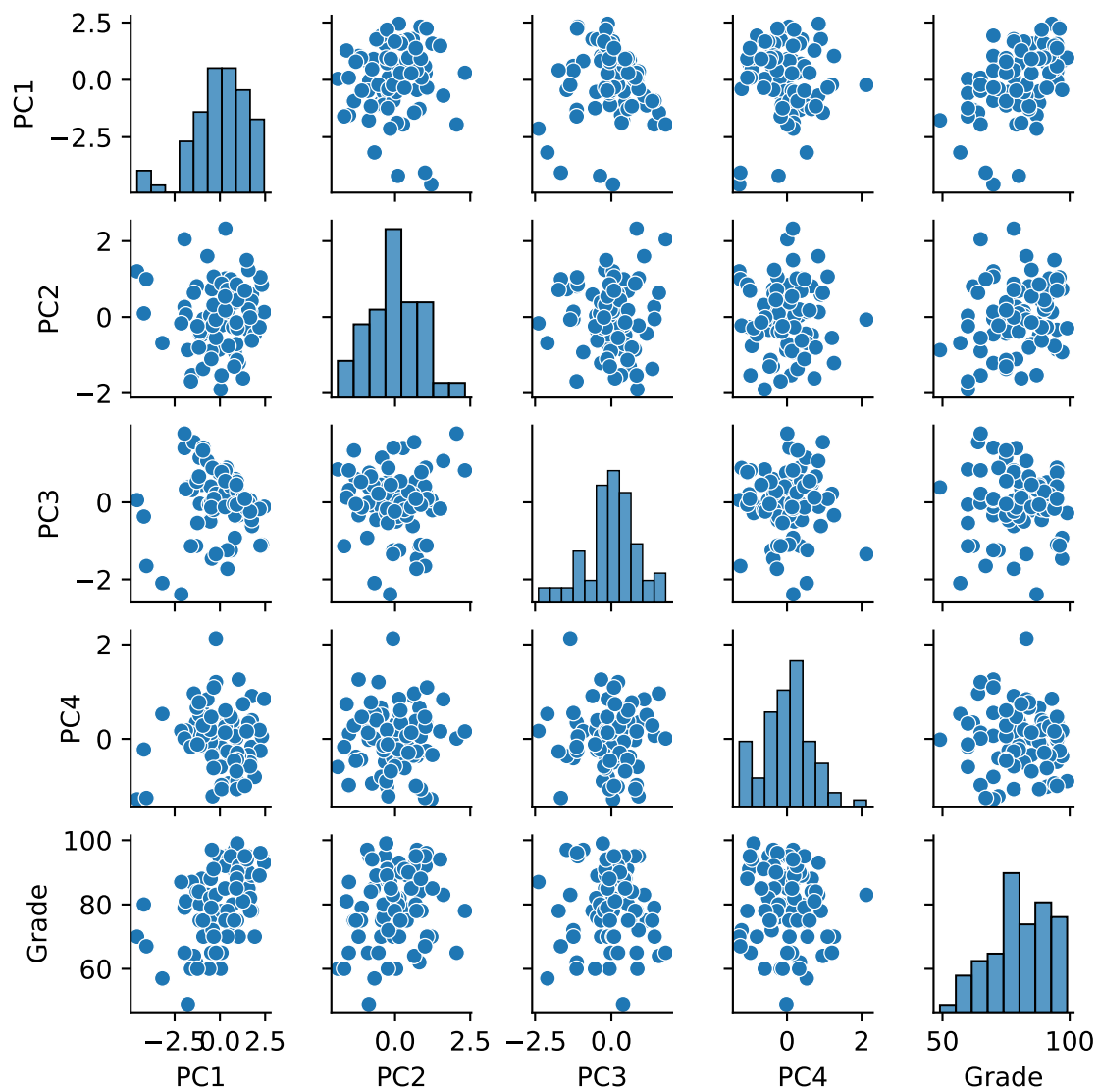


Figure 10.6: PCA scores of predictor variables in Grades example (see Multiple regression lecture notes).

Table 10.3: Coefficient of determination and adjusted coefficient of determination for regression of grades on original variables and on 2 or 4 PC scores.

	4 Original variables	4 PC scores	2 PC scores
R^2	0.289	0.289	0.282
R_a^2	0.251	0.251	0.263

This example demonstrates that PCA can be a useful preprocessing step for regression, by decorrelating the variables.

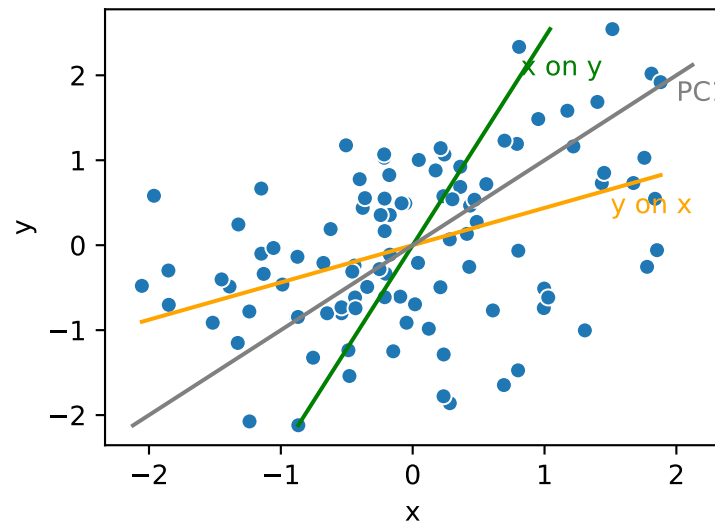


Figure 10.7: Regression of y on x , x on y and the first principal component of data with a correlation coefficient $r = 0.5$.

PCA and linear regression lines Thinking back to linear regression, we remember the distinction between the regression lines of y on x and x on y . In two dimensions there is now a 3rd line: the first principal component. As shown in Figure 10.7, this line goes right between the regression lines, and is probably what you would draw if asked to sketch the “line of best fit” to the data is. In fact, it *is* a line of best fit – it’s the line that minimises the sum of the squared distances from the data points to the line, rather than minimising the error in predicting y or x .

10.4 Derivation of PCA

Overview of derivation Here are the steps we’ll take in our derivation:

1. Define variance along the original axes
2. Project data onto rotated axes
3. Compute variance in these axes
4. Find direction of the axis that maximises variance of data projected onto it (1st principal component, PC1)
5. Interpret
6. Find the 2nd principal component (PC2)
7. Quantify what is lost by dimensionality reduction

This list may seem overwhelming, but it actually boils down to about 4 lines of code (assuming some helper functions), shown in Listing 10.1.

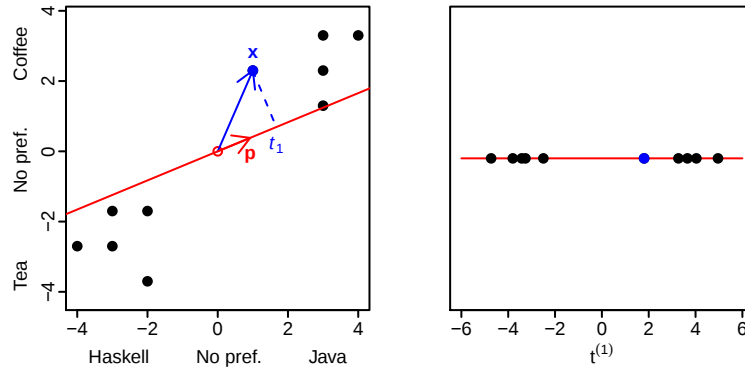
Listing 10.1: Listing of PCA. We are assuming the existence of helperfunctions `standardize()` and `sort_eigenvalues()`.

```
import numpy as np

def standardize(X)...

def sort_eigenvalues(lambda, P)...

def pca(X):
    """Given a data matrix X with n rows and D columns,
        return principal components (P) and
        eigenvalues (lambda)"""
    # Standardise the data X
    Z = standardize(X)
    # Compute the covariance matrix S
    S = np.cov(Z)
    # Find unsorted eigenvectors (P) and eigenvalues (lambda)
    lambdas, P = np.linalg.eig(S)
    # Sort the eigenvectors and eigenvalues
    # in order of largest eigenvalues
    lambdas, P = sort_eigenvalues(lambdas, P)
    # The eigenvalues (lambdas) are proportional
    # to the amount of variance explained
    for i in range(len(lambdas)):
        print('PC' + str(i+1) + ' explains ' +
              str(round((lambdas[i] / np.sum(lambdas))*100)) +
              '% of the variance.')
    return(lambdas, P)
```

Figure 10.8: Projection of a data point x onto a unit vector p .

Step 1: Defining variance along original axes We've already met a lot of the mathematical machinery we need in the multiple regression topic. We'll assume now that we have D variables $x^{(1)}, \dots, x^{(D)}$, and that we have transformed them into standardised versions of the variables $z^{(1)}, \dots, z^{(D)}$. We can arrange these standardised variables in an $n \times D$ data matrix,

$$\mathbf{Z} = \begin{pmatrix} z_{11} & \cdots & z_{1D} \\ \vdots & & \vdots \\ z_{n1} & \cdots & z_{nD} \end{pmatrix} = \begin{pmatrix} \mathbf{z}^{(1)} & \cdots & \mathbf{z}^{(D)} \end{pmatrix} = \begin{pmatrix} \mathbf{z}_1^T \\ \vdots \\ \mathbf{z}_n^T \end{pmatrix} \quad (10.8)$$

which it can be helpful to write in terms of the $D \times n \times 1$ vectors representing all the standardised data in each dimension ($\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(D)}$), or as the transposes of the $n \times D \times 1$ vectors ($\mathbf{z}_1^T, \dots, \mathbf{z}_n^T$) representing each standardised data point.

We've also met the covariance matrix, the $D \times D$ matrix that's derived from the data matrix:

$$\mathbf{S} = \begin{pmatrix} s_{11} & \cdots & s_{1D} \\ \vdots & & \vdots \\ s_{D1} & \cdots & s_{DD} \end{pmatrix} = \frac{1}{n-1} \mathbf{Z}^T \mathbf{Z} \quad (10.9)$$

The variance in the original axes is s_{11} , s_{22} etc. The covariance matrix in our toy example is:

$$\mathbf{S} = \begin{pmatrix} 9.7 & 8.0 \\ 8.0 & 7.7 \end{pmatrix}$$

Step 2: Project data onto a new axis We'll define a new axis by the **unit vector p** (Figure 10.8). The projection of a data point $\mathbf{z}_i$ onto this axis (its **component score**) is (as per Equation 10.1)

$$t_i = \mathbf{p}^T \mathbf{z}_i = p_1 z_{i1} + p_2 z_{i2} \quad (10.10)$$

At the moment, the vector p could be in *any* direction we choose.

Step 3: Compute variance in these axes The definition of the variance of the component scores t_i is:

$$s_t^2 = \frac{1}{n-1} \sum_{i=1}^n (t_i - \bar{t})^2 \quad (10.11)$$

We use Equation 10.10 to find the mean of the component scores:

$$\begin{aligned} \bar{t} &= \frac{1}{n} \sum_{i=1}^n t_i = \frac{1}{n} \sum_{i=1}^n (p_1 z_{i1} + p_2 z_{i2}) \\ &= p_1 \frac{1}{n} \sum_{i=1}^n z_{i1} + p_2 \frac{1}{n} \sum_{i=1}^n z_{i2} \\ &= p_1 \bar{z}^{(1)} + p_2 \bar{z}^{(2)} \\ &= 0 \quad \text{because the means of the standardised variables are 0} \end{aligned} \quad (10.12)$$

When we substitute $\bar{t} = 0$ and Equation 10.10 into Equation 10.11 we get the following:

$$\begin{aligned} s_t^2 &= \frac{1}{n-1} \sum_{i=1}^n (p_1 z_{i1} + p_2 z_{i2})^2 \\ &= \frac{1}{n-1} \begin{pmatrix} p_1 & p_2 \end{pmatrix} \begin{pmatrix} \sum_i z_{i1}^2 & \sum_i z_{i1} z_{i2} \\ \sum_i z_{i2} z_{i1} & \sum_i z_{i2}^2 \end{pmatrix} \begin{pmatrix} p_1 \\ p_2 \end{pmatrix} \\ &= \mathbf{p}^T \mathbf{S} \mathbf{p} \end{aligned} \quad (10.13)$$

Our old friend the covariance matrix has reappeared. Although we've demonstrated the derivation of this equation in 2 dimensions, the equation is still valid in D dimensions – the unit vector $\mathbf{p}$ would have D dimensions, and the covariance matrix $\mathbf{S}$ would be a $D \times D$ matrix.

Step 4: Find direction of axis that maximises variance of data projected onto it (1st principal component, PC1) Equation 10.13 gives the variance in the component scores for any direction of $\mathbf{p}$. We now want to find the direction of $\mathbf{p}$ that maximises that variance, under the constraint that $\mathbf{p}$ is of unit length, so $|\mathbf{p}| = 1$.

This is a constrained optimisation problem, which we can solve using Lagrange multipliers and differentiation. We won't show the details here, but the result is the following equation:

$$\mathbf{S} \mathbf{p} = \lambda \mathbf{p}$$

Hopefully you recognise this equation from a linear algebra course. It has two solutions, with different values of λ :

1. $\lambda = \lambda_1$, $\mathbf{p} = \mathbf{e}_1$, where λ_1 is the biggest **eigenvalue** of $\mathbf{S}$ and $\mathbf{e}_1$ is the associated **eigenvector**
2. $\lambda = \lambda_2$, $\mathbf{p} = \mathbf{e}_2$, where λ_2 is the second biggest **eigenvalue** of $\mathbf{S}$ and $\mathbf{e}_2$ is the associated **eigenvector**

we choose the **first principal component** $\mathbf{p}_1$ to be the eigenvector $\mathbf{e}_1$ with the largest eigenvalue λ_1 . λ_1 is the variance of the 1st component scores $\mathbf{t}^{(1)}$.

Step 5: Interpret Finding the eigenvalues and eigenvectors for our 2D example (described in [Principle of finding principal components](#)), we find that the first principal component is:

$$\mathbf{p}_1 = \begin{pmatrix} 0.75 \\ 0.66 \end{pmatrix} \begin{matrix} \text{Java} \\ \text{Coffee} \end{matrix} \quad \lambda_1 = s_{t^{(1)}}^2 = \mathbf{p}_1^T \mathbf{S} \mathbf{p}_1 = 16.5$$

The first component score $\mathbf{t}^{(1)}$ is the “Java-coffeeiness” of a student.

Step 6: Find the 2nd principal component (PC2) In 2D our job is already done, since there is only one direction perpendicular to $\mathbf{p}_1$, and eigenvectors (and therefore principal components) are always orthogonal to each other. It's the other eigenvector of $\mathbf{S}$, $\mathbf{p}_2 = \mathbf{e}_2$, with eigenvalue λ_2 , which is the variance of the 2nd component scores $\mathbf{t}^{(2)}$.

In D dimensions, the principal components are the D eigenvectors of the $D \times D$ matrix $\mathbf{S}$. It's helpful to introduce more matrix notation here. We arrange the principal components in the **rotation matrix**:

$$\mathbf{P} = \begin{pmatrix} \mathbf{p}_1 & \mathbf{p}_2 \end{pmatrix}$$

Step 7: Quantify what is lost by dimensionality reduction We can reverse the transformation from the scores to the original data

$$\mathbf{Z} = \mathbf{T} \mathbf{P}^T$$

If we drop the 2nd PC from $\mathbf{P}$ and the 2nd PC scores from $\mathbf{T}$, we can reconstruct a 1-dimensional version of the original data:

$$\hat{\mathbf{Z}}^{(1)} = \mathbf{t}^{(1)} \mathbf{p}_1^T$$

We can see (Figure 10.9) that the first principal component (the “Java-coffeeiness”) score of a student tells us **a lot** about them – but how much? Consider the **total variance**, the sum of the variances of the data:

$$\sum_{i=1}^D s_i^2 = \sum_{i=1}^D \lambda_i$$

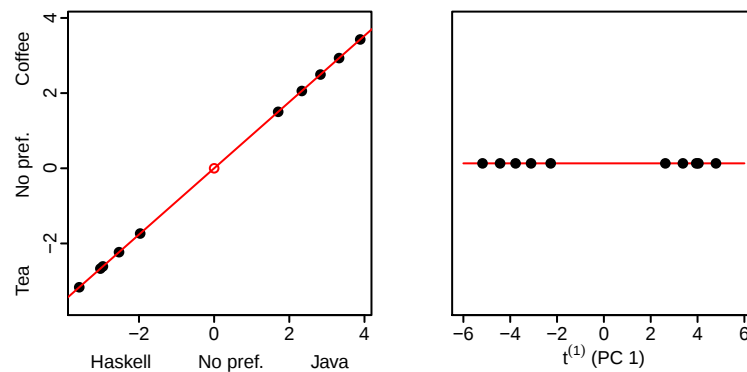


Figure 10.9: Projection into the original space.

It is equal to the sum of the eigenvalues of the covariance matrix. Thus the fraction of the total variance “explained” by the i th principal component is:

$$\frac{\lambda_i}{\sum_{j=1}^D \lambda_j}$$

In our toy example,

$$\frac{\lambda_1}{\lambda_1 + \lambda_2} = \frac{16.5}{16.5 + 0.61} = 96\%$$

Thus we can now be more precise about how much the first principal component (the “Java-coffeeiness”) score of a student tells us about them: 96% of the variance.

Related Python Lab: PCA

<https://github.com/Inf2-FDS/FDS-S1-08-pca>

In this lab you will implement PCA from scratch and compare your outcomes to the standard scikit-learn PCA. By the end of the lab you should be able to:

- explain why it is important to standardize data prior to PCA,
- implement PCA from scratch
- use the sklearn library to get the principal components from a dataset.

Part III

Introduction to Machine Learning

Chapter 11

Supervised learning: Classification with Nearest neighbours

 Further reading (not examinable)

Hastie et al. (2009) *The elements of statistical learning*, pp 463–471

11.1 Classification

The classification problem Suppose a bank has data on previous customers it has given loans to, including variables such as their income, housing status and employment status. Each of these sets of variables – also referred to as **feature vectors** – has a **label**, indicating whether the customer did or didn't pay back their loan. The bank might want to predict whether a new customer will be able pay back a bank loan from their features, i.e. to predict whether they belong to the class of customers who paid or the class of customers who didn't pay. This is an example of **classification**, which we define as the problem of predicting the correct category label ("class") for an unlabelled input feature vector.

Supervised and unsupervised learning Classification is an example of a supervised learning process. In a **supervised learning** process, there is a **training set** of data in which each data item has a number of **features** and a known **label**. The goal of supervised learning is to predict the label of an item that has not been previously seen from its features. In contrast, in **unsupervised learning** processes, the training set does not contain any labels, and the goal is to learn something about the structure of the data. We cover unsupervised learning in the chapter on 13.

Visualising the classification problem To visualise the classification problem, we'll use a toy example: the fruit data set, collected by Iain Murray (Murray, 2006). He bought pieces of fruit and measured their height and width (features) and noted the type of fruit (the label). Figure 11.1 visualises the data. In the context of the fruit, the classification problem is using this dataset to build a machine to predict the class of a piece of unidentified fruit automatically just by measuring its width and height. We will refer to the feature vector of this unidentified fruit as the **test point**.

Definition of a classifier To solve a specific classification problem, we construct a **classifier**. A classifier is a function that takes a feature vector $\mathbf{x}$ and returns a class c where c is a member of a set $\mathcal{C}$. In principle, we can construct a classifier in any way we want to, as long as it matches this definition. Regardless of how the classifier is constructed, it will have two important properties: **decision boundaries** and **classification errors**.

Decision boundaries Consider the problem of determining apples from pears. In this example case we have two features for each piece of fruit: its circumference (at the widest point) and its height, each measured in cm. Apples are 'more spherical' than pears which tend to be longer than they are wide. But some pears are relatively short, and some apples are taller than expected. In this case we have an input vector of the form

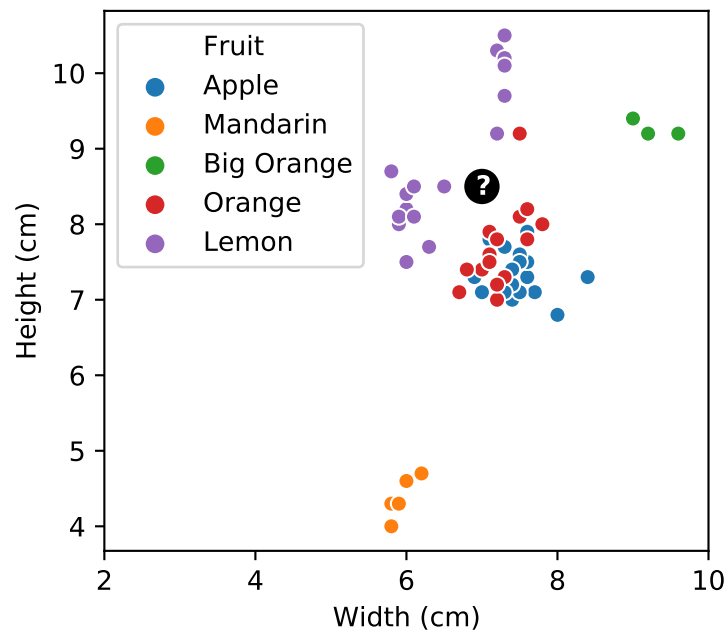


Figure 11.1: The supervised learning problem, as applied to fruit. We are given the labels of the fruit with various widths and heights. We are then presented with an unknown piece of fruit with given width and height (the test point, represented by the question mark). The task is then to predict what type of fruit it is.

$\mathbf{x} = (x^{(1)}, x^{(2)})^T$, where $x^{(1)}$ is the circumference and $x^{(2)}$ the height. The class c can take two values A or P (standing for apples and pears).

We have a set of training data: height and circumference measurements, along with class labels. In Figure 11.2 we plot this two-dimensional training data for the two classes. We can see that it is not possible to draw a straight line to separate the two classes.

We now have three new, unlabelled examples which we would like to classify (represented as stars in Figure 11.3):

- (16, 10): all the training data in the region of this point is classified as P , so we classify this point as P .
- (19, 6): looking at the training data it seems obvious to class this as A .
- (18, 7): it's not obvious in which class this example should be classified; the feature vector gives us evidence whether we have an apple or a pear, but does not enable us to make an unambiguous classification.

We can draw a straight line such that one side of it corresponds to one class and the other side to the other – as in the three possible lines shown in Figure 11.2. Such a line is called a **decision boundary**; if the data was three-dimensional, then the decision boundary would be defined by a plane. For one-dimensional data, a decision boundary can simply be a point on the real line. Intuitively it seems possible to find an optimal decision boundary, based on minimising the number of misclassifications in the training set.

Constructing classifiers using supervised learning To construct the classifier automatically we need:

1. a set of training data containing feature vectors and their labels
2. an algorithm that we train using the training data
3. **hyperparameters**, numbers that control how the algorithm learns and predicts

This is referred to as **supervised learning**, since the label for a training vector acts as supervision for the classifier when it is learning from training data.

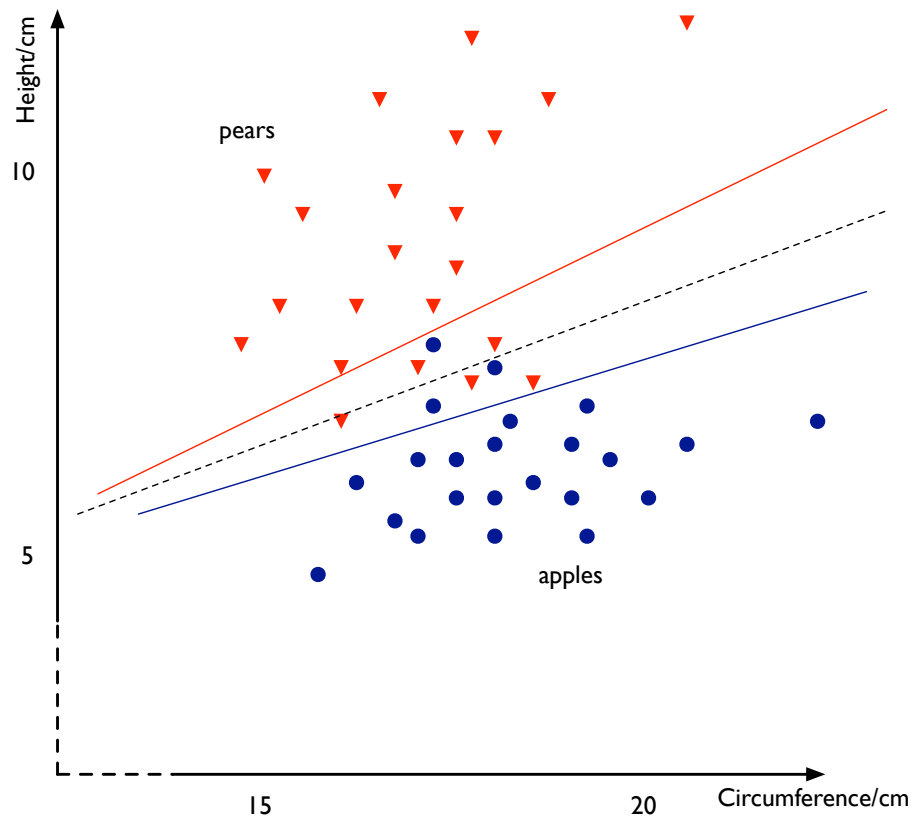


Figure 11.2: Training data for apples and pears. It is not possible to draw a straight line to perfectly separate the classes in this feature space. Three possible lines are drawn, but each results in a number of misclassifications.

11.2 Nearest neighbour classification

Principle of nearest neighbour classification Nearest neighbour classification (or one-nearest neighbour classification to be precise) has a very simple basis: to classify a test item, find the item in the training set which is closest and assign the test item to the same class as the selected training item. If there happens to be an identical item in the training set then it makes sense to assign the test item to the same class. Otherwise, the class of the member in the training set which is most similar to the test item is our best guess. We use a distance measure (e.g., Euclidean distance) to determine similarity. If we have a representation for which the distance measure is a reasonable measure of similarity, then the nearest neighbour method will work well.

Decision boundaries for nearest neighbour classification What do the decision boundaries look like for nearest neighbour classification? Each training data point defines a region around it; test points within this region will be classified to the same class as the training data point. These regions are illustrated for a simple case in Figure 11.4, where the boundaries of regions are shown as dotted lines. Each boundary is given as the perpendicular bisector of the line segment between the two corresponding data points. This partitioning formed by a set of data points is sometimes called a **Voronoi diagram** or a **Voronoi tessellation**.

Now we assume that each data point belongs to either of two classes, say, red or blue. To obtain the decision boundary we combine those boundaries which are between regions of different classes, as illustrated in Figure 11.5. The resultant boundary is referred to as being **piecewise linear**. Figure 11.6 shows the decision boundary and decision regions in the case of three classes.

Application of 1-nearest neighbour to a real dataset Figure 11.7 shows 1 nearest-neighbour classification applied to the fruit dataset. Note that we've standardised the variables, so that the data spreads out roughly equally in both directions. In common with other distance-based methods, we would like the results of clustering to be independent of the units we measure the variables in. It can be seen that the decision

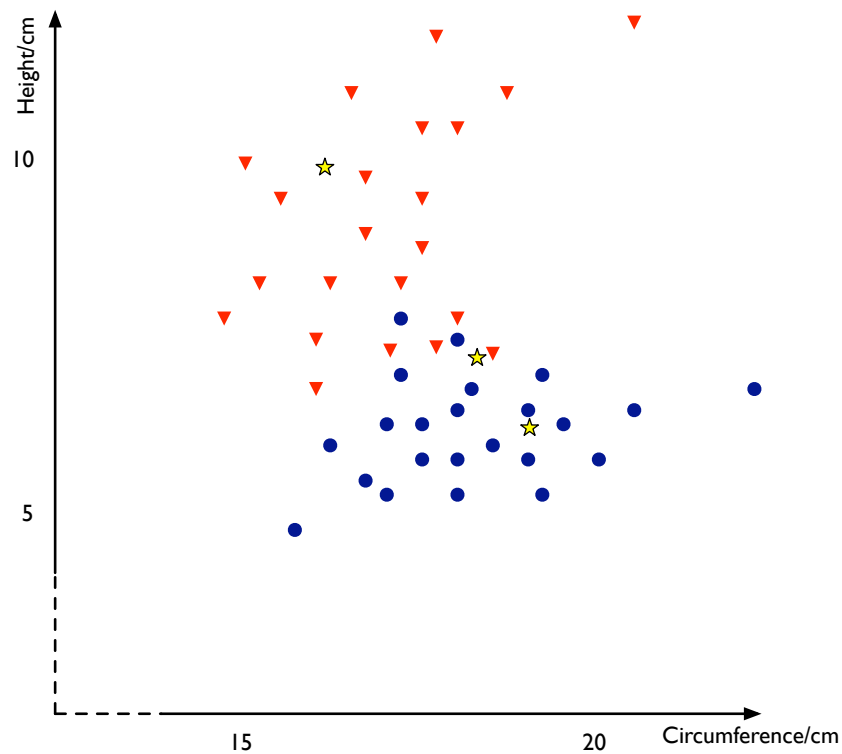


Figure 11.3: The training data for apples (blue circles) and pears (red triangles), together with three test points (yellow stars).

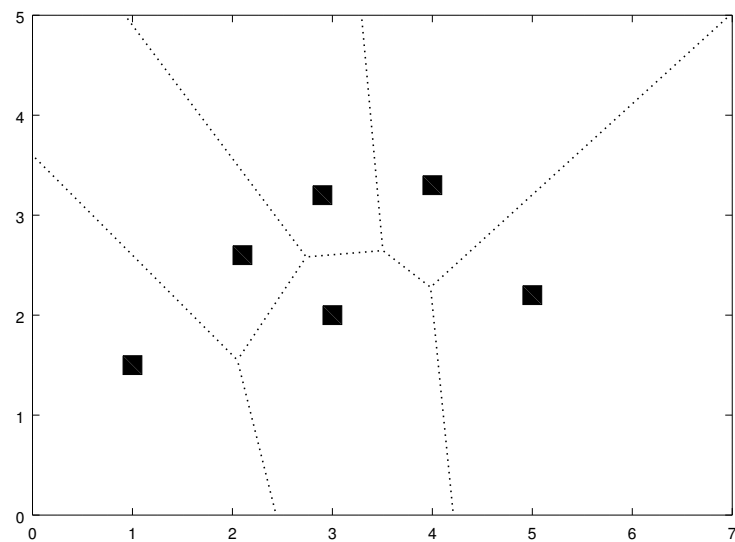


Figure 11.4: Decision boundaries by 1-NN for data points of distinct classes from each other

boundary is quite complex, with islands of apple amongst the oranges. We'll explore in the next section if this might be a problem.

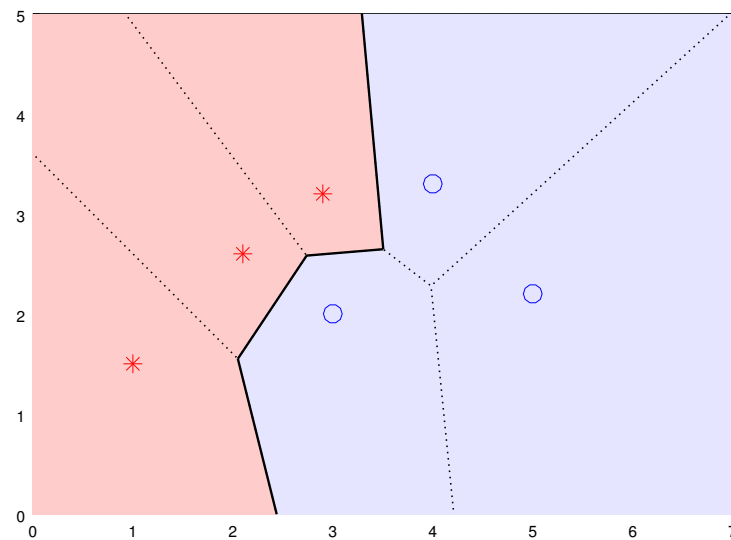


Figure 11.5: Decision boundary and decision regions for a 1-nearest neighbour classifier for a training dataset of two classes, where training samples of one class are shown with '*' in red, those of the other class are shown with 'o' in blue. The Euclidean distance is used as the distance measure in this example.

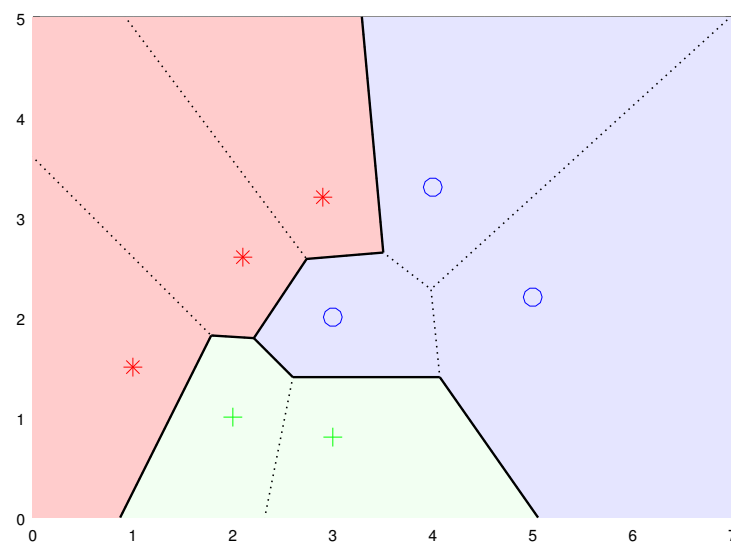


Figure 11.6: Decision boundary and decision regions for a 1-nearest neighbour classifier for three classes.

11.3 Evaluation

Classification error rate Having constructed a classifier, we would like to evaluate how well it works. One way to quantify how well a classifier is working is the number of items that it misclassifies – i.e., the number of times it assigns a class label $\hat{c}_i$ different from the true class label c_i . The classification error is often expressed as the percentage of the total number of items that is misclassified, the **classification error rate**, sometimes referred to as the “error rate”.

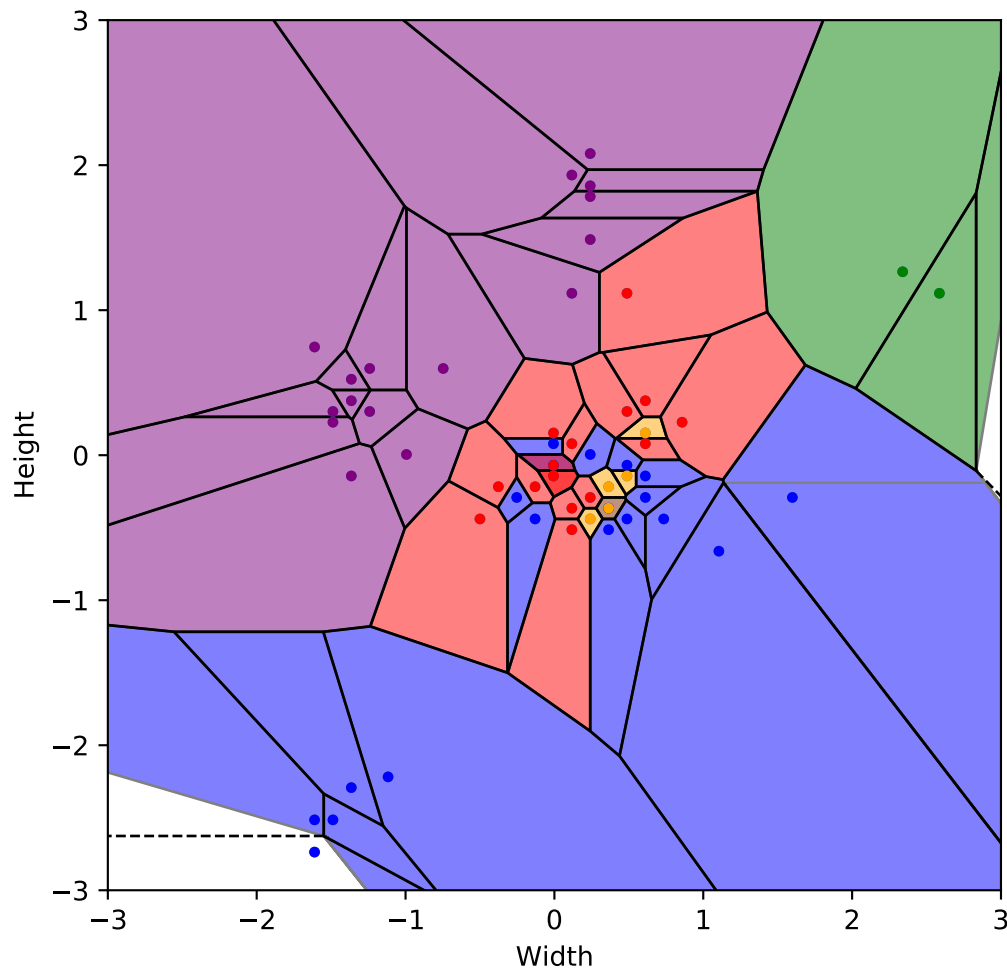


Figure 11.7: Decision regions for one nearest neighbour classification applied to the fruit dataset. The variables have been standardised to make the scales on both axes similar. Some regions are darker shade of blue or red. This indicates that there are 2 points labelled with “apple” or “orange” in the dataset with the same features. There is one region that is purple, amongst the blue and red region. There are two data points corresponding to this region with identical coordinates, one labelled with orange and one with apple. Colour indicates decision region for each type of fruit: Apple (blue), Mandarin (orange), Big Orange (green), Orange (red), Lemon (purple).

Classification error rate for one-nearest neighbour classification For one-nearest neighbour classification, the error rate when we consider members of the training set is 0, since the closest point in the training set to a member of the training set is itself¹.

Evaluating generalisation to unseen data This sounds very promising, until we remember that the job of the classifier is to classify data points that we haven't seen before. It may be that the classifier will not **generalise** to data that haven't seen. In order to estimate how well the classifier generalises, we can split our original dataset into a training set and a **testing set**. The training and testing sets are mutually exclusive, and a typical split might be 70% for training and 30% for testing. We train the classifier using the training set, and then evaluate the performance using the testing set. **We are not allowed to use the test set to train the classifier** – otherwise our estimate of its performance on the test set will be too optimistic.

In summary, the **training set error rate** is the percentage of misclassifications that the classifier makes on the training set after the learning algorithm has been applied. The **test set error rate** refers to errors made by the classifier on the testing set.

Training and testing set notation We'll use the notation: $\mathbf{x} = (x^{(1)}, x^{(2)}, \dots, x^{(D)})^T$ to denote a D -dimensional (input) feature vector, which has class label c . The **training set** is a set of n feature vectors and their class labels; the i 'th training item consists of a feature vector $\mathbf{x}_i$ and its class label c_i . The j 'th element of the i 'th feature vector is written x_{ij} .

¹Unless we have two data points with exactly the same features and different labels.

Chapter 12

k-Nearest neighbour classification and regression, setting hyperparameters, metrics and cross-validation

12.1 *k*-Nearest neighbour classification

Principle of *k*-nearest neighbour classification Rather than just using the single closest point, the *k*-nearest neighbour approach looks at the *k* points in the training set that are closest to the test point; the test point is classified according to the class to which the majority of the *k*-nearest neighbours belong.

Figure 12.1 repeats the apples and pears example, with the three test points at (16, 10), (19, 6), and (18, 7). For the first two items, the value of *k* is not really important: (16, 10) is classified as pear and (19, 6) is classified as apple, no matter how many nearest neighbours are considered. However, the third example above, (18, 7), is ambiguous, and this is reflected in the sensitivity of the classifier to the value of *k*:

- 1-nearest: classified as pear
- 2-nearest: tie (one apple and one pear are nearest neighbours). In this case, we could choose randomly between the two classes. Another option strategy would be to take the 1-nearest neighbour or the 3-nearest neighbour classification.
- 3-nearest: classified as apple
- 5-nearest: classified as pear
- 9-nearest: classified as apple

***k*-nearest neighbour algorithm** We can write the *k*-nearest neighbour algorithm precisely as follows, where and *d* is the distance metric (typically the Euclidean distance):

- For an unseen example *x*:
 - Compute the *n* distances $d_i = d(\mathbf{x}, \mathbf{x}_i)$ between *x* and the features of each training example $\mathbf{x}_i$, $i \in 1, \dots, n$.
 - Sort the distances from lowest to highest and find the indices $i_1, \dots, i_k$ of the *k* lowest values of d_i
 - Find the classes that belong to the closest points, i.e. $c_{i_1}, \dots, c_{i_k}$
 - Each of these represents a vote for a class. Count the votes for each class and return the one with the largest number.
 - If there is a tie, choose randomly, or look at the *k* + 1th neighbour to resolve the tie.

Decision boundaries produced by *k*-NN classification *k*-nearest neighbour classifiers make their decisions on the basis of local information. Rather than trying to draw a linear decision boundary across the whole space (as in Figure 11.2), *k*-nearest neighbour techniques make decisions based on a few local points. As such they can be quite susceptible to noise, especially if *k* is small: a small shift in the location of a test point can result in a different classification since a new point from the training dataset becomes the nearest neighbour. As *k* becomes larger, the classification decision boundary becomes smoother since

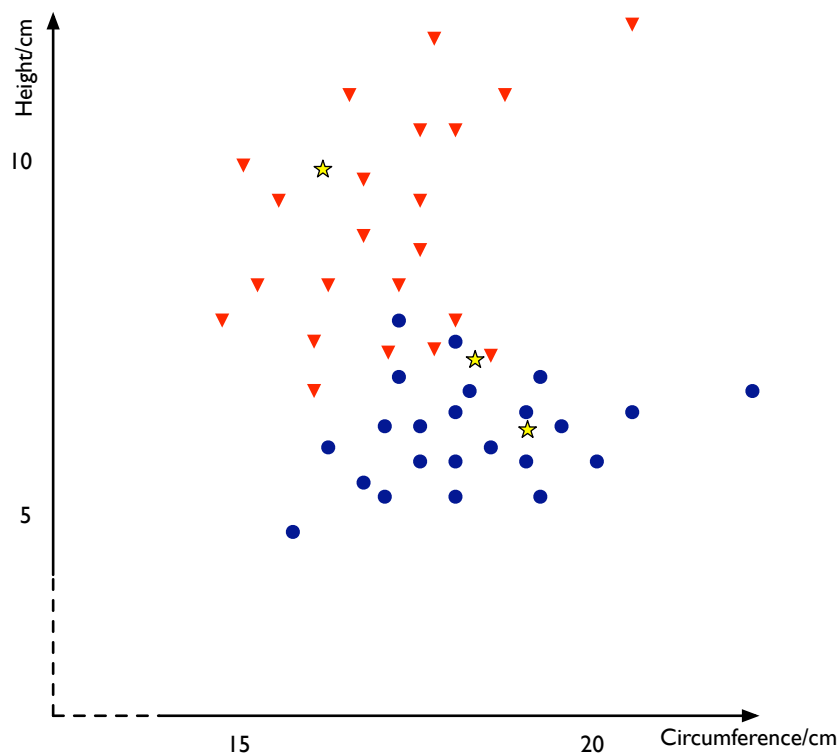


Figure 12.1: The training data for apples (blue circles) and pears (red triangles), together with three test points (yellow stars).

several training points contribute to the classification decision. Figure 12.2 illustrates the decision regions for various values of k for the 5-fruit example introduced in the previous chapter (raw data in Figure 11.1 and 1-nearest-neighbour classification in Figure 11.7).

Generalisation and regularisation In Figure 12.2 for low values of k the boundary between apples (blue) and oranges (red) is “noisy”: a small shift in the height and width of the fruit will lead to the classification training. The trained classifier is very flexible and therefore over-sensitive to the data it’s been training on, and we say that it is exhibiting **over-fitting** and **under-generalisation**.

As k increases, the decision boundaries get smoother, and we might think that the results will be exhibit better **generalisation** to unseen examples. As k increases further there could also the problem of **over-generalisation** or **under-fitting**. This problem isn’t seen clearly in Figure 12.2. However, if we made k very large, we would end up classifying everything as the fruit with the largest number of examples. Another example of over-generalisation might be the linear decision boundaries in Figure 11.2 in [Supervised learning: Classification with Nearest neighbours](#).

We can see that the decision regions with higher k in Figure 12.2 appear more regular. The process of changing the behaviour of a classifier so that it produces more regular or smoother output is known as **regularisation** and k is sometimes referred to as a **regularisation parameter**.

Over- and under-fitting (and their counterparts under- and over-generalisation) are issues for other supervised learning methods, for example when extending multiple regression with extra features ([Interaction terms and nonlinear fits](#)). In general, supervised machine learning models have hyperparameters that act as regularisation parameters. We leave more detailed work on regularisation parameters for later courses.

Choosing k The value k is **hyperparameter**: a number that we can choose to get the best performance from the algorithm. Figure 12.3 shows the classification error rate for various values of k on the training set and the testing set. As k increases the error on the training set initially increases rapidly, as explained in the last section. The testing error decreases a little and then starts rising around $k = 9$, indicating that a somewhat larger k helps generalisation. Both testing and training error then increase.

This graph suggests that we can look at the error on the testing set to set k . **But this would break the rule of using the test data to train the classifier, since our choosing the best hyperparameter k is part of**

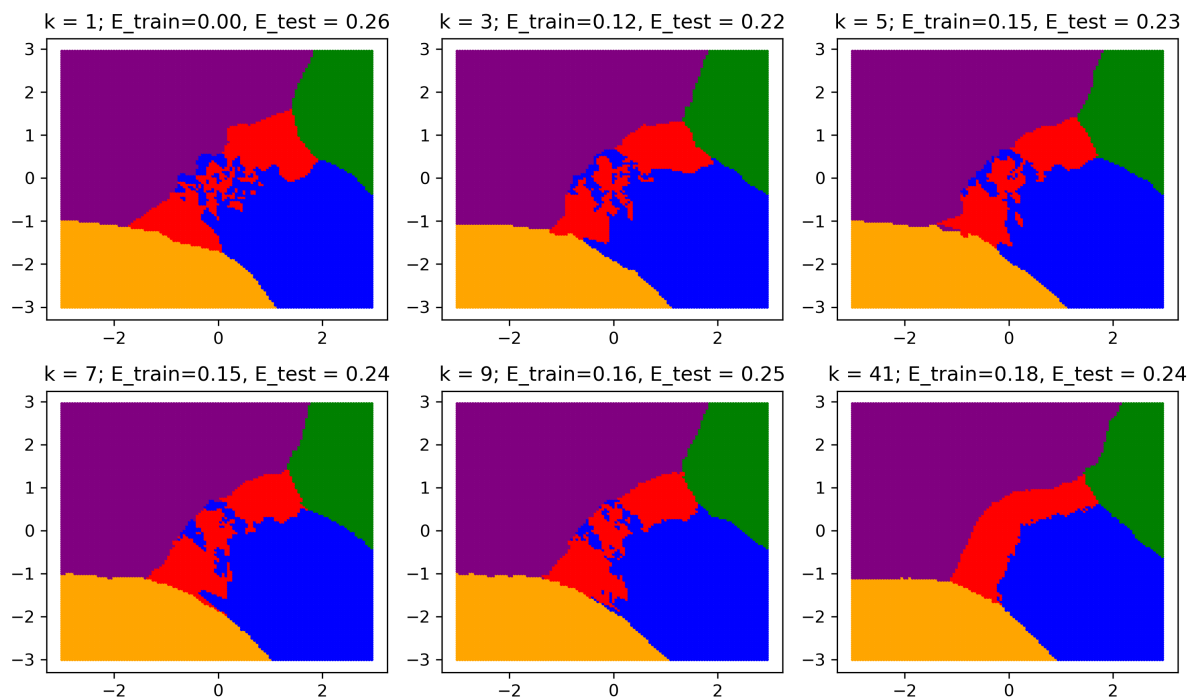


Figure 12.2: Decision regions, training error and testing error for various values of k applied to the 5-fruit data shown in Figure 11.1. Colour indicates decision region for each type of fruit: Apple (blue), Mandarin (orange), Big Orange (green), Orange (red), Lemon (purple).

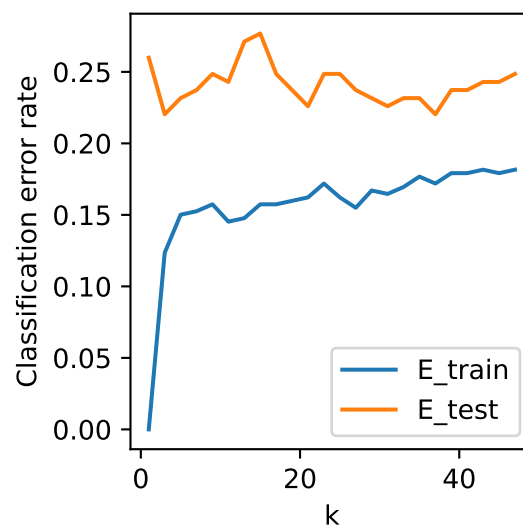


Figure 12.3: Classification error rate for various values of k .

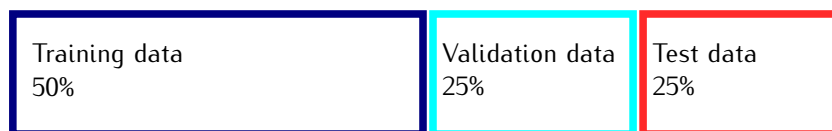


Figure 12.4: Typical training/validation/testing split in holdout cross-validation.

the training process. We have really been using the test data as **validation data**, that is, data used to help us validate our choice of hyperparameter by estimating what the error would be on an independent set of data drawn from the same population as the data available. This technique of using separate datasets for training and testing to predict the error is referred to as **cross-validation**.

Thus, we need to divide our dataset into 3 parts (Figure 12.4):

- Training data (about 50%): used to train the classifier for any particular value of k .
- Validation data (about 25%): used to compare performance of the trained classifier for different values of k .
- Testing data (about 25%): used to **report** the performance of the trained classifier with the one value of k that we have chosen.

The precise fractions of data are not crucial. However, it is important that the testing data is *only* used to report the performance, *not* to choose hyperparameters. A poor test score is probably an indication that the classifier will perform poorly on real world data. Here we have taken all the non-testing data available, and “held out” some of the data to form the validation set, leading to the name **holdout cross-validation** for this method. A more sophisticated way of undertaking validation is K -fold cross-validation (see later).

Computational efficiency of k -nearest neighbour classification k -nearest neighbour is very efficient at training time, since training simply consists of storing the training set.¹ Testing is much slower, since, in the simplest implementation, it would involve measuring the distance between the test point and every training point, which can make k -nearest neighbour impractical for large datasets.

Improving the efficiency of k -NN It is sometimes possible to store the training dataset in a data structure that enables the nearest neighbours to be found efficiently, without needing to compare with every training data point (in the average case). One such data structure is the k -d tree. However, these approaches are usually only fast in practice with fairly low-dimensional feature vectors, or if approximations are made.

12.2 Metrics

Accuracy In the chapter on [Supervised learning: Classification with Nearest neighbours](#), we introduced the classification error rate, the number of items misclassified as a fraction of the total number of items. We define **classification accuracy** as one minus the error rate, i.e. one minus the number of items misclassified divided by the number of items. When the error rate is zero, the classification accuracy is 1 or, equivalently, 100%.

Accuracy and unbalanced classes Accuracy seems to make sense as a metric – the fewer errors, the better the classifier. However, it can appear misleadingly high when there is a large difference in the number of items in each class, which we call **unbalanced classes**. For example, suppose we have a dataset containing 95% apples, 3% pears and 2% oranges. When we split into training and test sets, the split will be 95%/3%/2% in both the training and test sets. We could devise a classifier that classifies *any* item as an apple, regardless of its height and width. This classifier would have an accuracy of 95%. This type of dummy classifier is called a **baseline classifier**, since its performance sets a baseline against which to compare more “intelligent” classifiers.

¹In practice, responsible machine learning practitioners will try out different choices of k , and different distance measures, possibly optimising free parameters of a distance measure. Then training requires testing different choices, and becomes expensive.

		Predicted class	
		P	N
Actual class	P	TP	FN
	N	FP	TN

Figure 12.5: Confusion matrix. TP – number of true positives; FN – number of false negatives; FP – number of false positives; and TN – number of true negatives.

Sensitivity and selectivity as alternative metrics in two-class problems Classification problems often have two classes, for example someone has or has not repaid a loan, or someone does or doesn't have an illness. In these two-class problems, we regard one class as the “positive” outcome and the other as the “negative” outcome. Confusingly the “positive” outcome is usually the case that we are searching for, which may often be a negative thing – think of testing “positive” for Covid-19.

Two-class problems allow us to introduce other metrics, or sometimes pairs of metrics, that avoid the misleading impression given by accuracy with unbalanced classes. One common pair of metrics are **sensitivity** and **selectivity**, defined as:

Sensitivity Fraction of positives classified as positive

Selectivity (also known as **specificity**) Fraction of negatives classified as negative

Suppose now that instead of classifying fruit, we are classifying whether someone is “positive” for Covid, on the basis of their symptoms. We will assume that 2% of people truly have Covid. The dummy classifier described has a high accuracy (how high?). However, it classifies all the positive cases as negative, so the sensitivity is 0%. Conversely, all the negative cases are classified as negative, so the selectivity is 100%. This gives us a much clearer picture of the performance of the classifier than accuracy alone. An ideal classifier would have 100% selectivity and sensitivity.

Confusion matrix The fullest picture of the performance of a classifier on a two-class problem can be gained by looking at the confusion matrix, which compares the actual class and the predicted class (Figure 12.5). The cells of the matrix are the number of items classified as “true positives”, “false positive” etc. Sometimes we normalise by dividing every cell by the total number of items, so that the sum of all the cells is 100%.

We can define metrics in terms of the cells of the confusion matrix. For example, we have

$$\text{Sensitivity} = \frac{TP}{TP + FN} \quad \text{Selectivity} = \frac{TN}{TN + FP} \quad (12.1)$$

Other metrics can be constructed from the numbers in the cells of the confusion matrix; the appropriate metric to use will depend on the application.

12.3 k-Nearest neighbour regression

Principle of k-NN regression In the section on [Regression as prediction](#) we introduced a nonparametric regression method that predicted the value of the response variable (or target variable) y from a predictor variable x based on the mean value of the values y_i corresponding to region of fixed width of the x -axis around x . One downside of this method was that if there is no data within the fixed region of x , no prediction is possible.

We can get around this problem by using the principle of k -Nearest Neighbours. Suppose that $\mathcal{N}_k(x)$ is the set of the k closest predictor values out of the set $\{x_1, \dots, x_n\}$ to x . Then we can predict the value of y from x by taking the mean of the points in the k -nearest neighbourhood:

$$\hat{y} = \frac{1}{k} \sum_{i \in \mathcal{N}_k(x)} y_i \quad (12.2)$$

Alternatively we can weight the points in inverse proportion to distance:

$$\hat{y} = \sum_{i \in \mathcal{N}_k(x)} \frac{\frac{1}{d(x, x_i)} y_i}{\sum_{j \in \mathcal{N}_k(x)} \frac{1}{d(x, x_j)}} \quad (12.3)$$

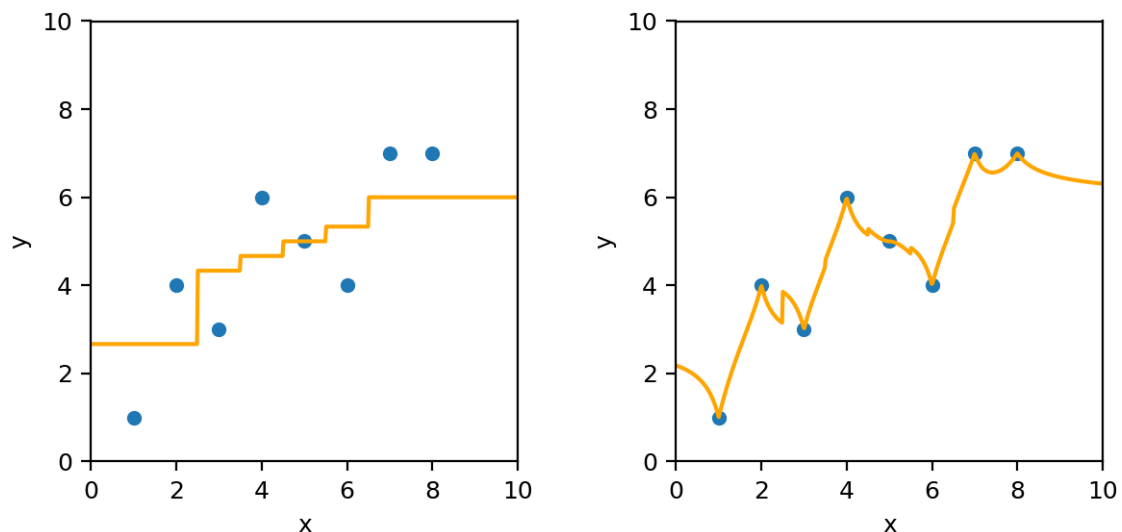


Figure 12.6: k -NN regression with $k = 3$ using equal weights (left) and weighting by inverse distance (right).

Application to data with one feature Figure 12.6 shows the predictions made by both versions of k -NN regression with $k = 3$ on some simple sample data. The mean method results in jumps in predictions, where the neighbourhood changes. Predictions in the inverse weighting method pass through every data point, but there are sharp discontinuities in the gradient.

Application to data with multiple features Equations 12.2 and (12.3) apply if x contains a set of features, i.e. is multidimensional, and we are trying to predict a scalar y for any value of x .

12.4 Limitations of supervised machine learning and cross-validation

Limitations and pitfalls of supervised machine learning The supervised machine learning paradigm can be summarised as:

1. Split data into training, validation and test sets
2. Use training and validation data to select
 - an algorithm
 - hyperparameters (if applicable)
3. Use test data to report performance of resulting predictor (classifier or regressor)

We've used k -nearest neighbours as our algorithm, but many other algorithms are available.

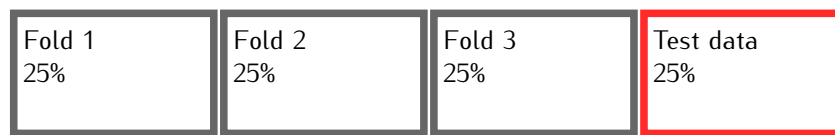
There are a number of issues when using this paradigm in the real world.

Data can change over time. For example typical Covid symptoms changed over the pandemic, so a classifier that worked well at the start of the pandemic wouldn't necessarily work well at the end of the pandemic. Perhaps it's been a particularly good year for apples in the year we've trained the fruit classifier, and they are larger then in other years.

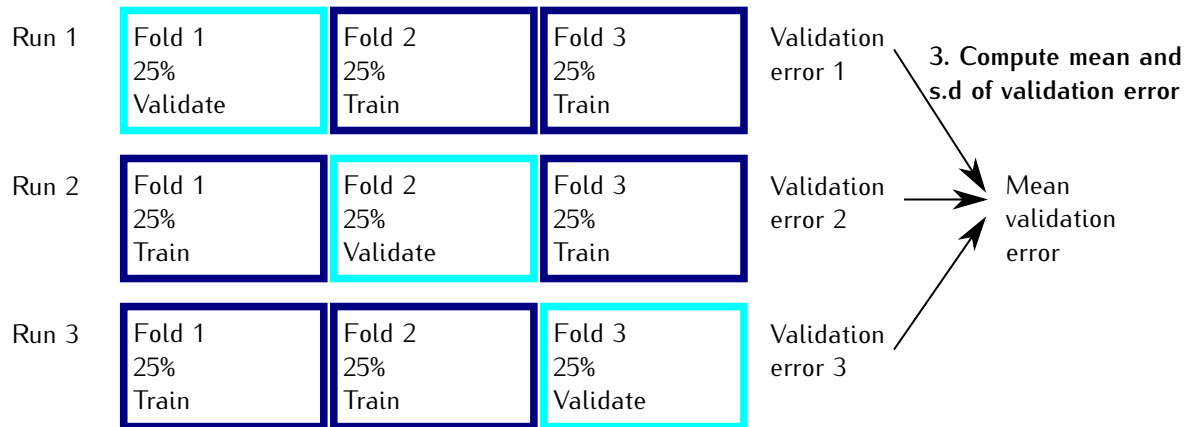
Selection of training, validation and test sets A related problem is that the training, validation and test sets should be drawn from the same distribution. The training, validation and test sets should be selected using random sampling, As discussed in the chapter on [Randomness, sampling and simulation](#), any form of non-random sample (for example taking the first 50% of data collected over time) could lead to the statistics of the training and test sets being different.

Limited amount of data Machine learning algorithms need more data to perform well. If you have 100 data points, there will be only 50 data points to use for training, 25 for validation and 25 for testing. As discussed in the chapter on [Randomness, sampling and simulation](#), these samples may not be representative of the wider population.

1. Split into testing data and non-testing data. Split non-testing data into 3 folds.



2. For each model and hyperparameter, train and compute validation error three times:



4. Pick hyperparameter with lowest mean validation error. Train on all non-testing data and report performance on test data.



Figure 12.7: K -fold cross-validation. See text for details.

K -fold cross-validation We've seen how we can split the data we're not using to report the final test result into training and validation data, and that this can be used to find the hyperparameter that gives the best performance. However, if we've a small amount of data, we're only using a third of the actual data that we're allowed to use to test the performance of each hyperparameter. The method of **K -fold cross-validation** allows us to use all the non-testing data for both training the classifier with each hyperparameter, and testing it.

Figure 12.7 gives an overview of how K -fold cross-validation works. There is a split into test data, and non-testing data, as previously. The difference is that the non-testing data is split into a number (here 3) of equally-sized **folds**, or blocks. For each value of the hyperparameter, the model is trained 3 times: first on the data in folds 2 and 3, then on the data in folds 1 and 3, and finally on the data in folds 1 and 2. Each trained model is tested on the data in the fold that wasn't used for training. The mean error from all three folds is then used as a *estimate* of what the test error will be in the final trained model. On the basis of this K -fold cross-validation, a hyperparameter is chosen – probably the one giving the lowest mean error. Finally, the classifier is trained again using the chosen hyperparameter on all the non-testing data.

In general, we don't have to have 3 folds. We can have a number K of folds, which gives rise to the term K -fold cross-validation. 5 and 10 are common values for K .

⚠ The meaning of K and k in nearest neighbours and cross-validation

It's unfortunate that the letter pronounced "K" is used for both cross-validation and nearest neighbours, as the letter means different things in the two contexts. In these notes we'll use uppercase K for K -fold cross-validation and lowercase k for k -nearest neighbours. However, this convention is not followed consistently beyond these notes.

Although the validation error gives an estimate of the expected test error, it is still necessary to compute the error on a separate test set. The estimate produced by cross-validation has an optimism bias, because

the data in each testing fold has been used to select the model hyperparameters.

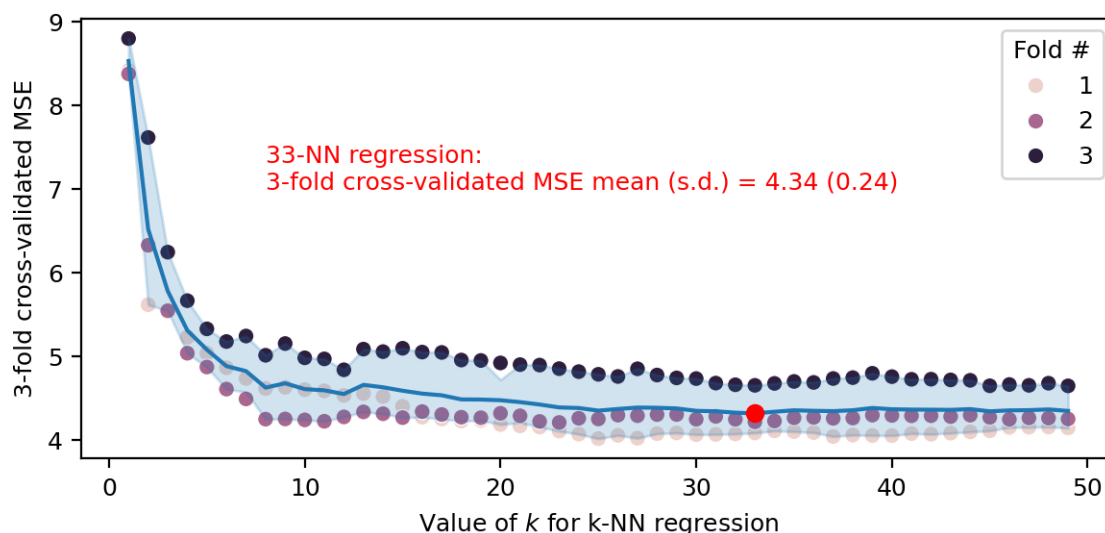
However, in one case it is OK to use all the data for training and cross-validation: if someone else is doing the testing using a part of the data that they have retained and shared with you – this is common in machine learning competitions. They will get the independent estimate, and you can use the data that you have to best train your models.

💡 Example of using K-fold cross validation with k -nearest neighbours on a regression problem

We will recast the parent-child heights data from the Chapter on [Linear Regression](#) as a machine learning problem. We will reserve 25% of the data as test data, and then use the remaining 75% for cross-validation, as shown in Figure 12.7. We decide that we'd like to try two algorithms:

- **k -NN regression**, with various values of the hyperparameter k
- **Linear regression** – there are no hyperparameters for simple linear regression

We need a metric to determine how good each algorithm/hyperparameter pair is. Since this is a regression problem, involving prediction of a continuous quantity, we can use the mean squared error (MSE), introduced in the section on [Numerical diagnostics](#). We first try k -NN with values of k from 1 to 50, and obtain the plot below, in which each point represents the mean MSE found on each of folds 1, 2 and 3 for each value of k , and the blue line represents the mean of the 3 MSEs for each value of k .

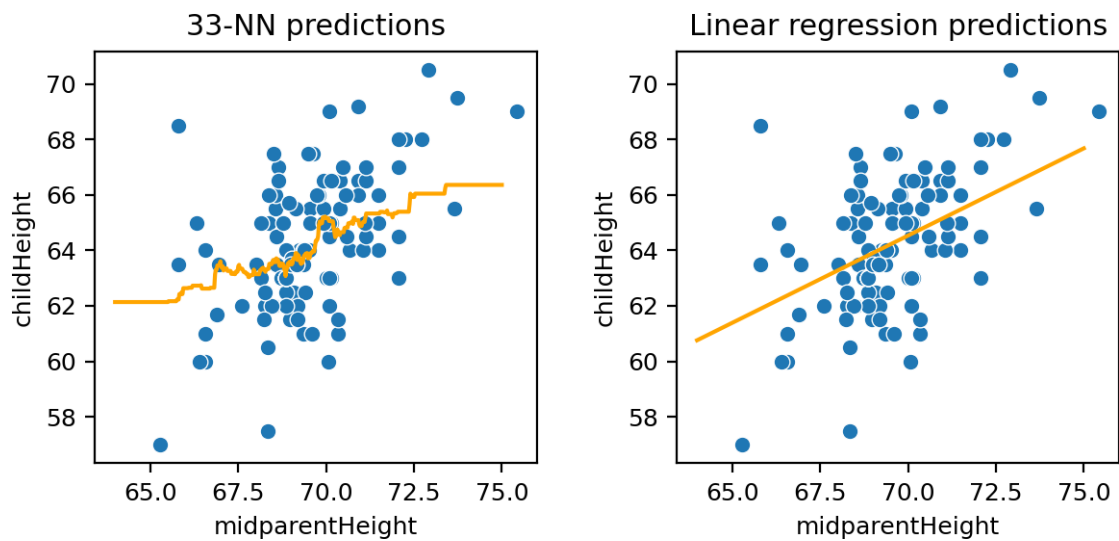


We can see that fold 3 generally has worse MSE than folds 1 and 2. This difference reflects the random sampling used to assign data to the folds – fold 3 was harder to predict. We regard the value of $k = 33$ with the lowest mean MSE as the best value, but there is practically very little difference between values of k above 10. We then train 33-nearest neighbours on all the test data, and take the mean MSE as indication of how well it will perform.

For linear regression, there are no hyperparameters to try out, but we can use cross-validation to generate an estimate of the MSE we should expect in testing. In this case, linear regression turns out to have a lower MSE, and therefore we might select it as our final method. However, for the purposes of illustration we will report test MSEs for both methods: the MSE for linear regression is lower than that for 33-NN, consistent with the indications of the cross-validation MSEs.

Model	3-fold cross-validated MSE mean (s.d)	Test MSE
33-NN	4.34 (0.24)	4.39
Linear regression	4.03 (0.16)	4.31

To get a sense for how the models are working, we plot the predictions against the test data below:



The nearest neighbours predictions are much bumpier, suggesting that despite the hyperparameter search, there is still some overfitting happening. Also, there are end effects at the left and the right, where the set of 33 neighbours is no longer changing.

⚠ Think about what is the appropriate model for the data

The example above shows that a more complex model (here k -Nearest neighbours) is not always the best model – it is often the case the linear regression can be a very helpful model, especially when we have good reason to believe that linear relationships are present – for example, we would expect the price of house to depend on the number of rooms and/or floor area.

However, for data where there are nonlinear relationships, especially complex ones, the greater flexibility of more complex models will allow them to fit the data better than a simpler model such as linear regression. That said, nonlinear terms in multiple regression can often be a good approach for dealing with nonlinear data – see [Interaction terms and nonlinear fits](#).

📺 Related Python Lab: k -nearest neighbours

<https://github.com/Inf2-FDS/FDS-S1-09-knn>

In this lab you will learn how to apply k -Nearest Neighbours (k -NN) to a dataset using the `scikit-learn` library. By the end of the lab you should be able to:

- explain how k can be chosen appropriately
- explain the importance between training, testing and validation data.

Chapter 13

Unsupervised learning: K -means

Further reading (not examinable)

- MacKay, D. J. C. (2003) *Information Theory, Inference and Learning Algorithms* pp. 284-288
- Xu, D. and Tian, Y. (2015). A comprehensive survey of clustering algorithms. *Annals of Data Science*, 2:165. <https://doi.org/10.1007/s40745-015-0040-1>

13.1 Clustering, unsupervised and supervised learning

Clustering aims to partition a dataset into meaningful or useful groups, based on distances between data points. In some cases the aim of cluster analysis is to obtain greater understanding of the data, and it is hoped that the clusters capture the natural structure of the data. In other cases cluster analysis does not necessarily add to understanding of the data, but enables it to be processed more efficiently.

As David MacKay put it:

Human brains are good at finding regularities in data. One way of expressing regularity is to put a set of objects into groups that are similar to each other. For example, biologists have found that most objects in the natural world fall into one of two categories: things that are brown and run away, and things that are green and don't run away. The first group they call animals, and the second, plants. (MacKay, 2003), p. 284

Our aim is to get algorithms to do what human brains do naturally.

Supervised learning Clustering can be contrasted with classification (Figure 13.1). As outlined in the chapter on [Supervised learning: Classification with Nearest neighbours](#), classification is a supervised learning process: there is a training set in which each data item has a label.

Unsupervised learning Clustering, on the other hand, is an **unsupervised learning** process in which the training set does not contain any labels, and where the goal is to learn something about the structure of the data. The aim of a clustering algorithm is to group such a dataset into clusters, based on the unlabelled data alone. In many situations there is no 'true' set of clusters. For example consider the twenty data points shown in Figure 13.2a. It is reasonable to divide this set into two clusters (Figure 13.2b), four clusters (Figure 13.2c) or five clusters (Figure 13.2d).

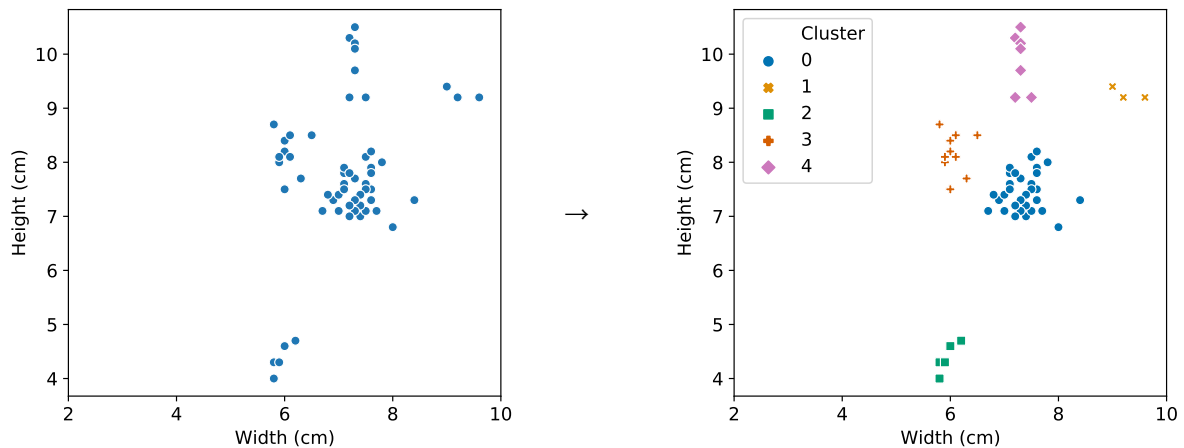
Reasons to cluster There are many reasons to perform clustering. Most commonly it is done to better understand the data (**data interpretation**), or to efficiently encode the dataset (**data compression**).

Data interpretation: Automatically dividing a set of data items into groups is an important way to analyse and describe it. Automatic clustering has been used to cluster documents (such as web pages), user preference data, and many forms of scientific observational data in fields ranging from astronomy to psychology to biology.

Clustering – a type of unsupervised learning procedure

Input: unlabelled data points. e.g. widths and heights of various unknown fruits

Output: each point is assigned to a cluster – which may correspond to the original fruits



Classification – a type of supervised learning procedure

Input: labelled data points. e.g. widths and heights of various fruits

Output: classifier that can predict the identity of an unlabelled data point

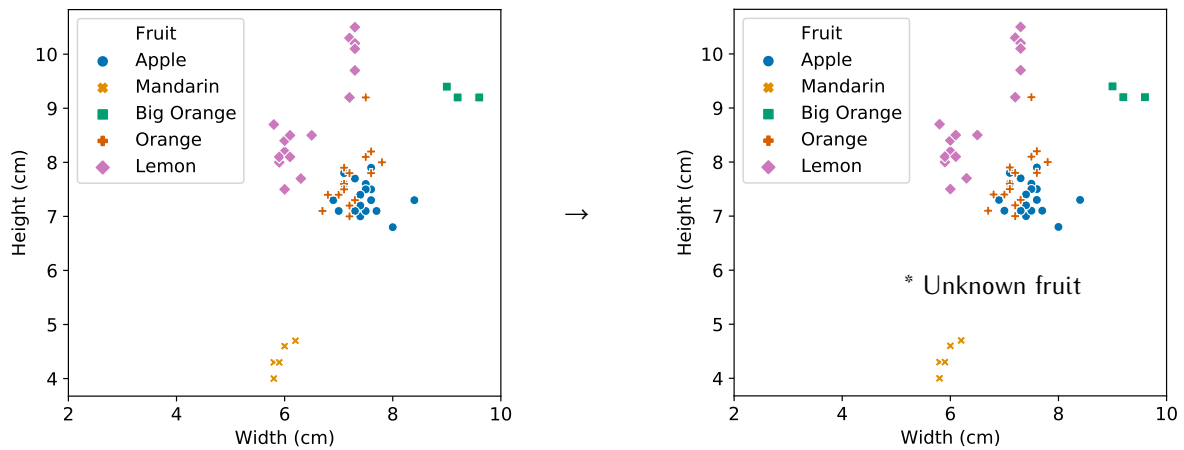


Figure 13.1: Unsupervised and supervised learning, exemplified by clustering and classification applied to Iain Murray's *oranges and lemons* dataset of the widths, heights and masses of various types of fruits.

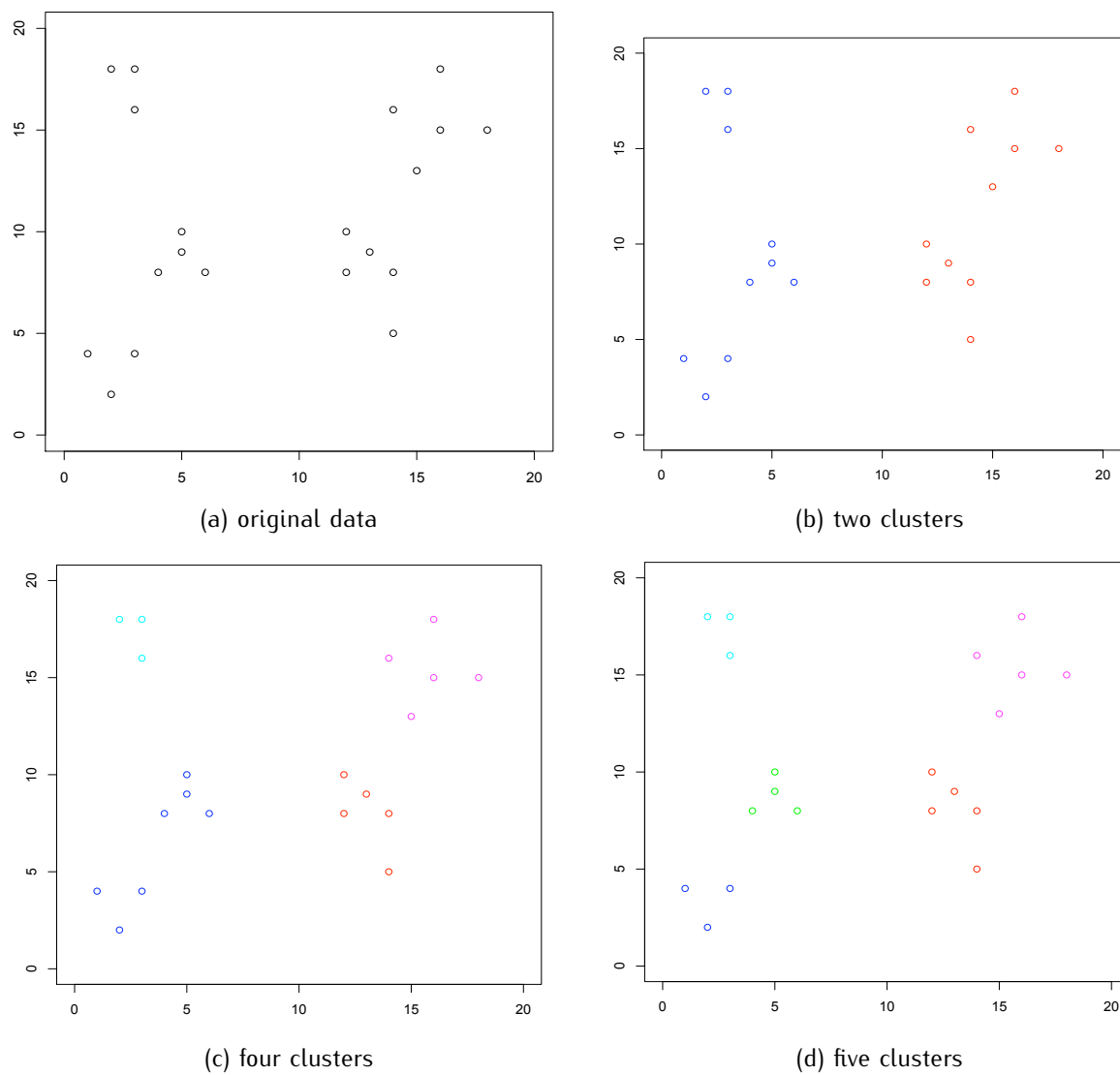


Figure 13.2: Clustering a set of 20 two-dimensional data points.

Data compression: Clustering may be used to compress data by representing each data item in a cluster by a single cluster **prototype**, typically at the centre of the cluster. Consider D -dimensional data which has been clustered into K clusters. Rather than representing a data item as a D -dimensional vector, we could store just its cluster index (an integer from 1 to K). This representation, known as **vector quantisation**, reduces the required storage for a large dataset at the cost of some information loss. Vector quantisation is used in image, video and audio compression.

13.2 Types of clustering

Types of clustering There are two main approaches to clustering: hierarchical and partitional. **Partitional clustering** does not have a nested or hierarchical structure. It simply divides the dataset into a fixed number of non-overlapping clusters, with each data point assigned to exactly one cluster. The most commonly employed partitional clustering algorithm, K -means clustering, is discussed in the next section. **Hierarchical clustering** forms a tree of nested clusters in which, at each level in the tree, a cluster is the union of its children (Figure 13.3). Whatever approach to clustering is employed, the core operations are distance computations: computing the distance between two data points, between a data point and a cluster prototype, or between two cluster prototypes.

Hierarchical clustering There are two main approaches to hierarchical clustering. In **top-down hierarchical clustering** algorithms, all the data points are initially collected in a single top-level cluster. This cluster is then split into two (or more) sub-clusters, and each these sub-clusters is further split. The algorithm continues to build a tree structure in a top-down fashion, until the leaves of the tree contain individual data points. An alternative approach is **agglomerative hierarchical clustering**, which acts in a bottom-up way. An agglomerative clustering algorithm starts with each data point defining a one-element cluster. Such an algorithm operates by repeatedly merging the two closest clusters until a single cluster is obtained.

13.3 K -means

Aim K -means clustering, a form of partitional clustering, aims to divide a set of D -dimensional data points into K clusters. The number of clusters, K , must be specified; it is not determined by the clustering: thus it will always attempt to find K clusters in the data, whether they really exist or not.

Algorithm Each cluster is defined by its cluster centre, and clustering proceeds by assigning each of the input data points to the cluster with the closest centre, using a Euclidean distance metric. The centre of each cluster is then re-estimated as the *centroid* of the points assigned to it. The process is then iterated. The algorithm, illustrated in Figure 13.4, is:

- *Initialise* K cluster centres, $\{\mathbf{m}_k\}_1^K$
- While not *converged*:
 - Assign each data vector $\mathbf{x}_i$ ($1 \leq i \leq n$) to the closest cluster centre;
 - Recompute each cluster mean as the mean of the vectors assigned to that cluster

To be more precise, we can express the assignment of a point i to a cluster k by defining the set of points in cluster k as $\mathcal{C}_k$. Using the notation $|\mathcal{C}_k|$ to indicate the number of points assigned to cluster k , we can then recompute the mean of cluster k using the following equation:

$$\mathbf{m}_k = \frac{1}{|\mathcal{C}_k|} \sum_{i \in \mathcal{C}_k} \mathbf{x}_i \quad (13.1)$$

Distance measure The algorithm requires a distance measure to be defined in the data space, and the Euclidean distance is often used. In this case, if $\mathbf{x}$ and $\mathbf{y}$ are two points in the data space, the distance function $d(\mathbf{x}, \mathbf{y})$ is given:

$$d(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\| = \sqrt{\sum_{j=1}^D (x_j - y_j)^2} \quad (13.2)$$

where $\|\cdot\|$ denotes the Euclidean norm (i.e. L^2 -norm) of a vector.

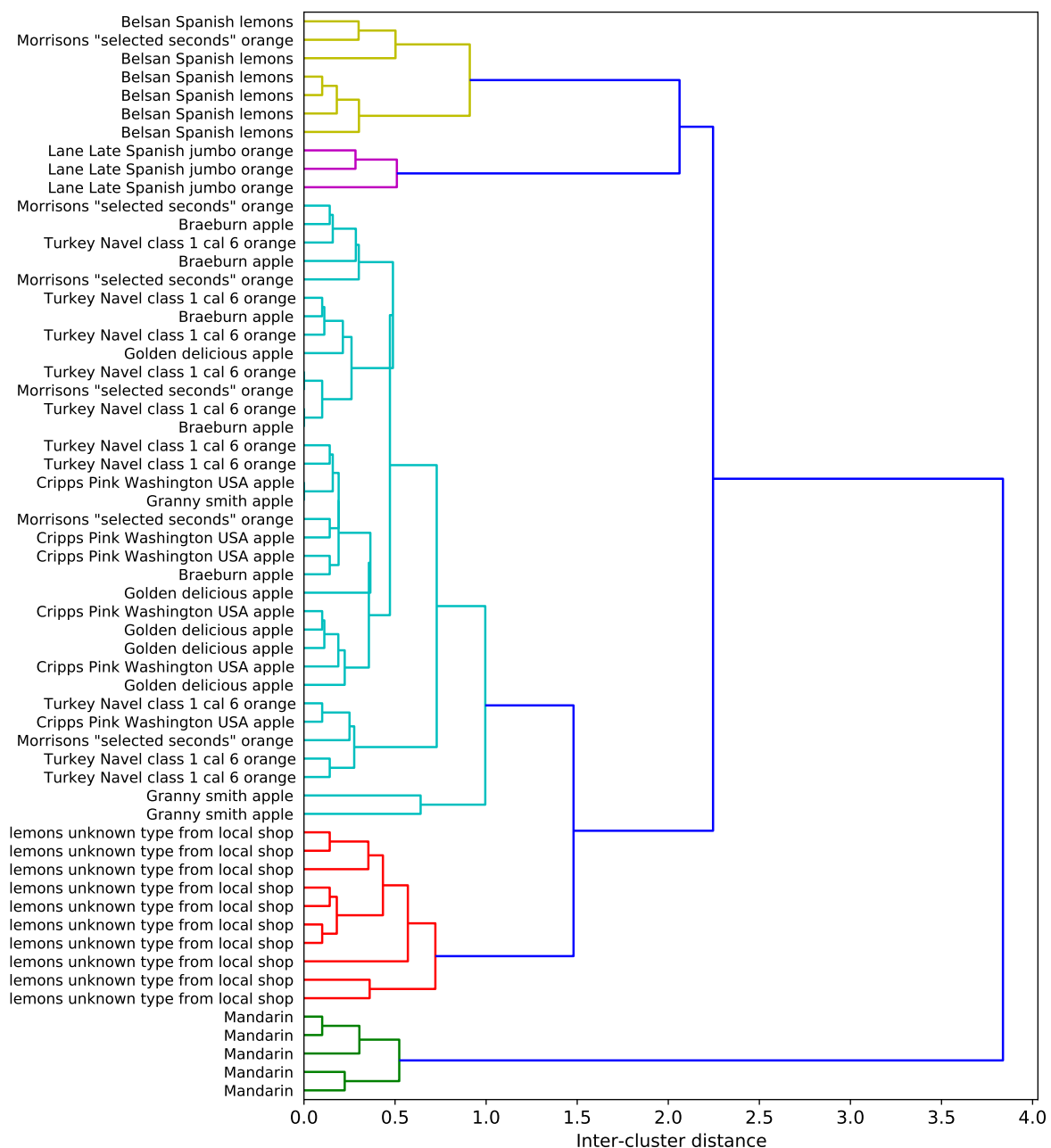


Figure 13.3: Hierarchical clustering of the fruit dataset. The x-axis indicates the distance between clusters, which in this clustering example is the centroid of all the points belonging to the cluster. For example the two Mandarins at the very bottom are about 0.2 units apart, and form a cluster. In turn the centre of this cluster is about 0.45 units from the centre of the next closest cluster (the three other Mandarins). We can regard all sub-trees lower than a threshold inter-cluster distance of our choosing as clusters. For example, if we chose a threshold of 1.0, we would see 5 clusters, indicated by the green, red, cyan, magenta and yellow sub-trees.

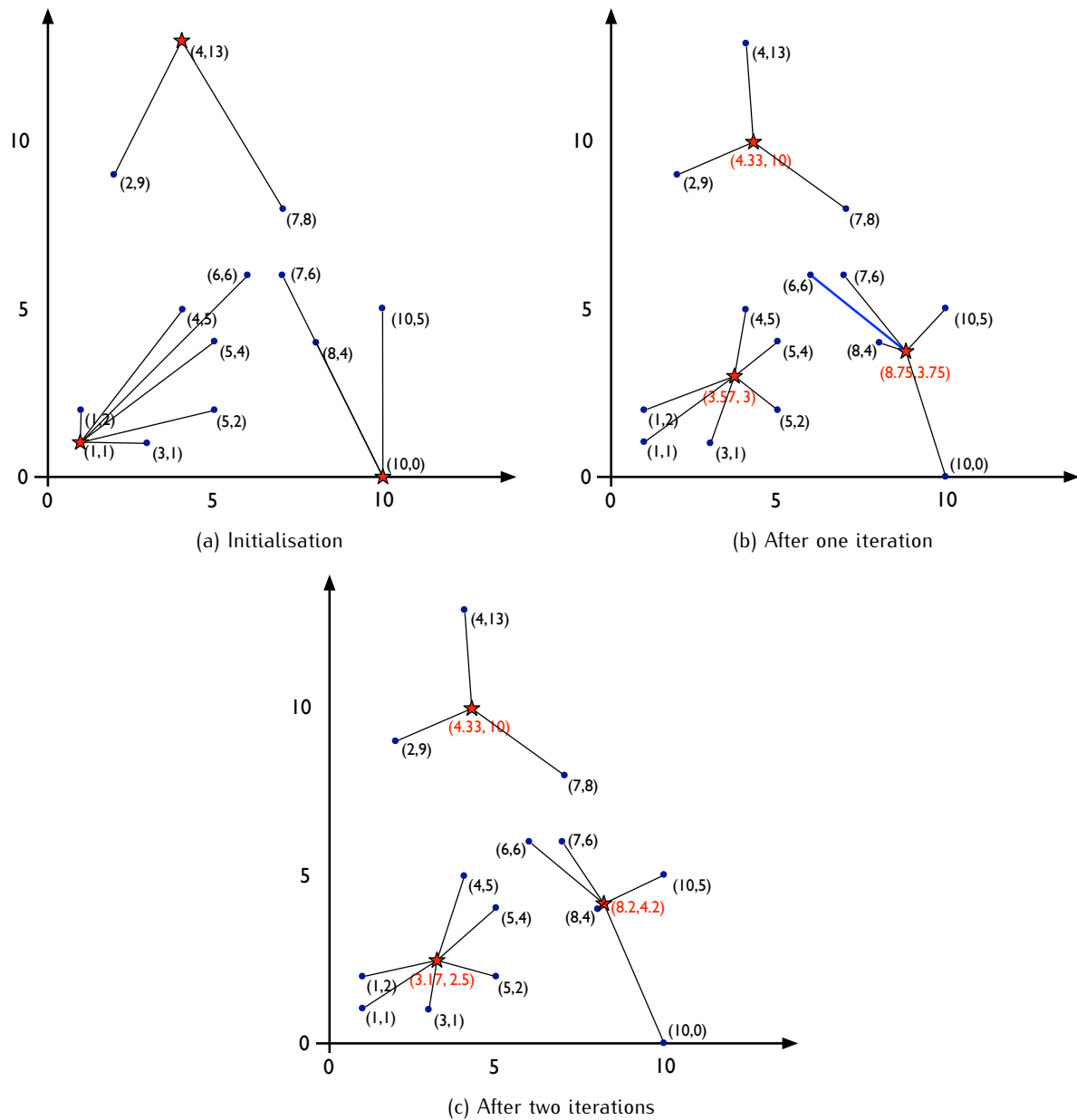


Figure 13.4: Example of K -means algorithm applied to 14 data points, $K = 3$. The lines indicate the distances from each point to the centre of the cluster to which it is assigned. Here only one point $(6,6)$ moves cluster after updating the means. In general, multiple points can be reassigned after each update of the centre positions.

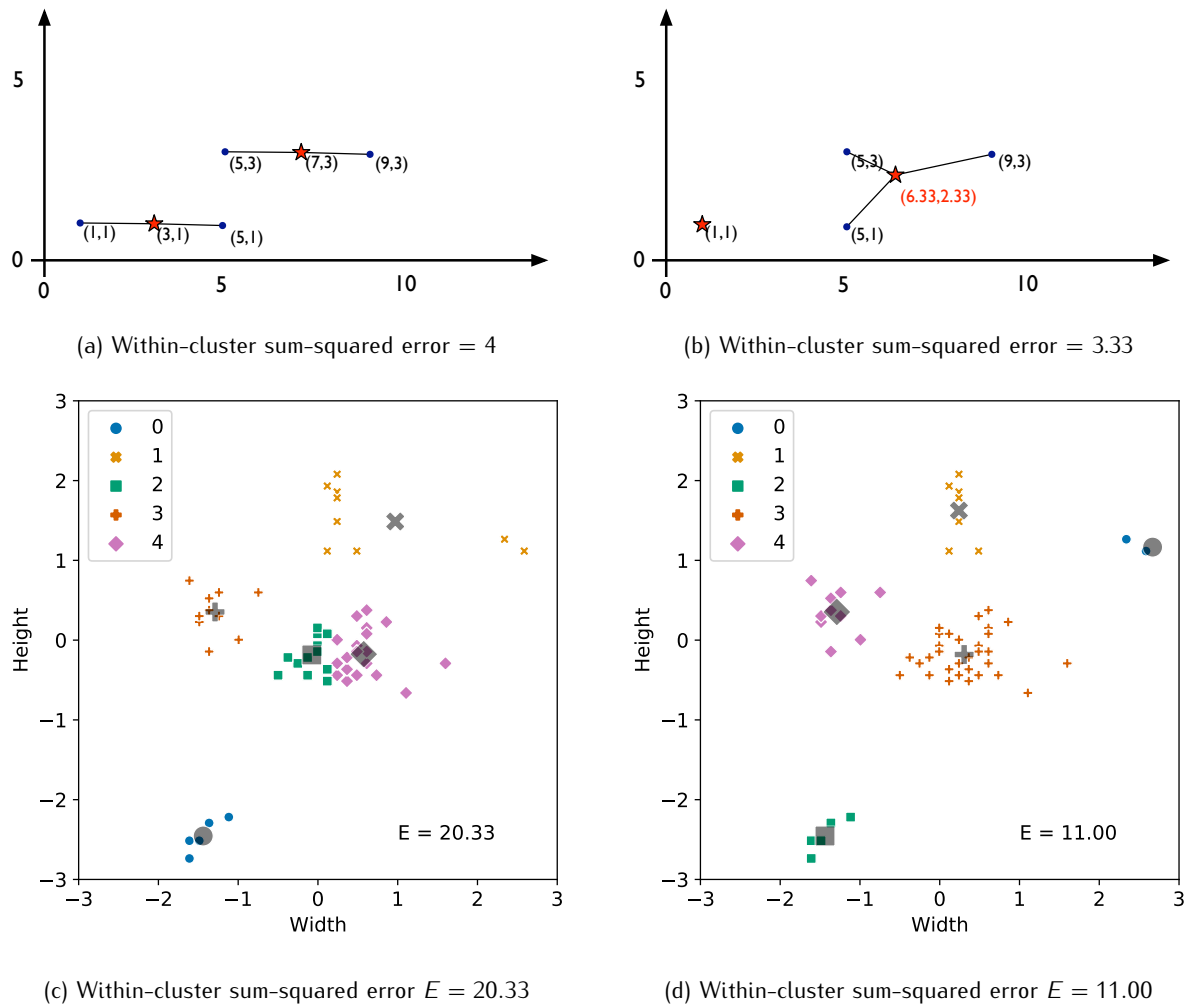


Figure 13.5: (a) and (b) Two different converged clusterings for the same dataset, but starting from different initialisations. (c) and (d) Two different converged clusterings for fruit dataset, but starting from different initialisations.

Initialisation The initialisation method needs to be further specified. There are several possible ways to initialise the cluster centres, including:

- Choose random data points as cluster centres
- Randomly assign data points to K clusters and compute means as initial centres
- Choose data points with extreme values
- Find the mean for the whole dataset then perturb into K means

All of these work reasonably, and there is no ‘best’ way. However, as discussed below, the initialisation has an effect on the final clustering: different initialisations lead to different cluster solutions.

Convergence The algorithm iterates until it converges. Convergence is reached when the assignment of points to clusters does not change after an iteration. An attractive feature of K -means is that convergence is guaranteed. However, the number of iterations required to reach convergence is not guaranteed. For large datasets it is often sensible to specify a maximum number of iterations, especially since a good clustering solution is often reached after a few iterations. Figure 13.4 illustrates the K -means clustering process. Figure 13.5 illustrates how different initialisations can lead to different solutions.

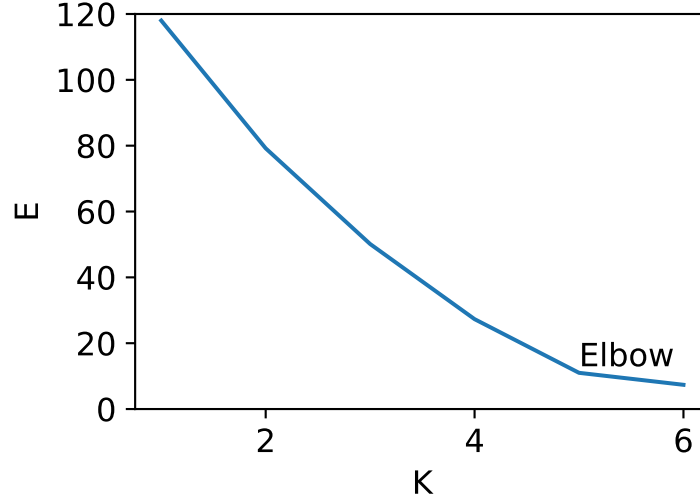


Figure 13.6: Scree plot of K -means clustering applied to the fruit dataset. The elbow is around $K = 5$, implying that 5 is a reasonable number of clusters.

13.4 Evaluation and application of K -means clustering

Multiple solutions Depending on the initialisation of the cluster centres, K -means can converge to different solutions; this is illustrated in Figure 13.5 in which the same dataset of 4 points has two different clusterings, depending on where the initial cluster centres are placed. Both these solutions are **local minima** of the error function, but they have different error values. For the solution in Figure 13.5a, the error is:

$$E = \frac{((4 + 4) + (4 + 4))}{4} = 4.$$

The second solution (Figure 13.5b) has a lower error:

$$E = \frac{0 + (32/9 + 20/9 + 68/9)}{4} = \frac{10}{3} < 4.$$

Figures 13.5c and 13.5d show two clusterings of the fruit dataset – the first clustering is noticeably worse, with one cluster centre in-between two distinct groups of points.

To get around the problem of multiple solutions, it is common to repeat K -means algorithms from multiple randomly-chosen starting points, and then take the solution that gives the lowest value of a within-cluster mean squared error E , that we will define shortly.

Mean squared error function To compare two different clusterings each with K clusters we can use the **within-cluster mean squared error** function¹, which can be thought of as measuring the **scatter** of the data points relative to their cluster centres. Thus, if we have two sets of K clusters, it makes sense to prefer the one with the smallest mean squared error: the clustering with the lowest scatter of data points relative to their cluster centres.

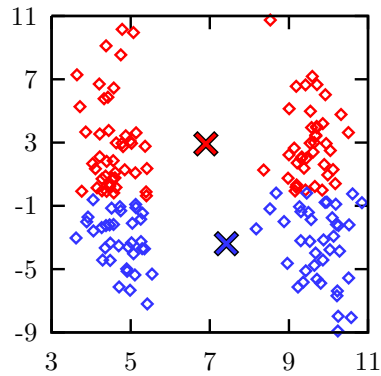
Let us define the set of points in the k th cluster as $\mathcal{C}_k$. Then we can write the mean squared error as

$$E = \frac{1}{n} \sum_{k=1}^K \sum_{i \in \mathcal{C}_k} \|\mathbf{x}_i - \mathbf{m}_k\|^2, \quad (13.3)$$

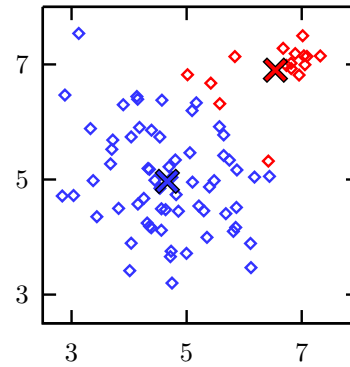
where $\mathbf{m}_k$ is the centre of cluster k , n is the number of data points in total, and $\|\cdot\|$ denotes the Euclidean norm (i.e. L^2 -norm) of a vector.

Another way of viewing this error function is in terms of the squared deviation of the data points belonging to a cluster from the cluster centre, summed over all centres. This may be regarded as a variance measure for each cluster, and hence K -means is sometimes referred to as **minimum variance clustering**.

¹Commonly in K -means, the **within-cluster sum of squared errors**, also known as the inertia, is used instead of the mean squared error. It is simply the mean squared error, but without the $1/n$ factor.



(a) Differences in scale on the axes cause nonsensical clusters.



(b) A smaller cloud (red points) attracts points from a larger cloud (blue points).

Figure 13.7: Failures of K -means. Figures by Iain Murray.

At the start of this chapter we said that the aim of clustering was to discover a structure of clusters such that points in the same group are close to each other, and far from points in other clusters. The mean squared error only addresses the first of these: it does not include a between-clusters term.

Failures of K -means Like all variance-based approaches this criterion is dependent on the scale of the data. This if variables are on different scales, nonsensical clusters can arise (Figure 13.7a). The range of the y -coordinates is around 20, whereas the range on the x -axis is about 8. Thus, a point at the top of the left-hand cluster is closer to a point at the top of the right-hand cluster than it is to a point at the bottom of the left-hand cluster. Standardising variables is often a sensible step to take solve this problem – here it would reduce the distances in the y -direction relative to the x -direction.

Another failure mode of K -means is that large clouds can pull small clusters off-centre (Figure 13.7b). The reason for this is that we've not modelled the variance of each cluster – we've implicitly assumed they are the same. The solution is to allow the variance of the clusters to vary. Gaussian mixture models do this, but are beyond the scope of this course.

How to decide on K We can use a **scree plot** (Figure 13.6), plotting the mean squared error E against the number of clusters K . We are looking for the **elbow**, where the error stops falling dramatically. Scree plots are also used to inform decisions about the number of dimensions to retain in dimensionality reduction methods (Dealing with high dimensions – PCA).

Other variants of K -means We have discussed the **batch** version of K -means. The **online** (sample-by-sample) version in which a point is assigned to a cluster and the cluster centre is immediately updated is less vulnerable to local minima.

There are many variants of K -means clustering that have been designed to improve the efficiency and to find lower error solutions.

The curse of dimensionality As the number of dimensions D increases, the distances between points increases. This has the effect of making the intracluster distances on a similar scale to the intercluster distances, and so clustering becomes less reliable. This is one manifestation of the **curse of dimensionality**. To counteract the curse, a dimensionality reduction method such as Principal Components Analysis (PCA) can be applied to the data before it is clustered (see chapter Dealing with high dimensions – PCA).

Related Python Lab: K -means

<https://github.com/Inf2-FDS/FDS-S2-03-kmeans>

In this lab you will implement the K -means algorithm from scratch. By the end of the lab you should be able to:

- implement and explain the different steps involved in K -means

- explain what the benefits of clustering algorithms are
- explain what partitional clustering algorithms are
- and be able to use sklearn's K-means algorithm.

Part IV

Statistical inference

Chapter 14

Randomness, sampling and simulation



Recommended reading

- *Modern Mathematical Statistics with Applications*, Sections 6.1 and 6.2

14.1 Introduction to statistical inference

Inferential statistics In the chapter on [Descriptive statistics](#), we looked at the difference between a sample and a population. We also considered a number of statistics that could apply to the population and to the sample: the mean, variance, standard deviation and median. **Statistical inference** is the process of drawing conclusions about quantities that are not observed ([Gelman et al., 2004](#)).

One example of statistical inference is inferring the mean of quantity in a population from a sample of that population. For example, suppose we'd like to know what the mean weight of a wild cat is. We've weighed (observed) a sample of 10 wild cats from a population of 400 and therefore know the mean weight in this sample and sample standard deviation of the weight. What can we conclude about the mean weight in the population, 390 of which we have not observed?

Inferential statistics has been around for much longer than the word "statistics". The 9th-century book "Manuscript on Deciphering Cryptographic Messages" written in Arabic by Al-Kindi (educated in Baghdad) shows how to decipher encrypted messages by frequency analysis, i.e. counting the frequency of particular letters.

Inferential statistics tasks Inferential statistics can seem like a toolbox full of tools with confusing names such as "standard error of the mean", "*t*-test", " χ^2 test", and "bootstrap", and a confusing set of rules about what to use each tool for. We're going to try to give you an idea of what task each tool is useful for, and how it works. There are three main tasks we will consider:

1. Estimation
2. Hypothesis testing
3. Comparing two samples (A/B testing)

Estimation The estimation task addresses the question "how big is some number that we can't measure directly, and how certain are we about our answer"? More formally, in the **estimation** task, we are trying to infer a statistic of a population from data from a sample of that population. Inferring the mean weight of the population of wild cats from a sample is an example of estimation. Another example is opinion polling, where pollsters try to estimate the voting intentions of a population of voters from a sample.

There are two sub-tasks in estimation:

- **Point estimation:** deriving the best possible estimate of the population statistic from the sample data. In the case of measuring the mean weight of a population of animals, the point estimate of the population mean may be simply the sample mean. For example if the sample mean of a population of

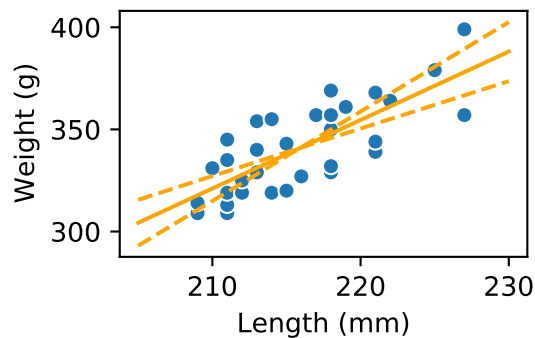


Figure 14.1: Uncertainty in regression line. The best estimate regression line (solid) and the lines at the edge of the 95% confidence interval (dashed lines). Note this plot is simplified, since the uncertainty in the intercept is not represented.

squirrels is 320 g, our best estimate of the population mean is 320 g. In the case of opinion polling, the point estimator may be more complex, perhaps giving more weight to certain demographics, depending on how likely they are to vote or be truthful about their preferences.

- **Confidence interval estimation:** deriving an indication of how confident we can be in that estimate from the sample data and the size of the sample. The confidence interval can be expressed as a range around the point estimate. For example, we'll be able to say that we are 95% confident that the interval [304, 336]g contains the population mean of the weight. Where the confidence interval is symmetric around the point estimate, we may write it as the point estimate plus or minus half the confidence interval, e.g. 320 ± 16 g. In the polling example, the confidence interval is called the "margin of error". We'll learn how to calculate confidence intervals using various methods, but it's worth remembering that the bigger the sample size, the smaller the confidence interval.

Another example of an unobserved quantity is linear regression coefficients. We already know how to find (point) estimates of them, using the formulae we covered earlier in the course. In a linear regression of the weight of a sample of squirrels on their length (Figure 14.1), we will be able to say that the best estimate of the slope of the regression line is 3.35 g/mm, but we are 95% confident that the slope lies in the interval [2.32, 4.38] g/mm.

Hypothesis testing Hypothesis testing answers "Yes/No" questions, for example "is chocolate good for you?". More formally, in hypothesis testing, we are trying to ascertain which of two or more competing theories are the best explanation of the data. For example, in 1965 a court case was brought against the state of Alabama (Swain versus Alabama, 1965) due to there being no Black members of the jury in a trial.¹ Part of the case concerned the fact that at one stage of the jury selection, 8 Black people were chosen for a jury panel of 100 people, but the fraction of Black people in the population was 26%. Our question is "Is the jury selection system biased against Black people?".

Comparing two samples (also known as A/B testing) Here we have two samples that have been treated differently, and we want to either test if the groups are different, or estimate how different they are. For example, to find out the effectiveness of a vaccine, we select a sample of volunteers from the population randomly, divide them randomly into two groups, give the vaccine to one group (Treatment group) and give the other group a placebo (Control Group). In the vaccine group 3 volunteers catch the disease, but in the placebo group 95 volunteers catch the disease. Is the vaccine effective? How much would we expect the vaccine to cut the risk of catching the disease if we give it to the whole population?

In the context of user testing, often in web applications, this is called A/B testing. A famous example was at Amazon, where a developer had the idea of presenting recommendations as customers add items to a shopping cart (Kohavi et al., 2007). Amazon managers forbid the developer to work on the idea, but the developer disobeyed orders and ran a controlled experiment on customers, by splitting them into two groups ("A" and "B"), one which had recommendations shown and one which didn't. The group which had recommendations shown bought more, and displaying recommendations quickly became a priority for Amazon.

¹We found the example of Swain versus Alabama in Adhikari et al. (2020), and follow their treatment.

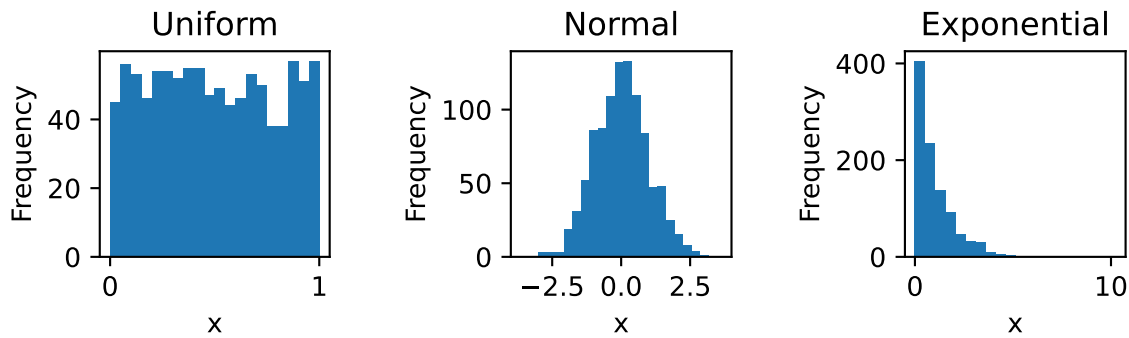


Figure 14.2: Histograms of 1,000 samples taken from normal, uniform and exponential distributions.

Two approaches to statistical inference We are going to learn a number of techniques for undertaking point and interval estimation, hypothesis testing and comparing samples. We will also think carefully about the interpretation of these techniques. There are two main approaches to undertaking statistical inference tasks, statistical simulations and statistical theory:

1. **Statistical simulations:** Here we use repeated random sampling to carry out the statistical inference procedures. The advantages of statistical simulation procedures are they often require fewer assumptions about the data and/or hypothesis than statistical theory, and they require somewhat less theory to understand. However, they can be compute-intensive, and care is still needed in their use.
2. **Statistical theory:** Here we use the properties of various well-known theoretical distributions to draw inferences about our data. We need to check that the assumptions behind the distribution match the statistical question we are trying to answer. For example, a distribution of delays to flights is likely to be highly right-skewed, so we shouldn't assume a normal distribution when dealing with it. Typically, the process is not compute-intensive: very often it amounts to arithmetic and then reading of a quantity from a distribution table. These procedures come as standard in a number of stats packages, including R and Python's statsmodels.

A number of fundamental concepts underpin both the statistical theory and statistical simulations.

Plan for this part of the course The plan for the statistical inference topics in this part of the course will be:

1. Fundamental theory (the rest of this chapter):
 - We'll learn how we can use statistical simulations to generate samples from a model and compute statistics for each of these samples to give a **sampling distribution**.
 - Learn about the distribution of the mean of repeated samples from a model. This will lead us to the Central Limit Theorem, which can help us to estimate the uncertainty in our estimate of the mean, i.e. confidence intervals, and the Law of Large Numbers, which also helps with estimation.
2. [Estimation](#) and [Confidence intervals](#)
3. [Hypothesis testing](#) and [p-values](#)
4. [A/B testing](#)

14.2 Sampling, statistics and simulations

Sampling A prerequisite for statistical simulations is being able to sample from probability distributions and from sets of discrete items, including observed data.

Random sample Following the definition given in *Modern Mathematical Statistics with Applications*, in **random sample** of size n from either a probability distribution or a finite population of N items, the random variables $X_1, \dots, X_n$ comprising the sample are

1. all independent
2. have the same probability distribution.

Strictly speaking, conditions 1 and 2 mean the sample is an independently and identically distributed (i.i.d.) random sample. When we use the term “random sample” in these lecture notes, we will mean i.i.d. random sample. This definition is important because a number of mathematical results hold for i.i.d. random samples.

i Random number generation and pseudo-random numbers

The OED definition of random includes the phrase “completely unpredictable in details”, corresponding to the everyday notion of “random” meaning “unpredictable”. Physical devices, such as dice and coins, have been used to generate random numbers for thousands of years, and there are now hardware **random number generators**, which use various physical phenomena to generate random numbers that are completely unpredictable.

Programming languages typically contain software routines to generate random numbers, like those contained in the Python `random` and `np.random` modules. However, the sequence of numbers generated by these software routines are deterministic; when started from a particular value of a number called the random seed, the generator will repeat the same sequence of numbers, meaning the numbers are predictable in principle. In practice, the properties of the numbers produced by a good random number generator satisfy various tests that would be satisfied by a true random number generator, such as there not being correlations between adjacent numbers in the sequence. We therefore refer to the numbers produced by software random number generators as **pseudorandom numbers**.

Setting a random seed with can be helpful, as it allows results to be reproduced. However, it can be risky to set the seed, since if it is reset in between creating two samples, we may end up with the same sample twice or multiple times.

Sampling from probability distributions Sampling from a probability distribution requires generating random numbers. A standard random number generator produces numbers within an interval (e.g. $[0, 1]$) with uniform probability for each number, i.e. it samples from a uniform distribution. We can demonstrate the distribution of a standard random number generator by drawing many samples and plotting a histogram (Figure 14.2). We adapt these functions to sample from any univariate distribution, e.g. a normal distribution or an exponential distribution (Figure 14.2).

Sampling from a set of discrete items We can also sample from a population of discrete items. We can select n items from a set of N items at random either **without replacement** or **with replacement**.

If we sample without replacement, it is as though we are drawing items of various types (e.g. coloured balls) at random out of a bag, and not replacing them. We can only draw (sample) up to N items. Also as we remove items from the bag, the probabilities of drawing a particular type (colour) changes. Thus samples drawn without replacement are not independently and identically distributed, and so are not consistent with our definition of a random sample.

If we sample with replacement, we put the item back in the bag, before making our next draw – we can carry on doing this for ever. The probability of picking a particular item is the same on every draw (identically distributed), and one draw does not affect another (independent), so sampling with replacement does give a random sample according to our definition.

If we are sampling without replacement a number of items n that is much smaller than the population size N , the probabilities will not change much from one draw to the next. In practice, if $n/N < 0.05$, sampling without replacement is approximately the same as sampling with replacement.

We could construct an algorithm for random sampling either with or without replacement from a uniform random number generator, but these functions are provided in packages such as `numpy.random.choice` in Python.

A particular application of sampling from a set of discrete items is creating a sample of a larger dataset.

Relative frequency We have encountered **relative frequency** when discussing how to construct histograms (Section on [Distributions of numeric variables](#)). In the context of sampling it is helpful to give a definition related to the probability theory covered in [Discrete Maths and Probability](#). To set the scene, suppose we draw one ball from a bag containing 2 red, 2 blue and 1 yellow ball with replacement n times. We regard each draw as one replication of an experiment. The result of each experiment (replication) is an event

$E \in \{\text{red, blue, yellow}\}$. Each event contains a subset of the event space of five possible outcomes (the five balls). We can then define:

$$\text{Relative frequency of event } E = \frac{\text{number of replications in which } E \text{ occurs}}{\text{number of replications}} = \frac{n_E}{n} \quad (14.1)$$

If we assume that each theoretical outcome (i.e. each ball) is equally likely to be picked, and replace “replication” with “outcome” in the formula above, the above formula would give the theoretical probability of event E . If the events E are disjoint, and cover the sample space, then the sum of the relative frequencies add to 1, as would the theoretical probabilities of the set of events. These similarities between relative frequency and theoretical probability are reflected in an alternative name for relative frequency, **empirical probability**.

In the context of sampling without replacement from a large population (for example in an opinion poll) we can think of each observation (e.g. surveyed individual) as a replication, and the sample size as the number of replications.

Non-random samples from a population We can also imagine ways of sampling that are not systematically random, which could lead to sampling or selection bias in our results (see section on [Bias](#)). For example, we might have a list of the daily takings in a restaurant. We could take the first n days. But suppose that the dataset has been sorted in terms of takings? We would then have days with low takings at the start of the list, so the statistics of the sample would not resemble the statistics of the population. We could try taking every 7th day in the list – but if the list is in date order we will always be sampling from one day of the week, e.g. Mondays. Mondays may be different from other days, and thus lead to sampling bias in our estimate of the takings. Random sampling ensures that we don’t have this type of problem.

Samples of convenience When we are collecting data, it might be tempting to sample from the data that we can collect conveniently. For example, a polling company may find it easier to contact people who have more time to answer the phone, which may tend to be retired people. If we don’t correct for this sort of bias, it’s called a **sample of convenience**. One way of combating convenience sampling is **stratified sampling**, in which the sampling is targeted so that the proportions of attributes of the sample matches the proportions in the population.

Definition of a statistic Before going further, it’s helpful to have the definition of **statistic**: “A **statistic** is any quantity whose value can be calculated from sample data.” (*Modern Mathematical Statistics with Applications*, chapter 6). We probably recognise the mean, variance and median as statistics by this definition. But we’ve also derived other quantities from sample data, such as the correlation coefficient and regression coefficients – they are also statistics. We will follow *Modern Mathematical Statistics with Applications* and denote a statistic using an uppercase letter, to indicate that it is a random variable, since its value depends on the particular sample selected. E.g. $\bar{X}$ represents the mean and S^2 the variance.

Simulations and sampling Before considering inferential statistics proper, we will focus on running statistical simulations, i.e. using a computer program to make predictions from probabilistic models of real-world processes. For example, the probabilistic model of tossing a coin multiple times is that the tosses are independent and that the probability of a head is $1/2$ (or perhaps another value, if with think the coin is loaded). The statistical simulation generates a sequence of heads and tails.

To run a statistical simulation we need to decide on:

- The statistic of interest ($\bar{X}$, S , etc.)
- The population distribution (e.g. normal with particular mean and variance) or set of discrete items
- The sample size (denoted n)
- The number of replications k

The simulation procedure is then:

1. For i in $1, \dots, k$
 - (a) Sample n items from the population distribution or set of discrete items
 - (b) Compute and store the statistic of interest for this sample
2. Generate a histogram of the k stored sample statistics

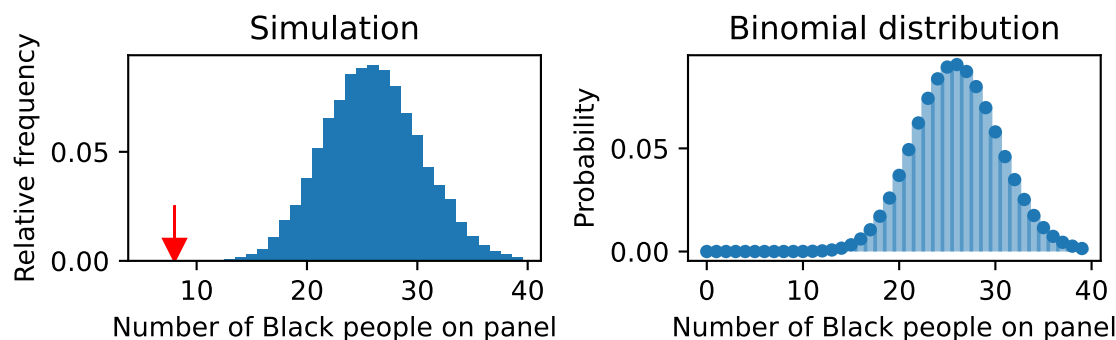


Figure 14.3: Results of statistical simulation of the panel size in Swain versus Alabama (1965). **Left:** The blue histogram what fraction of 10 000 replications of the simulation produced jury panels of 100 with the given number of Black people on them. (It is therefore shows the relative frequency of a jury with the given number of Black people on it.) The red arrows indicates the number of Black jurors in Swain versus Alabama (1965). **Right:** The theoretical probability distribution, a Binomial distribution with $p = 0.26$ that corresponds to the statistical simulation.

Example of hypothesis testing using a simulation experiment To demonstrate the utility of the statistical experiment we've introduced, let's look again at the example in which 26% of the population is Black and 8 Black people are selected to be on a jury panel of 100 people. The null hypothesis H_0 is "The jury panel was chosen at random from the population". We can map the null hypothesis onto the general framework above as follows:

- The statistic of interest is T_0 , the number of Black people in a sample of $n = 100$ panel members
- The population distribution is a Bernoulli distribution with the sample space {Black, Non-Black} in which $P(\text{Black}) = 0.26$.
- The sample size is $n = 100$
- The number of replications $k = 10\,000$

We follow the procedure described in the previous section to give the results shown in the "Simulation" plot in Figure 14.3. Coding this up will be an exercise for you in the Lab. We can see that none of the 10 000 simulations of the null hypothesis produced a jury with 8 members, suggesting that we should reject the null hypothesis in favour of an alternative one. This looks like a clear-cut case; in the chapter on [Hypothesis testing and \$p\$ -values](#), we'll consider in more detail how to interpret the results when the data is less distinct from the simulations.

Deriving the sampling distribution Note that in this example, we didn't have to go to the trouble of running a simulation experiment. We might have noticed that the total number of Black people will be distributed according to a binomial distribution with $n = 100$ and $p = 0.26$, as shown in the "Binomial distribution" plot in Figure 14.3.

14.3 Distributions of statistics of small samples from probability distributions

Example of sampling from probability distributions In the previous example, we've sampled a total number of successes from a Bernoulli distribution. We'll now look at what happens when we sample the mean, standard deviation and median from the normal, uniform and exponential distributions by running the following simulations:

- Statistics of interest: Either
 - (a) mean $\bar{X}$
 - (b) standard deviation S
 - (c) median $\tilde{X}$

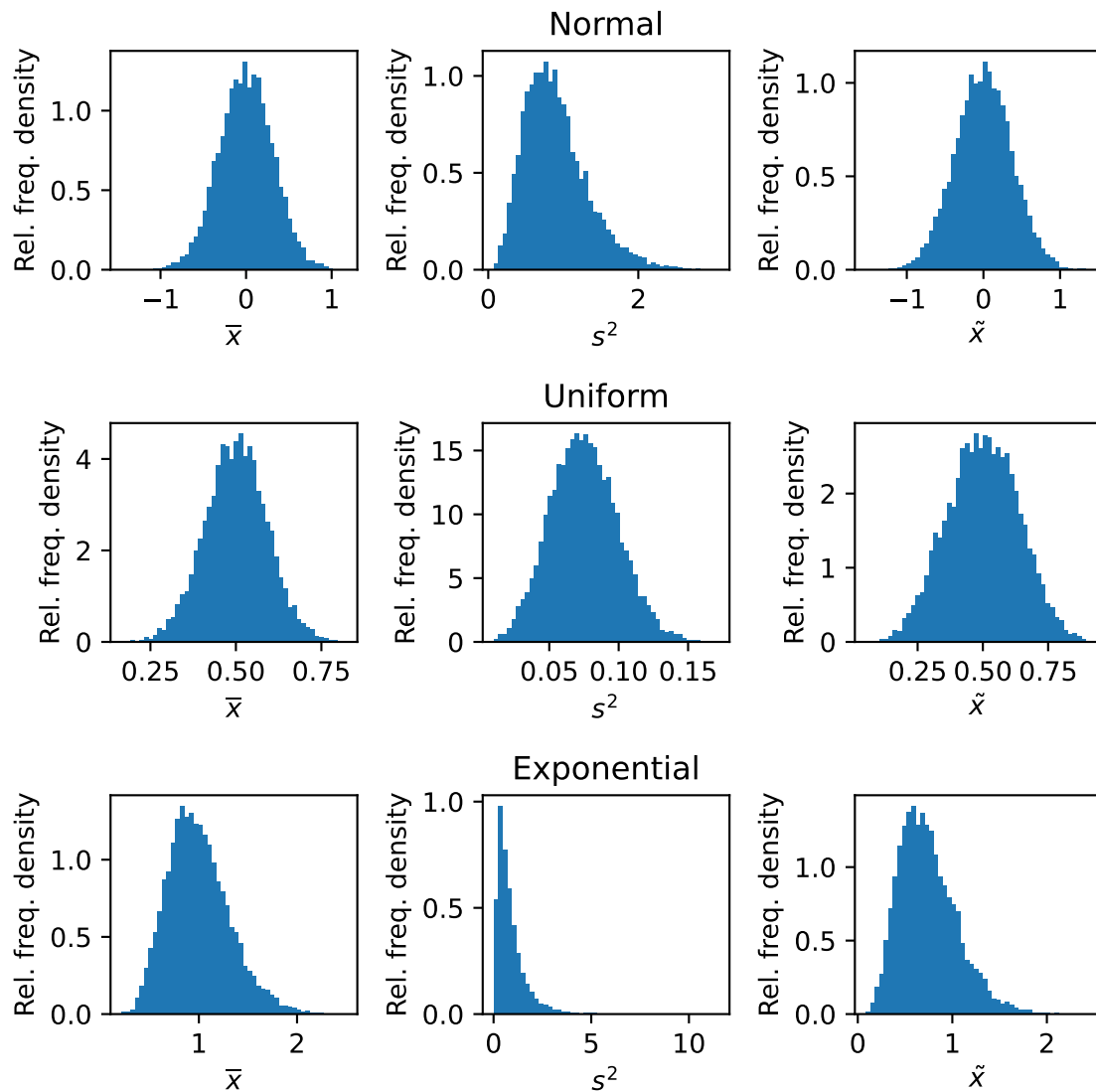


Figure 14.4: Sampling distribution generated by 10 000 simulations of the mean $\bar{x}$, variance s^2 and median $\tilde{x}$ of 10 samples drawn from a normal distribution (top row), uniform distribution (middle row) and exponential distribution (bottom row).

- Population probability distribution: Either
 - (a) Normal distribution with mean 0 and variance 1
 - (b) Uniform distribution on $[0, 1]$
 - (c) Exponential distribution $p(x) = e^{-x}$.
- Sample size $n = 10$
- Number of replications $k = 10\,000$

Figure 14.4 shows the sampling distributions of the three statistics (columns) generated from each of the three population distributions (rows). There are a number of points to notice:

Sample mean (first column), all distributions The distribution of the mean is narrower than the original distribution in every case. This is because some of the variability in the individual samples is averaged out. The standard deviation of this distribution is called the **standard error of the mean**.

Sample mean of normal distribution The distribution looks to be normal – it turns out that this is easy to prove.

Sample mean of uniform distribution The distribution is symmetric and looks to be near-normal.

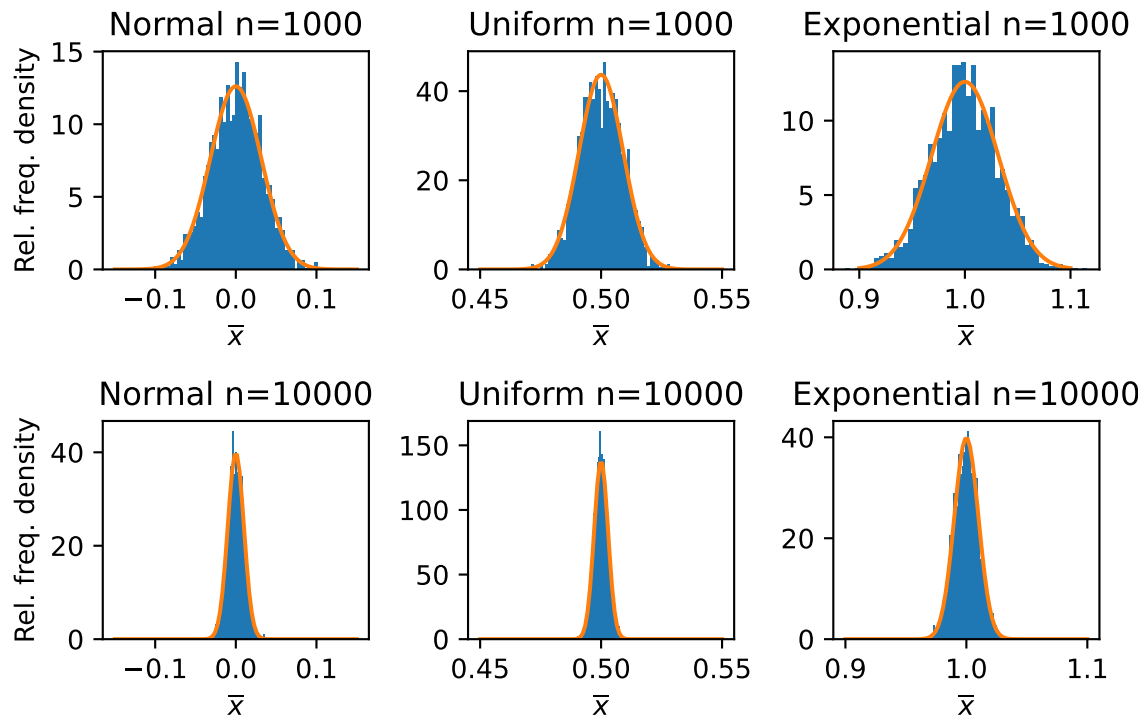


Figure 14.5: Distributions of means from samples of size $n = 1000$ (top row) and $n = 10000$ (bottom row) drawn from the normal, uniform and exponential distributions shown in Figure 14.2. The blue histograms show the histograms obtained from $k = 2000$ simulations. The orange curves are normal distributions with mean equal to the mean of the original distribution and variance $\sigma_{\bar{x}}$ equal to σ^2/n , where σ^2 is the variance of the original distribution.

Sample mean of exponential distribution The distribution is clearly skewed, but less so than the original exponential distribution.

Sample variance (second column) All these distributions are skewed, reflecting the fact that it's very unlikely to get 10 samples that are all very close together, and therefore have low variance. It turns out that there is a theoretical distribution (the χ^2 distribution) that describes the shape of sample variance from the normal distribution.

Median (third column) The main point to draw from this column is that we can use the simulation method to produce a distribution for any statistic, regardless of how easy it would be to calculate a theoretical distribution for it.

As we will see later, we could generate the sampling distribution of the mean and the variance analytically rather than by simulation. However, it is not always possible to compute sampling distributions of the desired statistics analytically, and we can always run statistical simulations.

14.4 The distribution of the sample mean of large samples

The distribution of the sample mean A particularly common statistic of interest is the sample mean. It therefore makes sense to understand how the distribution of the sample mean depends on the distribution from which we sample and the number of samples we take.

We've already seen in Figure 14.4 that the sampling distribution of the mean of 10 items from a normal distribution is itself a normal distribution, though with smaller variance. However, the sample mean distributions for an exponential distribution in our simulation, was not normal. We can repeat the simulation experiments for the three distributions, but with larger sample sizes of $n = 1000$ and $n = 10000$ (Figure 14.5). What we see is remarkable: *the distributions of the sample means are all normal, regardless of whether they came from a normal, uniform or exponential distribution*. Perhaps less remarkably, we also see that as the sample size gets larger the distributions get narrower.

These simulations and observations give us the intuition for two very important statistical laws that apply to many non-normal distributions, as well as normal ones:

- The Central Limit Theorem
- The Law of Large Numbers

Central Limit Theorem Here is an informal statement of the **Central Limit Theorem** (CLT):

The distribution of the mean or sum of a random sample of size n drawn from any distribution will converge on a normal distribution as n tends to infinity.

In the case of the sample mean, its expected value is the same as the mean of the population distribution, and its expected variance is a factor of n lower than the population variance.

In the case of the sample sum, its expected value is the same as the product of the sample size n and the expected value of the distribution, and its expected variance is n times the variance of the population distribution.

We denote the expected variance of the mean $\sigma_{\bar{X}}^2$ and we denote the standard deviation of the mean $\sigma_{\bar{X}}$, called the **standard error of the mean**, often abbreviated as SEM. It's important to note that the SEM is *not* the same as the standard deviation of the original distribution. According to the statement above, an estimate of the SEM is:

$$\hat{\sigma}_{\bar{X}} = \frac{\sigma}{\sqrt{n}}. \quad (14.2)$$

We can verify that this statement holds in the case of sampling a mean in Figure 14.5 by computing the means and SEM from the simulations and comparing with the expected values of μ (population mean) and $\sigma_{\bar{X}} = \sigma/\sqrt{n}$.

The Swain versus Alabama jury selection example demonstrates the CLT applied to a total $T_0 = \sum_{i=1}^n X_i$, where $X_i = 1$ indicates a Black member of the population was selected, and $X_i = 0$ indicates non-Black. The distribution is a Bernoulli distribution with population mean $\mu = p = 0.26$, the probability of picking a Black person. We can see from Figure 14.3 that the mean of the total is $n\mu = 100 \times 0.26 = 26$, and the variance is approximately $\sigma_{T_0}^2 \approx n\sigma^2 = np(1-p) = 19.24$, as expected for a Bernoulli distribution, giving a standard deviation of 4.38. Furthermore, the distribution is approximately normal.

The Law of Large Numbers Here is an informal statement of the **Law of Large Numbers**:

In the limit of infinite n , the expected value of the sample mean $\bar{X}$ tends to the population mean μ and the variance of the sample mean $\bar{X}$ tends to 0.

Note that sometimes the Law of Large Numbers is referred to as the “law of averages”. This can lead to confusion. The law of averages is sometimes called the “Gambler’s fallacy”, i.e. the idea that after a run of bad luck, the chance of good luck increases. If the events that are being gambled on are independent of each other (e.g. successive tosses of the same coin), the probability of a head will be the same regardless of how many tails have preceded it.

In the second row of Figure 14.5 we can see that the distribution for $n = 10000$ is narrower than the distribution for $n = 1000$, and that the sample means converge on the population means. The Law of Large Numbers says that we could, in principle, continue this process by choosing an n as large as we would like to make the variance as small as desired.

Formal statement of the Central Limit Theorem (*Modern Mathematical Statistics with Applications* 6.2) Let $X_1, \dots, X_n$ be a random sample from a distribution with mean μ and variance σ^2 . Then, in the limit $n \rightarrow \infty$ the standardised mean $(\bar{X} - \mu)/(\sigma/\sqrt{n})$ and standardised total $(T_0 - n\mu)/(\sqrt{n}\sigma)$ have a normal distribution. That is

$$\lim_{n \rightarrow \infty} P\left(\frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \leq z\right) = P(Z \leq z) = \Phi(z)$$

and

$$\lim_{n \rightarrow \infty} P\left(\frac{T_0 - n\mu}{\sqrt{n}\sigma} \leq z\right) = P(Z \leq z) = \Phi(z)$$

where $\Phi(z)$ is the cumulative distribution function (cdf) of a normal distribution with mean 0 and s.d. 1. Thus, when n is sufficiently large, $\bar{X}$ has an approximately normal distribution with mean $\mu_{\bar{X}} = \mu$ and variance $\sigma_{\bar{X}}^2 = \sigma^2/n$, and the distribution of T_0 is approximately normal with mean $\mu_{T_0} = n\mu$ and variance $\sigma_{T_0}^2 = n\sigma^2$. We can also say that the standardised versions of $\bar{X}$ and T_0 are **asymptotically normal**.

Formal statement of the (Weak) Law of Large Numbers (*Modern Mathematical Statistics with Applications* 6.2)

Let $X_1, \dots, X_n$ be a random sample from a distribution with mean μ and variance σ^2 . As the number of observations n increases, the expected value of the sample mean remains $E[\bar{X}] = \mu$, but the expected variance $V[\bar{X}] = E[(\bar{X} - \mu)^2] \rightarrow 0$. We say that “ $\bar{X}$ converges in mean square to μ ”.

More formally, the probability that the difference between the sample mean and population mean is greater than an arbitrary value ε is

$$P(|\bar{X} - \mu| \geq \varepsilon) \leq \frac{\sigma^2}{n\varepsilon^2}$$

for any value ε . Thus, as $n \rightarrow \infty$, the probability approaches 0, regardless of the value of ε . A proof of this statement relies on **Chebyshev’s inequality**, and can be found in *Modern Mathematical Statistics with Applications* 6.2.

Note that this is the statement of the Weak Law of Large Numbers. There is also a strong law, which has somewhat more stringent requirements on convergence. All distributions that obey the strong law also obey the weak law, but some distributions only obey the weak law and some obey neither law. A discussion of this topic is beyond the scope of this course; the distributions that do not obey the distribution tend to be “weird”, e.g. having infinite variance.

Frequentist versus Bayesian statistics

You may have heard of the difference between Frequentist and Bayesian statistics. The two systems have different philosophical bases, but, in simpler cases, often produce similar results. Roughly speaking, the differences between the two are:

Frequentist The population is a fundamental concept. There is just one possible value of the population mean and variance, i.e. the one that exists in the population. In estimation, we are trying to estimate these quantities, and in hypothesis testing, we are trying to compare our sample with this population.

Bayesian A fundamental concept is the model of the likelihood of the data given parameters (such as the mean). The parameters themselves are uncertain. Conceptually, the population itself is generated from the model, so a number of combinations of parameters and luck may have generated the particular value of (say) the mean observed in a population. Before we have seen any data, we have an initial idea about the distribution of the parameters (the prior). The inference process involves using the data to update this prior distribution to give a distribution of the parameters given the data.

For around a century, there has been controversy about which approach is best. Broadly speaking, we will be using Frequentist approaches in this course. At the level we are working at here, it will give very similar results to Bayesian approaches. The important thing is to understand the meaning and interpretation of our inference.

Related Python Lab: Randomness, sampling and simulations

<https://github.com/Inf2-FDS/FDS-S2-01-randomness-sampling-simulations>

In this lab you will learn about functions that generate random numbers, sampling, and how they can be used to run statistical simulations. By the end of this lab you should be able to:

- sample from probability distributions using numpy functions
- sample from discrete sets of items
- run statistical simulations to compute the distribution of a statistic
- identify samples of convenience and problems with them, and
- code more complex statistical processes so that they can be run/sampled from.

Chapter 15

Estimation

Recommended reading

- *Modern Mathematical Statistics with Applications*, Sections 7.1

15.1 Point estimation

Estimation In the chapter on [Descriptive statistics](#), we made the distinction between populations and samples. In many cases we may want to know a property of the population, and it may be only feasible to collect data for a sample of that population. For example, we may wish to know the mean (or median) weight of the population of Scottish wildcats, or the voting intentions of a population. In these examples a population of finite size exists, even if we don't know exactly how large the population is.

There are other populations that are less well-defined: for example the weight of cheese added to a pizza by a pizza-chef varies from pizza to pizza. The population of pizzas created by the chef is never really complete – what we are interested in is the mean and standard deviation of the distribution of the weight of cheese that chef adds to the pizzas. It seems reasonable to assume that the cheese weight is normally distributed, since (a) it is a continuous quantity and (b) there is no particular reason to think that the distribution is skewed. The normal distribution is a model of the data.

Parameters In both of these examples we call the population mean and standard deviation **parameters**. In the case of the pizza chef, the mean and standard deviation are parameters of normal distribution, which describe its centre and width. In some distributions, e.g. an exponential distribution $p(x) = \exp(-\lambda x)$ (for $x > 0$), the mean and standard deviation are not separate parameters; they are both equal to the inverse of the parameter λ .

Estimation problems The Oxford English Dictionary defines **estimation** as “the process of forming an approximate notion of (numbers, quantities, magnitudes, etc.) without actual enumeration or measurement.” In other words, we would like to get an approximate idea of population parameters without looking at the entire population. There are two main estimation problems:

1. What is the best way of using the sample to construct a **point estimator** for each population parameter?
2. How do we construct a **confidence interval** to indicate how accurate we expect that point estimator to be?

The answers to these questions will depend on the distribution of the data, and is quite a complex area. We will give an overview of some of the issues here, but not go into depth.

Generic notation for parameters and estimators To help with making some definitions general, we will refer to a generic parameter by the Greek letter ϑ and its estimator by $\hat{\vartheta}$. We'll also use the hat notation for the specific parameters. For example, $\hat{\mu} = \bar{X}$ means “the point estimator of the population mean μ is the sample mean $\bar{X}$ ”.

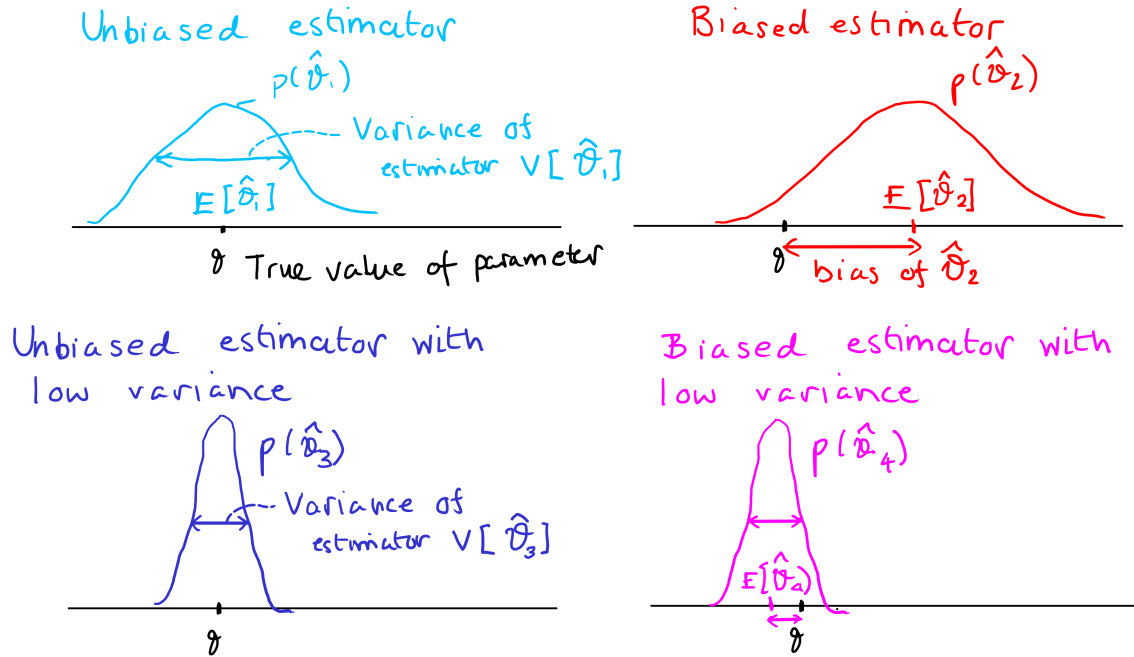


Figure 15.1: Schematic diagram of bias and variance. In each plot the true value of the parameter is indicated by the black point ϑ . The estimator $\hat{\vartheta}_1$ (top left, light blue) is an unbiased estimator, meaning that its expected value is equal to ϑ ; the distribution is centred on the true value ϑ . In the top right plot (red) $\hat{\vartheta}_2$ is a biased estimator; its expected value $E[\hat{\vartheta}_2]$ is greater than the true value, and the difference between the two is the bias. In the bottom left plot, the distribution of the estimator $\hat{\vartheta}_3$ (dark blue) is unbiased and has lower variance, meaning that it is more tightly clustered on the true value. In the bottom right plot, $\hat{\vartheta}_4$ is a biased estimator, but has lower variance than $\hat{\vartheta}_1$. It may be preferable, as it has lower mean squared error (MSE).

More than one estimator for a parameter In some cases, we can have more than one estimator for a parameter. For example, both the mean and the median are estimators of the centre μ of a symmetric distribution, so we can write $\hat{\mu} = \bar{X}$ as well as $\hat{\mu} = \tilde{X}$.

i Capture-recapture method and estimator

Suppose we want to estimate the number of squirrels N in a population. We can do this with a clever method called capture-recapture:

1. Capture n of the squirrels, tag them so that they can be identified if caught again, then release them.
2. Wait for the squirrels to move around.
3. Recapture K of the squirrels and record the number k of these recaptured squirrels that have tags.
4. The estimator of the number of squirrels in the population is

$$\hat{N} = \frac{nK}{k}$$

This should work if the capturing and recapturing processes are random. If this is the case, the expected proportion of tagged squirrels in the whole population n/N is equal to the proportion in the recaptured sample k/K , hence the estimator.

15.2 Estimation bias and variance

Bias and variance Estimators have two properties: bias and variance, which we've summarised graphically in Figure 15.1.

In general, when we estimate a parameter ϑ using an estimator $\hat{\vartheta}$ we will be wrong. We define the error of any particular estimation as $\hat{\vartheta} - \vartheta$. We define the **bias of an estimator** as:

$$\text{bias} = E[\hat{\vartheta} - \vartheta] = E[\hat{\vartheta}] - \vartheta \quad (15.1)$$

The bias tells us, on average, by how much the estimate is too high or too low. If, on average, an estimator gets the right result, i.e. $E[\hat{\vartheta}] - \vartheta = 0$ for *any* value of ϑ we say that the estimator is **unbiased**.

We would like each estimate of the parameter to be as close to the true value of the estimator as possible. A measure of close we can expect estimates to be to the true value is the **mean squared error of an estimator**, which is defined as:

$$\text{MSE} = E[(\hat{\vartheta} - \vartheta)^2] \quad (15.2)$$

It turns out that the mean squared error (MSE) can be decomposed into the **variance of the estimator** and the squared bias:

$$\text{MSE} = E[(\hat{\vartheta} - \vartheta)^2] = V[\hat{\vartheta}] + (E[\hat{\vartheta}] - \vartheta)^2 = \text{variance} + (\text{bias})^2 \quad (15.3)$$

We might think that we should always prefer unbiased estimators that have minimal MSE (or, equivalently zero bias and minimal variance). It turns out that for some parameters of some distributions the unbiased estimators do not have minimal MSE, and that by adding bias we can reduce the variance, thereby reducing the MSE. For an example, see Example 7.4 in *Modern Mathematical Statistics with Applications*.

Example: point estimators for the mean of a normal distribution with known variance We'll first consider the simplest, and rather artificial case, namely a random sample of n observations $X_1, \dots, X_n$ from a normal distribution in which we *know* the variance parameter σ^2 independently of the data. An obvious choice for a point estimator of the mean parameter μ is the sample mean, i.e. $\hat{\mu} = \bar{X}$. For a normal distribution, for *any* value of n , the standardised variable

$$Z = \frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \quad (15.4)$$

has a standard normal distribution (this can be proved quite easily). From this we can derive that $E[\bar{X}] - \mu = 0$, which means that the bias is zero. We can also see that

$$\text{MSE} = E[(\bar{X} - \mu)^2] = \sigma^2/n \quad (15.5)$$

Since the bias is 0, the variance is equal to the MSE. Here the mean squared error is the square of the standard error of the mean (SEM), defined in the chapter on [Randomness, sampling and simulation](#). We can see that as we increase the number of samples n , the MSE decreases, which makes sense.

Example of a senseless biased estimator Note that an estimator does not have to be unbiased or have minimal variance. For example, we could try to estimate the mean with $\bar{X} + 1$. There would then be a bias of 1 and the MSE would be higher. This particular addition of bias makes no sense, but there are cases (see *Modern Mathematical Statistics with Applications*, Section 7.1) where it can make sense.

Example of a biased estimator from Machine Learning As noted in the section on [K-fold cross-validation](#), cross-validation can be used to estimate the value of a metric (e.g. accuracy) when a classifier is tested on unseen data. However, if the cross-validation data has been used to choose a hyperparameter, the cross-validated estimate of the metric is biased. We can regard the value of the metric measured from unseen data as the quantity being estimated ϑ , and the cross-validated value of the metric as the estimator $\hat{\vartheta}$.

i Derivation of unbiased estimator of the population variance (not examinable)

We can now understand the reasoning for why the unbiased estimator of the variance σ^2 has an $n - 1$ in the divisor (see [Why the divisor \$n - 1\$ in the sample variance?](#)) The estimator of the mean is:

$$\hat{\mu} = \bar{X} = \frac{1}{n} \sum_{i=1}^n X_i \quad (15.6)$$

The naive estimator of the variance would be:

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \hat{\mu})^2 \quad (15.7)$$

We rearrange this expression to collect terms that comprise two independent random variables (e.g. X_i and X_j) or one random variable squared (i.e. X_i^2):

$$\begin{aligned} \hat{\sigma}^2 &= \frac{1}{n} \sum_{i=1}^n (X_i - \hat{\mu})^2 = \frac{1}{n} \sum_{i=1}^n \left(X_i - \frac{1}{n} \sum_{i=1}^n X_i \right)^2 \\ &= \frac{1}{n} \sum_{i=1}^n \left(\frac{n-1}{n} X_i - \frac{1}{n} \sum_{j \neq i} X_j \right)^2 \\ &= \frac{1}{n} \sum_{i=1}^n \left(\frac{(n-1)^2}{n^2} X_i^2 - 2 \frac{n-1}{n^2} X_i \sum_{j \neq i} X_j + \frac{1}{n^2} \left(\sum_{j \neq i} X_j \right)^2 \right) \\ &= \frac{1}{n} \sum_{i=1}^n \left(\frac{(n-1)^2}{n^2} X_i^2 - 2 \frac{n-1}{n^2} X_i \sum_{j \neq i} X_j + \frac{1}{n^2} \left(\sum_{j \neq i} X_j^2 + \sum_{j \neq i} \sum_{k \neq j, i} X_j X_k \right) \right) \end{aligned} \quad (15.8)$$

We then compute the expectation of this estimator and use the properties of expectations to bring it into a form where we can compare it to the actual population variance $\sigma^2 = E[X^2] - (E[X])^2$:

$$\begin{aligned} E[\hat{\sigma}^2] &= \frac{1}{n} \sum_{i=1}^n \left(\frac{(n-1)^2}{n^2} E[X_i^2] - 2 \frac{n-1}{n^2} E[X_i \sum_{j \neq i} X_j] + \frac{1}{n^2} \sum_{j \neq i} E[X_j^2] + \frac{1}{n^2} \sum_{j \neq i} \sum_{k \neq j, i} E[X_j X_k] \right) \\ &= \left(\frac{(n-1)^2}{n^2} + \frac{n-1}{n^2} \right) E[X_i^2] + \left(-2 \frac{(n-1)^2}{n^2} + \frac{(n-1)(n-2)}{n^2} \right) E[X_i] E[X_j] \\ &= \frac{n(n-1)}{n^2} E[X_i^2] - \left(\frac{(n-1)^2}{n^2} + \frac{n-1}{n^2} \right) E[X_i] E[X_j] \\ &= \frac{n-1}{n} \left(E[X_i^2] - (E[X_i])^2 \right) \\ &= \frac{n-1}{n} \sigma^2 \end{aligned} \quad (15.9)$$

Therefore, this estimator has a bias:

$$E[\hat{\sigma}^2] - \sigma^2 = \frac{n-1}{n} \sigma^2 - \sigma^2 = -\frac{1}{n} \sigma^2 \quad (15.10)$$

indicating that it underestimates the true variance, but amount of the underestimate decreases as n gets larger.

To create an unbiased estimator, we can multiply this estimator by $n/(n-1)$, which gives the estimator:

$$\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \hat{\mu})^2 \quad (15.11)$$

15.3 Standard error

Estimated standard error of an estimator In addition to the point estimate, we would like a measure of how far the point estimate may be from the true value of the parameter. Assuming we have an unbiased (or near unbiased) estimator, we already have a measure that tells us this – it’s the variance of the estimator $V[\hat{\theta}]$, which will be equal (or nearly equal) to the MSE. To give a measure that is in the same units as the mean, we tend to take the square root of the variance, to give us the **standard error of an estimator**, which we denote $\sigma_{\hat{\theta}} = \sqrt{V[\hat{\theta}]}$.

Going back to the example of estimating the mean of a normal distribution with known variance, Equation 15.5 gives the MSE of the estimator. However, the MSE is defined in terms of the population variance σ^2 , which above we assumed that we know – but in real life we only have the *estimate* $\hat{\sigma}^2$ from the sample. However, we can replace any parameters in the formula for the variance with their estimates to give the **estimated standard error of an estimator**, which we denote $\hat{\sigma}_{\hat{\theta}}$.

Relationship between standard error of the mean and standard deviation of the distribution In the chapter on [Randomness, sampling and simulation](#), we encountered the standard error of a particular estimator, namely the standard error of the mean (SEM), denoted $\hat{\sigma}_{\bar{X}}$. More generally, we could denote the SEM as $\hat{\sigma}_{\hat{\mu}}$, since that implies that we could be using an estimator other than the sample mean to estimate the mean parameter.

It is important to be clear about the difference between the terms “standard deviation” and “standard error of the mean”. The standard deviation tells the variability of the population, distribution or sample. If the standard deviation is describing the population or distribution it’s a parameter, and we denote it σ ; if the standard deviation is derived from the sample, it’s a statistic, and we denote it s .

In the artificial case where we know the standard deviation of the population σ independently of the data, we know from the [Central Limit Theorem](#) that the standard error of the mean is related to the standard deviation and the sample size by:

$$\sigma_{\hat{\mu}} = \frac{\sigma}{\sqrt{n}} \quad (15.12)$$

This relationship shows that to make the estimator twice as accurate (i.e. halving the SEM) we need to quadruple the size of the data n . Figure 15.2 uses statistical simulations to demonstrate that the SEM behaves as predicted for both a normal distribution (top row) and exponential distribution (bottom row).

Estimated standard error of the mean In the much more realistic case where we don’t know the standard deviation of the population, the standard error of the mean is $\hat{\sigma}_{\hat{\mu}} = S/\sqrt{n}$, where S is the sample standard deviation. Since S varies depending on the sample, it’s a random variable, and therefore the estimated SEM itself is a random variable.

Figure 15.3 shows the distribution of the SEM in statistical simulations of sampling $n = 10$ or $n = 100$ samples from a normal and exponential distribution. For $n = 10$ it can be seen that the SEM varies a lot around the theoretical value. This means that a particular sample might give us a much higher or lower SEM than the true value. For $n = 100$, the estimated SEM is distributed much more tightly around the theoretical value – it’s therefore safer to use the estimated SEM as a substitute for the true SEM.



Related Workshop: Statistical problems 1

<https://opencourse.inf.ed.ac.uk/inf2-fds/course-materials/semester-2/week-3/task>

The aim of this task and workshop is to apply some of the techniques from inferential statistics, in particular standard errors and confidence intervals.

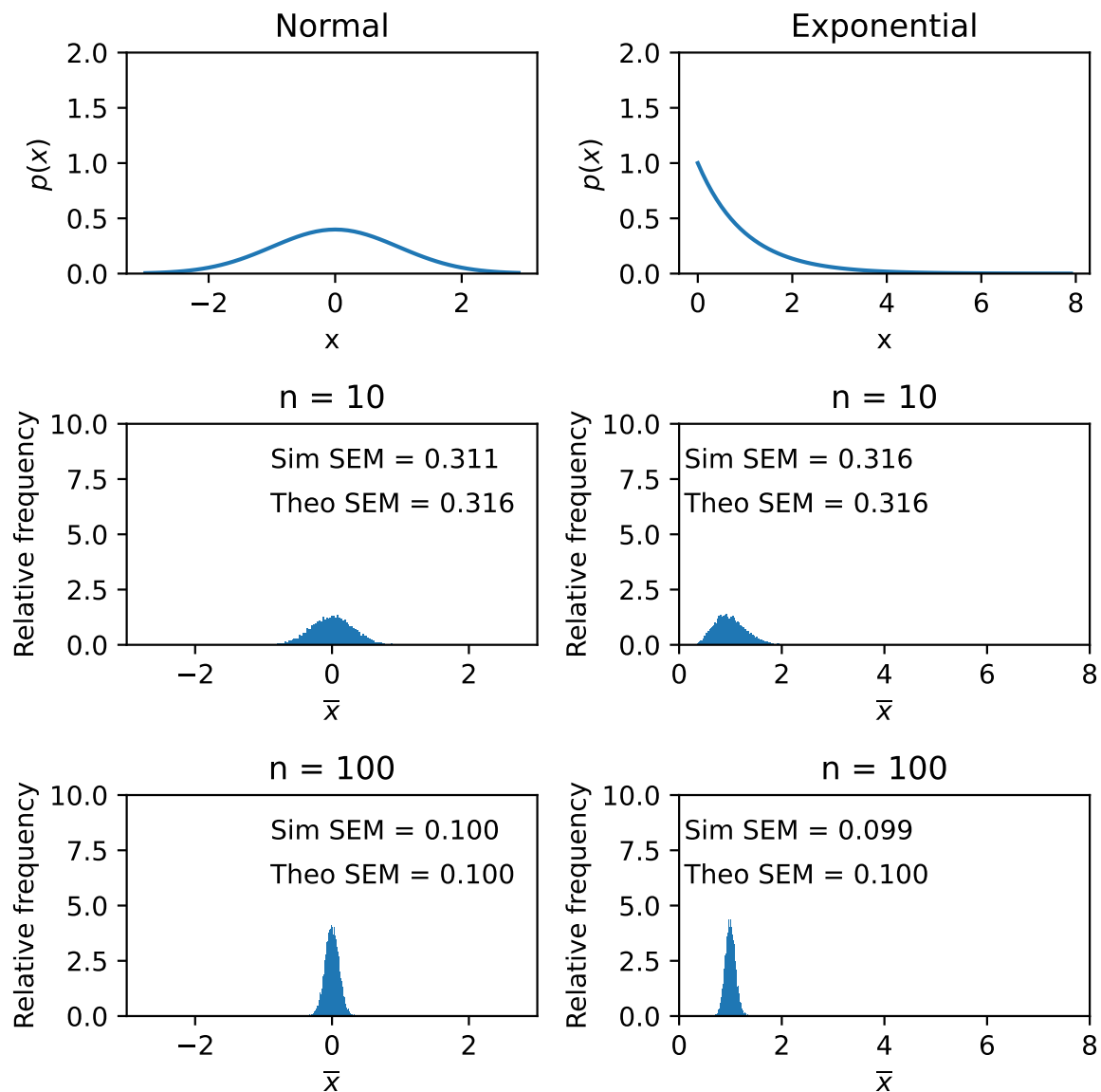


Figure 15.2: Standard error of the mean as the sample size n increases. Top row: normal probability distribution with $\mu = 1$ and $\sigma = 1$ (left) and exponential probability distribution with $\lambda = 1$ (right). Middle row: Histograms showing distribution of the sample mean $\bar{X}$ in 10,000 simulations of sampling 10 random numbers from a normal and exponential distributions. The simulated standard error of the mean (Sim SEM) is calculated from the simulated means. The theoretical standard error of the mean (Theo SEM) is calculated from the formula $\sigma_{\bar{x}} = \sigma/\sqrt{n}$. Bottom row: same as middle row, but with $n = 100$. As n increases, the SEM decreases. Again the simulated SEM is very similar to the theoretical SEM. Note that the distribution of the mean of the samples from the exponential distribution for $n = 100$ is almost normal; for $n = 10$ it is somewhat skewed right.

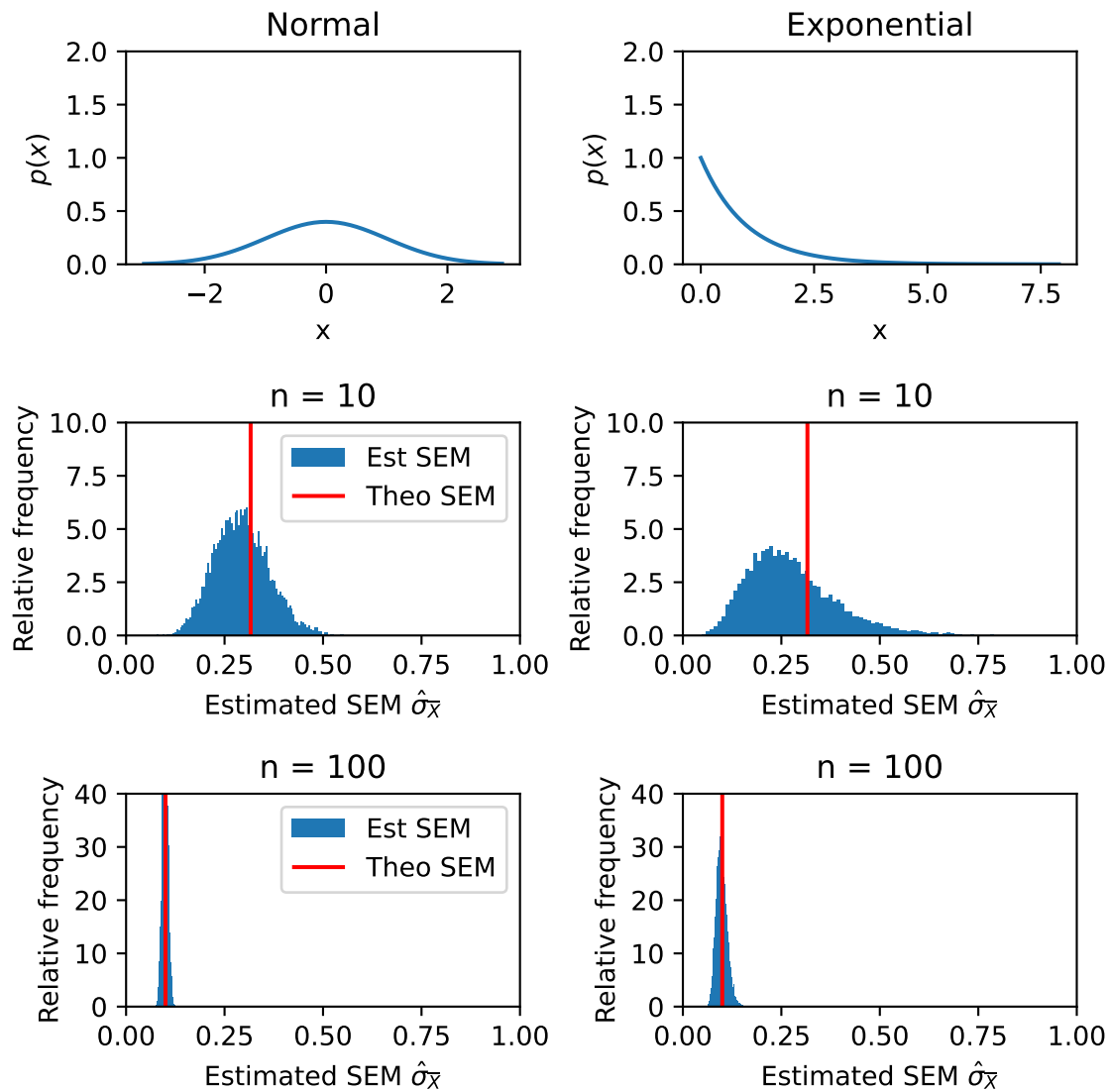


Figure 15.3: Distribution of estimated standard error of the mean. Top row: normal probability distribution with $\mu = 1$ and $\sigma = 1$ (left) and exponential probability distribution with $\lambda = 1$ (right). Middle row: Histograms showing distribution of the SEM calculated by from each of 10,000 simulations of sampling 10 random numbers from a normal and exponential distributions. Bottom row: same as middle row, but with sample size $n = 100$.

Chapter 16

Confidence intervals

Recommended reading

- *Modern Mathematical Statistics with Applications*, Sections 8.1–8.3 and 8.5

16.1 Principle of confidence intervals

Illustration: confidence intervals for the mean From the [Central Limit Theorem](#) we know that for large samples, the distribution of the mean is normal, and that the estimated standard error of the mean should be close to the standard error of the mean. We can then ask “if we looked at an interval around our estimate for the mean, how often would the true value be contained in that interval”?

Figure 16.1 gives an illustrated answer to this question. Each blue horizontal line corresponds to one sample of size n from a population, and shows a range of estimates for the population mean based on that sample – in other words a **confidence interval**. We can see that the true value of the mean (black vertical line) is contained in most of the intervals, but not all of them.

Size of confidence interval We have chosen the length of the intervals to ensure that, if we carried on estimating the mean and the interval, about 95% of intervals would contain the true mean. To determine this length, we use the z critical values of the standardised normal distribution with zero mean and variance 1 (Figure 16.2), which we refer to as the **z -distribution**. We define the **z critical value** z_α as the value of z in a normal distribution which has the area α under the curve to its right. If we want the intervals to contain the true mean 95% of the time, we need to make sure that the mean is within the central 95% of the distribution. This implies that we need 2.5% of the area under the curve to the right of the upper bound, so we look up $z_{0.025}$ in a statistical table or a function in a stats package and find that $z_{0.025} = 1.96$ – we will show how to do this later. The z critical value of 1.96 tells us that the length of the lines on the side of each estimate of the mean should all be 1.96 times the standard error of the mean (SEM).

We may want to be more or less certain of whether the mean is contained in a confidence interval. In this case we can look up the z critical value for our chosen level of confidence. We can also decide to express the confidence interval in terms of the multiples of the SEM. For example, a confidence interval of plus or minus one SEM corresponds to a 68% confidence interval.

Reminder It is worth remembering that these simulations are artificial in the sense that we can repeat many samples. In real life we only get one sample, which does or does not contain the true value – but we don’t know.

Looking up a z critical value To look up a z critical value, you can use the python `scipy` package. For example to find $z_{0.2}$ you would use:

```
from scipy.stats import norm
alpha = 0.2
print(norm.isf(alpha))
```

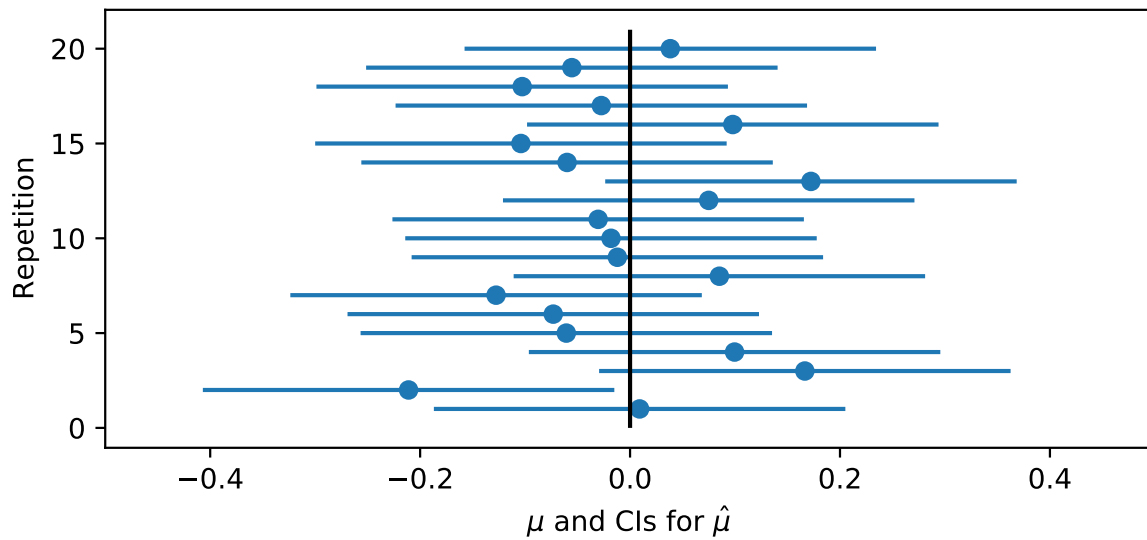


Figure 16.1: Principle of confidence intervals. We repeat a simulation using a sample size of $n = 100$ to estimate the sample mean of a normal distribution with mean 0 and standard deviation 1. The black vertical line indicates the true mean, the blue dots indicate the sample means, and the blue horizontal lines indicate the 95% confidence intervals obtained in each of the 20 repetitions. It can be seen that 19 of the confidence intervals do contain the population mean, but one of them does not.

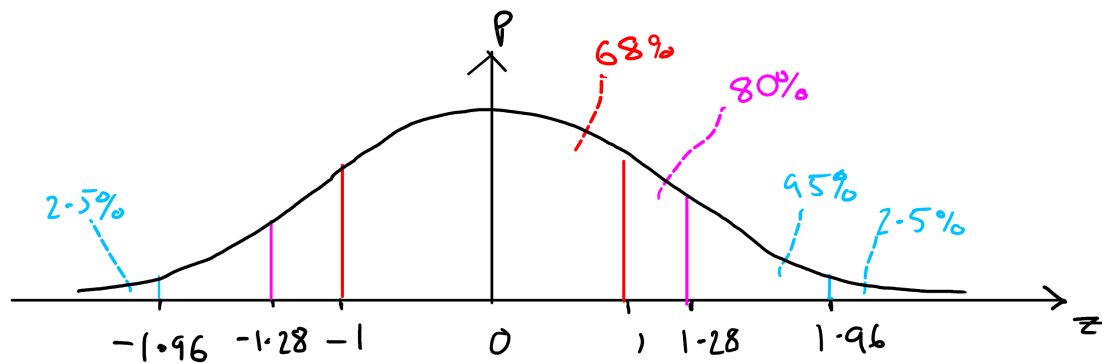


Figure 16.2: Confidence intervals of the z -distribution. The intervals containing various amounts of probability mass under a standardised normal distribution with mean 0 and variance 1 are shown. The 95% confidence interval (blue) is $[-1.96, 1.96]$ and has 2.5% of the probability mass in each tail. The 80% confidence interval is $[-1.28, 1.28]$. The amount of probability mass contained in one standard deviation is 68%. In general for a confidence interval of $100(1 - \alpha)\%$, the upper and lower boundaries are determined by the z critical value $z_{\alpha/2}$. E.g. With the 95% confidence interval, $\alpha = 0.05$ and there is $\alpha/2 = 2.5\%$ of the area of the curve above the upper boundary of the confidence interval.

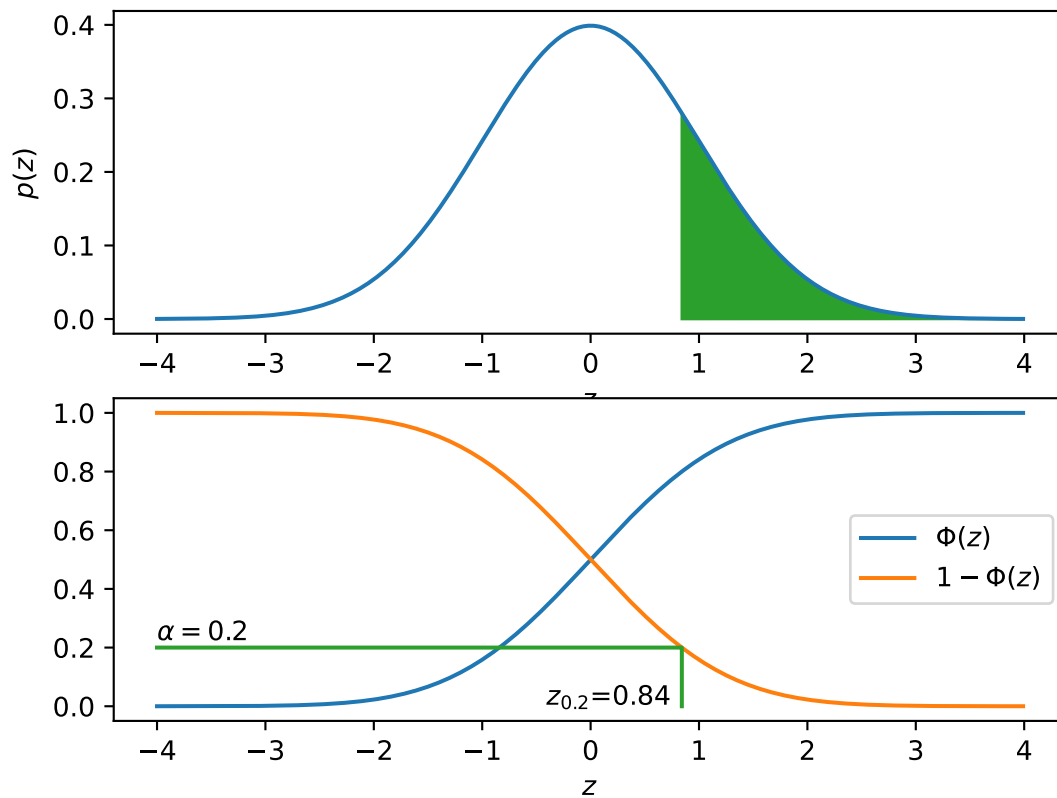


Figure 16.3: Concept of the z critical value. Top: z_α is the value of z such in a normal distribution such that the area under the curve to the right of z_α (green) is equal to α . i.e. $\alpha = \int_{z_\alpha}^{\infty} p(z) dz$. Bottom: the blue curve shows the cumulative distribution function $\Phi(z) = \int_{-\infty}^z p(z) dz$. The orange curve shows the “survival function” $\text{sf}(z) = 1 - \Phi(z)$. The survival function of z is exactly the area to the right of z under the pdf. Therefore, we want to look up the inverse survival function to determine z_α from α , as indicated by the green lines.

Table 16.1: Abbreviated table of critical values for t and z distributions.

α v	0.100	0.050	0.025	0.010	0.005	0.001
1	3.078	6.314	12.706	31.821	63.657	318.309
2	1.886	2.920	4.303	6.965	9.925	22.327
3	1.638	2.353	3.182	4.541	5.841	10.215
4	1.533	2.132	2.776	3.747	4.604	7.173
5	1.476	2.015	2.571	3.365	4.032	5.893
6	1.440	1.943	2.447	3.143	3.707	5.208
7	1.415	1.895	2.365	2.998	3.499	4.785
8	1.397	1.860	2.306	2.896	3.355	4.501
9	1.383	1.833	2.262	2.821	3.250	4.297
10	1.372	1.812	2.228	2.764	3.169	4.144
20	1.325	1.725	2.086	2.528	2.845	3.552
30	1.310	1.697	2.042	2.457	2.750	3.385
40	1.303	1.684	2.021	2.423	2.704	3.307
∞	1.282	1.645	1.960	2.326	2.576	3.090

The function name `isf` stands for **inverse survival function**. As illustrated in Figure 16.3, it's the inverse of one minus the cumulative distribution function (cdf).

You can also look up z critical values in statistical tables, such as the ones in the appendices of *Modern Mathematical Statistics with Applications*. Table 16.1 shows an abbreviated example of such a table. The final row (labelled ∞ , for reasons to be explained later on) contains the z critical values for the values of α shown in the table header. The meaning of the rows will be explained later (see [Confidence intervals for the mean from small samples](#)).

16.2 Definition of confidence intervals

Definition of confidence intervals We define a **confidence interval** as an interval $(\hat{\vartheta} - a\hat{\sigma}_{\hat{\vartheta}}, \hat{\vartheta} + b\hat{\sigma}_{\hat{\vartheta}})$ that has a specified chance $1 - \alpha$ of containing the parameter, and where the positive numbers a and b defining the lower and upper bounds of the interval depend on α . The smaller α is, the larger the values of a and b can be for the statement to hold. A common value for α is 0.05 (i.e. 5%), which gives a 95% confidence interval. However, we could set $\alpha = 0.2$, which would give a narrower 80% confidence interval. Often a and b are equal, but we have given them distinct symbols for full generality.

We can express the definition in terms of a probability statement as follows:

$$P\left(\hat{\vartheta} - a\hat{\sigma}_{\hat{\vartheta}} < \vartheta < \hat{\vartheta} + b\hat{\sigma}_{\hat{\vartheta}}\right) = 1 - \alpha \quad (16.1)$$

In this probability statement, the upper and lower bounds of the interval are random variables, since they are based on the estimators and the estimated standard error, which are themselves random variables derived from the sample.

Expression in terms of random variable in fixed interval We can rearrange the definition of the confidence interval in terms of a standardised variable $(\hat{\vartheta} - \vartheta)/\hat{\sigma}_{\hat{\vartheta}}$:

$$P\left(-b < \frac{\hat{\vartheta} - \vartheta}{\hat{\sigma}_{\hat{\vartheta}}} < a\right) = 1 - \alpha \quad (16.2)$$

Because this standardised variable is derived from the sample, it fits our definition of a statistic. Furthermore, it is composed of *two* statistics, the estimator $\hat{\vartheta}$ and the estimated standard error $\hat{\sigma}_{\hat{\vartheta}}$.

16.3 Method of estimating confidence interval for the mean of a large sample

Methods of estimating confidence intervals There are two main methods of estimating confidence intervals:

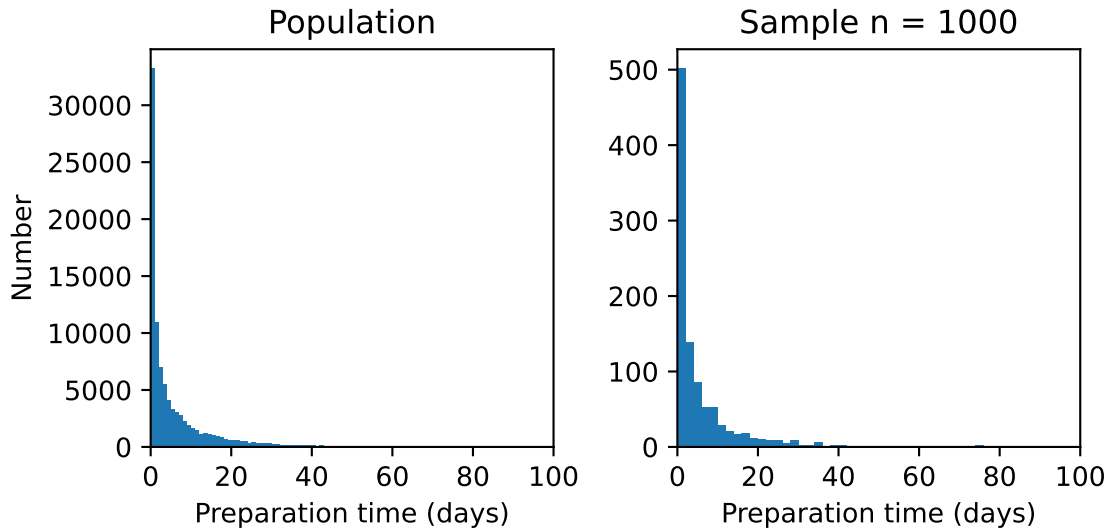


Figure 16.4: Distribution of time from making a reservation to the reservation time (“preparation time”) in restaurants using the “air” booking system in Japan in the period January 2016–April 2017.

1. Under some assumptions about the distribution of the data X_i and the number of samples n we can derive the distribution of $(\hat{\vartheta} - \vartheta)/\hat{\sigma}_{\hat{\vartheta}}$, which will then tell us the values of $-b$ and a at the $100\alpha/2$ th centile and the $100(1 - \alpha/2)$ th centile.
2. More generally we can use a type of statistical simulation called a bootstrap estimator to derive the confidence interval.

We’ll demonstrate the first approach by continuing with our simplified example of a normal distribution with known parameters. In the following section we’ll then cover the bootstrap estimator.

Example: confidence interval for the mean of a normal distribution with known variance In the example of sampling from a normal distribution with known population variance σ introduced in the section on [Standard error](#), the standard error of the mean is $\hat{\sigma}_{\hat{\mu}} = \sigma/\sqrt{n}$. Because the population variance σ is known, it’s not a random variable, and therefore the SEM $\hat{\sigma}_{\hat{\mu}}$ isn’t a random variable either. The standardised variable in Equation 16.2 is therefore

$$\frac{\hat{\vartheta} - \vartheta}{\hat{\sigma}_{\hat{\vartheta}}} = \frac{\hat{\mu} - \mu}{\hat{\sigma}_{\hat{\mu}}} = \frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \quad (16.3)$$

and only contains one random variable, $\bar{X}$. This makes it quite easy to deal with, since we know that this distribution is a standard normal distribution, so we can define a 95% confidence interval by setting a and b to be values at which the cumulative distribution function (cdf) is equal to 2.5% (= $\alpha/2$) and 97.5% (= $1 - \alpha/2$). In this case the values $a = b = 1.96$ satisfy these conditions. Generally we set a and b symmetrically, so that there is equal weight in the “tails” of the distribution (Figure 16.2).

Confidence intervals for the mean of a large sample The central limit theorem states that the distribution of the sample mean of a “large” sample from any distribution should be normal. How large the sample needs to be depends on the distribution, but Figure 15.2 demonstrates that sample means of $n = 100$ samples from an exponential distribution already appear to fairly normally distributed with the SEM as predicted to be the standard deviation of the exponential distribution divided by $\sqrt{n}$. This means that we can use the procedure above to find confidence intervals.

Furthermore, for large samples, the sampling distribution of the estimated standard error of the mean $\hat{\sigma}_{\bar{X}}$ narrows so much that the estimate that we take from the sample is likely to be very close to the theoretical value $\sigma_{\bar{X}}$ (Figure 15.3). So for a large sample Equation 16.2 becomes

$$\frac{\hat{\vartheta} - \vartheta}{\hat{\sigma}_{\hat{\vartheta}}} = \frac{\hat{\mu} - \mu}{\hat{\sigma}_{\hat{\mu}}} \approx \frac{\bar{X} - \mu}{S/\sqrt{n}} \quad (16.4)$$

Table 16.2: Summary statistics of the population and a sample of preparation times, generated by the pandas describe function.

	Population	Sample
count	92378.00	1000.00
mean	8.30	8.06
std	25.65	27.72
min	0.00	0.00
25%	0.21	0.17
50%	2.08	1.96
75%	7.88	6.92
max	393.12	364.96

Confidence intervals for the mean of an empirical distribution Up until now, we have considered estimating parameters from theoretical probability distributions, such as the normal distribution or the exponential distribution. We'll now consider how we can estimate the parameters of an empirical distribution, i.e. real-world data, from a sample of that distribution – remember that the CLT works for *any* distribution.

💡 Confidence interval of restaurant order times

As an example, we will take the population of times between making a reservation and the time of the reservation itself in Japanese restaurants using the “air” booking system. The full population contains 92378 times (Figure 16.4, left) and we've created a random sample of 1000 of these times (Figure 16.4, right). In real life, if we had the full set of data, there would not be any point in creating this random sample of times, but we do so here to demonstrate how well we can estimate confidence intervals. From now on imagine that the sample of 1000 times is all that we have available to us. It's important to notice that the distribution of the sample resembles the population distribution, even though it is rougher.

Table 16.2 shows the summary statistics for the population and the sample. We can see that the estimates for the mean, standard deviation and centiles from the sample are all similar to the true population values. From the table we can see that the population mean is $\mu = 8.30$ days and the sample mean is $\bar{x} = 8.06$ days. The sample mean would be different if we'd happened to have taken a random different sample.

From the summary statistics from the sample, we have $\bar{x} = 8.06$ days and the standard deviation $s = 27.72$ days. Our estimator for the mean is $\hat{\theta} = \bar{x} = 8.06$ days. Our estimator for the standard error of the mean is $\hat{\sigma}_{\hat{\theta}} = s/\sqrt{n} = 27.72/\sqrt{1000} = 0.88$ days. The 95% confidence interval for the mean in days is therefore $(\hat{\theta} - 1.96\hat{\sigma}_{\hat{\theta}}, \hat{\theta} + 1.96\hat{\sigma}_{\hat{\theta}}) = (6.34, 9.78)$.

Reporting confidence intervals When reading scientific papers, there are various ways of reporting confidence intervals:

- $M=8.06$, $CI=6.34-9.78$. Here “M” stands for mean and “CI” stands for confidence interval.
- 8.06 ± 1.72 (95% confidence interval)
- 8.06 ± 0.88 (± 1 SEM). This is a 68% confidence interval, though the confidence interval isn't specified in terms of area under the curve.
- 8.06 ± 1.76 (± 2 SEM).

16.4 Bootstrap estimation of confidence intervals

Principle of a bootstrap estimator Suppose we want to estimate the standard error of an estimator. In the chapter on [Randomness, sampling and simulation](#), we have already seen distributions of statistics of repeated small and large samples from theoretical distributions. In the case of computing the sample mean from each sample we obtained the standard error of the mean by computing the standard deviation of the sampling distribution. We could have also computed the standard error in the variance or the median of a



Figure 16.5: Bootstrapping: Baron Münchhausen pulls himself and his horse out of a swamp by his pigtail. Public domain image from [Wikipedia's article on bootstrapping](#).

particular number of samples, by computing the standard deviation of the distributions of these statistics shown in Figure 14.4.

What happens if we don't have a theoretical distribution, but we do have a sample of data that we think is representative of the population from which it's drawn? We appear to be in an impossible position, since we don't have a theoretical distribution to sample from.

A **bootstrap estimator** resolves the problem of a lack of theoretical distribution. It treats the sample that we have available as a population, and resamples from it to give the sampling distribution of the estimator. From the sampling distribution we can compute the standard error of the estimator. It feels as though we shouldn't be able to treat the sample as a population, but it works because if we have a large enough sample, the distribution of the sample will resemble the population itself.

The name "bootstrap estimator" arises because it appears to do something physically impossible, such as "pulling ourselves up by our own bootstraps". (Equivalently we could pull ourselves up by our pigtail, Figure 16.5).

Bootstrap procedure for finding a confidence interval for the mean We will start with a large sample n from the data, which has a mean $\bar{x}$. By large, we mean large enough that the sample resembles the population distribution. Of course, this is not possible to know exactly, so the larger the better. We decide to take B bootstrap samples. Common numbers are 1000 or 5000, or 10000. More samples are generally better, but bootstrapping can be computationally expensive, and fewer samples can also give reasonable results.

Here is the procedure:

- For j in $1, \dots, B$
 - Take sample x^* of size n from the sample *with replacement*
 - Compute the sample mean of the new sample $\bar{x}_j^*$
- To compute the bootstrap confidence interval, we find the centiles of the distribution at $100\alpha/2$ and $100(1 - \alpha/2)$. We can do this by arranging the sample means $\bar{x}_j^*$ in order from lowest to highest, and pick $\bar{x}_j^*$ at $k = \alpha(B + 1)/2$ to be the lower end of the CI and pick $\bar{x}_j^*$ at $k = B - \alpha(B + 1)/2$ to be the upper end of the CI.
- We can also compute the bootstrap estimator of the variance of the mean:

$$s_{\text{boot}}^2 = \frac{\sum_{j=1}^B (\bar{x}_j^* - \bar{x})^2}{B - 1}$$

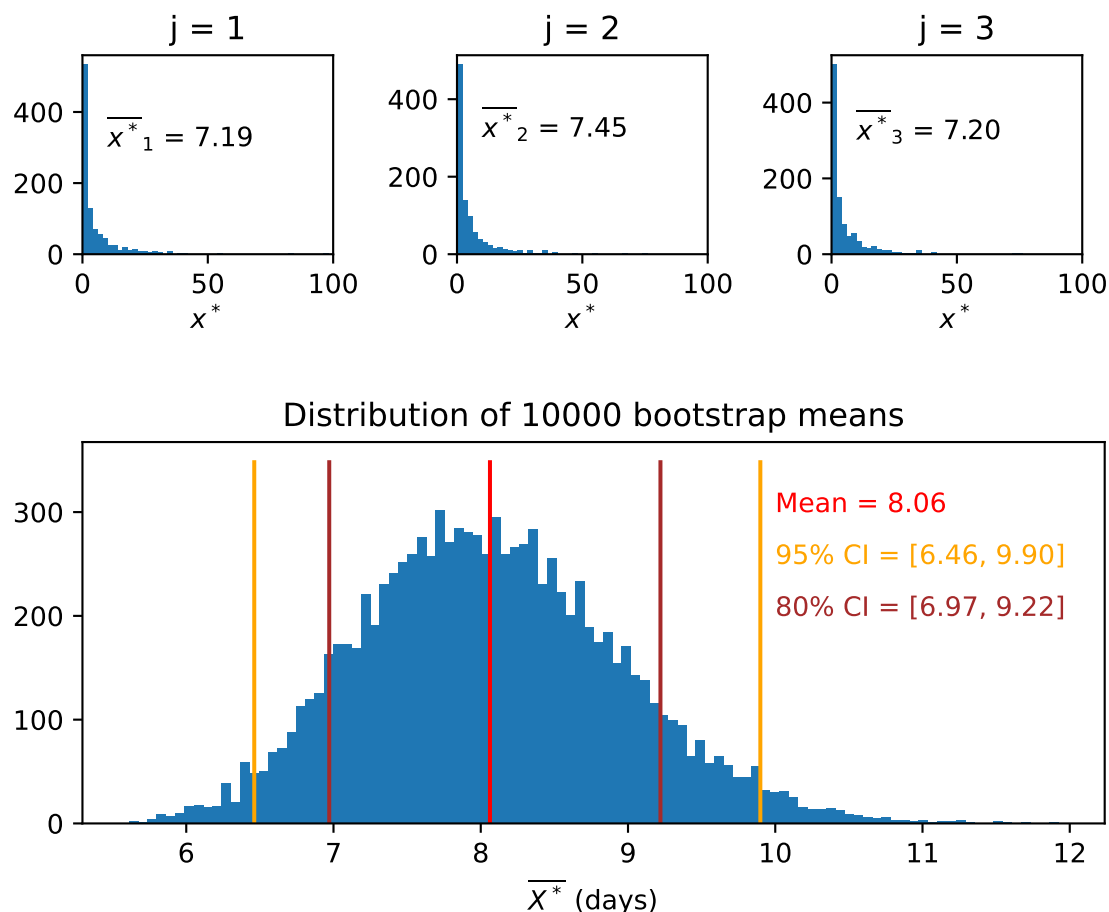


Figure 16.6: Demonstration of bootstrap mean applied to restaurant reservation time data (Figure 16.4). The top row shows the distributions obtained from the first 3 of 10000 bootstrap samples. Although the distributions are similar to each other, they are not exactly the same, and the sample mean of each is different. The bottom figure is the distribution of all 10000 of these bootstrap sample means. The mean of the original sample is shown, as is the 95% and 80% confidence intervals.

The advantages of the bootstrap procedure are that we can use it for any estimator, e.g. the median, and that we do not need to make any assumptions about the distribution of the estimator.

Example of bootstrap estimator applied to the mean We'll now apply the bootstrap estimator to give us a confidence interval for the mean (Figure 16.6). For each of our 10,000 bootstrap samples, we'll resample 1000 samples *with replacement* from our sample of 1000. Each of these samples will be a distribution (top row of Figure 16.6), from which we can compute the 10,000 bootstrap means. Then we'll plot the distribution of the bootstrap means (bottom row of Figure 16.6) and find the 95% and 80% confidence intervals. In this case we can see both the 95% and 80% confidence intervals contain the *population* mean (8.30, Table 16.2). However, if we replicate the experiment with a different initial random sample of 1000, in around 5% of cases we should expect that the 95% confidence interval does not contain the mean.

Comparison of bootstrap confidence intervals with normal approximation The 95% confidence interval obtained via the bootstrap procedure is (6.46, 9.90) days, which is very similar to the confidence interval obtained by the normal approximation, (6.34, 9.78) days. The bootstrap interval is slightly shifted to the right, suggesting that the normal approximation is quite accurate at a sample size of $n = 1000$.

General formulation of bootstrap estimator A great advantage of the bootstrap is that we can easily apply it to statistics other than the mean. Here is the general procedure for estimating the confidence interval for a generic estimator $\hat{\theta}$:

- For j in $1 \dots B$
 - Take sample x^* of size n from the sample *with replacement*
 - Compute the sample statistic of the new sample $\hat{\vartheta}_j^*$
- Then compute the bootstrap estimator of the variance of the statistic:

$$s_{\text{boot}}^2 = \frac{\sum_{j=1}^B (\hat{\vartheta}_j^* - \hat{\vartheta})^2}{B - 1}$$

- To compute the bootstrap confidence interval, we find the centiles of the distribution at $100\alpha/2$ and $100(1 - \alpha/2)$.

This procedure works well for measures of centrality such as the median, and for the variance. It doesn't work so well for statistics of extremes of the distribution, such as the maximum or minimum.

16.5 Interpretation of confidence intervals

Interpretation of confidence intervals Although we have only computed confidence intervals in a simple artificial example, we are already at a stage where we can consider how to interpret confidence intervals. From Equation 16.1 we can see that confidence intervals are a random interval – whenever we take a new sample, we will end up with a new interval, as illustrated in Figure 16.1. The interpretation (according to the frequentist interpretation of statistics) is that if we performed a long run of experiments (i.e. repeatedly took samples) the parameter (the mean in this case) would be in around 95% of the confidence intervals.

How big should a confidence interval be? Should we choose the 95% confidence interval or the 80% confidence interval? The answer to this question depends on the problem. For example, suppose we have a machine that makes tens of thousands of ball bearings for aircraft jet engines every day. Each ball bearing needs to have a diameter of $2 \pm 0.0001\text{mm}$ for the engine to work safely. We measure the diameter of a sample of the ball bearings every day. Because this is a safety-critical application, we need to have high confidence (say 99.999%) that the ball bearings are in the range $2 \pm 0.0001\text{mm}$. This might require a large sample size, but it's worthwhile because the consequences of getting it wrong could be catastrophic.

On the other hand, suppose we are estimating the number of red squirrels in a population so that we know how much red-squirrel friendly food to put out for them over winter. We might want to leave out a bit more than we expect they need, we're happy to accept a 10% chance that the true number of squirrels might be greater than the upper end of a confidence interval, so we compute the 80% confidence interval, and put out enough food for the number of squirrels at the upper end of the interval. There's a 10% chance that we might not be providing for enough squirrels, but it's not as catastrophic as in the aircraft situation (depending on how much you value red squirrels compared to humans).

Upper and lower confidence bounds In this case, we're not worried about our estimate being too low, so we only need to compute the upper confidence bound – we would quote a mean number of squirrels and an upper limit.

16.6 Confidence intervals for the mean from small samples

Small samples We'll now consider estimating confidence intervals for the mean based on a "small" (usually $n < 40$) sample of data that appears to be distributed normally and whose variance we are not given – we can only estimate it from the data.

For example, suppose we want to estimate the mean weight of a population of female squirrels from a sample of $n = 32$ squirrels (Wauters and Dhondt, 1989). The sample mean is $\bar{x} = 341.0\text{g}$ and the estimated standard error of the mean is $\hat{\sigma}_{\bar{x}} = 3.9\text{g}$.

We can imagine that if we had taken a different sample of $n = 32$ squirrels, we would have found both a different sample mean and a different estimate for the standard error of the mean. Thus, the standardised variable $(\bar{X} - \mu)/\hat{\sigma}_{\bar{X}}$ itself contains *two* random variables, $\bar{X}$ and $\hat{\sigma}_{\bar{X}}$, derived from the sample. As we are estimating the standard deviation, rather than knowing it, the normal approximation for the distribution of the mean begins to break down.

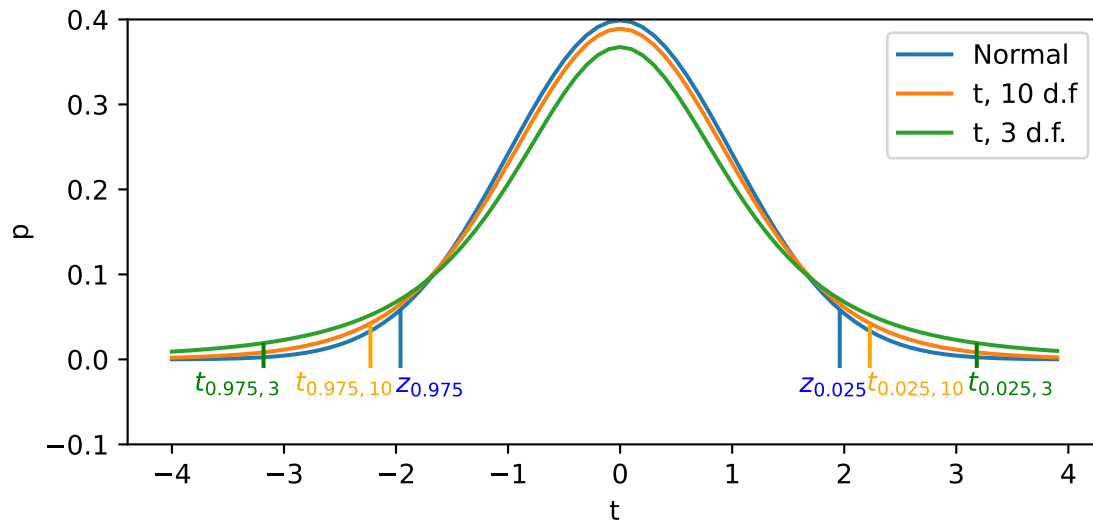


Figure 16.7: The t -distribution for 3 degrees of freedom and 10 degrees of freedom, with normal distribution for comparison. 2.5% t critical values and z critical values are shown.

The t -distribution We could use the bootstrap estimator to estimate confidence intervals. However, in this special case, there is another option. There's a theorem that states that when $X_1, \dots, X_n$ is a random sample of size n from a normal distribution with mean μ , the random variable

$$T = \frac{\bar{X} - \mu}{\hat{\sigma}_{\bar{X}}} \quad (16.5)$$

is distributed as a **t -distribution** with $n - 1$ degrees of freedom, where the t -distribution with ν degrees of freedom has the probability density function depicted in Figure 16.7, which is given by the equation:

$$p_{\nu}(t) = \frac{1}{\sqrt{\pi\nu}} \frac{\Gamma((\nu+1)/2)}{\Gamma(\nu/2)} \frac{1}{(1+t^2/\nu)^{(\nu+1)/2}} \quad (16.6)$$

where $\Gamma(x)$ is a gamma function. We will not prove this theorem here; in *Modern Mathematical Statistics with Applications* Section 6.4 there is the sketch of a proof.

The t -distribution is very similar in shape to the z -distribution: it is bell-shaped, symmetrical, and centred on 0. However, for small numbers of degrees of freedom, the t -distribution has longer tails than the z -distribution. This means that the tails of a t -distribution contain a greater fraction of the weight of the distribution than do the tails in a z -distribution. We define the **t critical value** $t_{\alpha, \nu}$ as the value of t in a t -distribution with ν degrees of freedom which has the area α under the curve to its right.

For small degrees of freedom, the t critical values are considerably bigger than the z critical values of the normal distribution (Figure 16.7). As the number of degrees of freedom increases, the t -distribution approaches a z -distribution. The distribution with 40 degrees of freedom (not shown in the figure) looks very similar to a z -distribution.

Looking up a t critical value To look up a t critical value, you can use the python `scipy` package. For example to find $t_{0.025, 10}$ you would use:

```
from scipy.stats import t
alpha = 0.025
nu = 10
t_cv = t(nu).isf(alpha)
print(t_cv)
```

You can also look up t critical values and z critical values in statistical tables, such as the ones in the appendices of *Modern Mathematical Statistics with Applications*. Table 16.1 shows an abbreviated example of such a table. Each row contains t critical values for degree various levels of α . The final row, with infinite number of degrees of freedom, is the z critical values for these values of α . The full tables include values for more degrees of freedom.

Using the t -distribution to derive a confidence interval The $100(1 - \alpha)$ percent confidence interval around a mean $\bar{x}$ of a sample of n values with estimated SEM $\hat{\sigma}_{\bar{x}}$ derived using a t -distribution is:

$$(\bar{x} - t_{\alpha/2, n-1} \hat{\sigma}_{\bar{x}}, \bar{x} + t_{\alpha/2, n-1} \hat{\sigma}_{\bar{x}}) \quad (16.7)$$

Note that we have used the t critical value $t_{\alpha/2, \nu}$. Here the number of degrees of freedom is one less than the sample size ($\nu = n - 1$). Also, we have divided α by 2 because we are wanting upper and lower bounds to the confidence interval. It might be that we only need an upper bound, as we considered when we were estimating squirrel numbers earlier. In this case we would just quote $\bar{x} + t_{\alpha, n-1} \hat{\sigma}_{\bar{x}}$. This is still a $100(1 - \alpha)$ confidence interval, since the interval from $-\infty$ to the upper bound contains $100(1 - \alpha)$ of the area under the t -distribution.

To continue the squirrel example, suppose we want to find a 95% confidence interval for the weight. The 95% confidence interval implies $\alpha = 0.05$ and $\nu = n - 1 = 31$. We would then look up the $t_{0.025, 31} = 2.040$ and substitute it into Equation 16.7 along with the sample mean and estimated SEM, and then use this to generate the confidence interval, which we could quote as $\hat{\mu} = 341.0 \pm 8.0\text{g}$ (95% confidence interval, $n = 32$). This is a bit wider than the interval we would obtain using the corresponding critical value of a normal distribution $z_{0.025} = 1.96$.

Related Python Lab: Estimation of confidence intervals with the bootstrap

<https://github.com/Inf2-FDS/FDS-S2-02-estimation-bootstrap>

In this lab you will use statistical simulations to undertake bootstrap estimation of confidence intervals. By the end of this lab you should be able to:

- code the bootstrap estimator for a number of estimators
- validate statistical coding by comparing the output of functions with known results
- interpret the output
- compare the output with confidence intervals obtained by other methods

Related Workshop: Statistical problems 1

<https://opencourse.inf.ed.ac.uk/inf2-fds/course-materials/semester-2/week-3/task>

The aim of this task and workshop is to apply some of the techniques from inferential statistics, in particular standard errors and confidence intervals.

Chapter 17

Hypothesis testing and p -values

Recommended reading

- [XKCD comic strip on multiple testing – funny!](#)
- [*A hypothesis is a liability* Yanai and Lercher \(2020\) – thought-provoking and amusing article](#)

17.1 Principle of hypothesis testing

Hypothesis testing helps us to answer yes/no questions, such as “is chocolate good for you?” or “is a jury selection procedure biased?” There are two aspects to hypothesis testing:

1. Deciding on whether a **hypothesis** or **model** is compatible with data from observational studies and randomised experiments.
2. If the hypothesis is compatible with the data, investigating the mechanisms specific to the data, e.g. the biological effect of chocolate on the body or the process by which a jury panel was selected.

In the course we are going to focus on the statistical aspects (Aspect 1), but it’s worth remembering that the question is not answered once we’ve completed this step – it should prompt further investigation of the question rather than ending the inquiry ([Yanai and Lercher, 2020](#)). Furthermore, as [Yanai and Lercher \(2020\)](#) illustrate rather amusingly, it is important to explain data before undertaking hypothesis testing – a good visualisation can reveal features of the data that a hypothesis test can’t.

Method of hypothesis testing At the core of hypothesis testing are the **null hypothesis** and the **alternative hypothesis**:

The null hypothesis H_0 : The claim that we initially assume to be true, formalised as a statistical model. e.g. “The jury panel was chosen by random selection from the population in a district.”

The alternative hypothesis H_a : The claim that is contradictory to H_0 , typically not formalised as a statistical model. E.g. “The jury panel was chosen by some other, unspecified, method.”

The aim of hypothesis testing is to either reject or not reject the null hypothesis. Note that we do not “accept” the null hypothesis as true, we are just saying that it’s not been proved to be false.

Test procedure The procedure to carry out a hypothesis test, which we call the **test procedure**, consists of:

1. Deciding on a **test statistic**, which is a function of the sample data, e.g. the number of Black people in a jury panel.
2. Determining what the distribution of the test statistic would be if it arose from the null hypothesis statistical model.
3. Either:

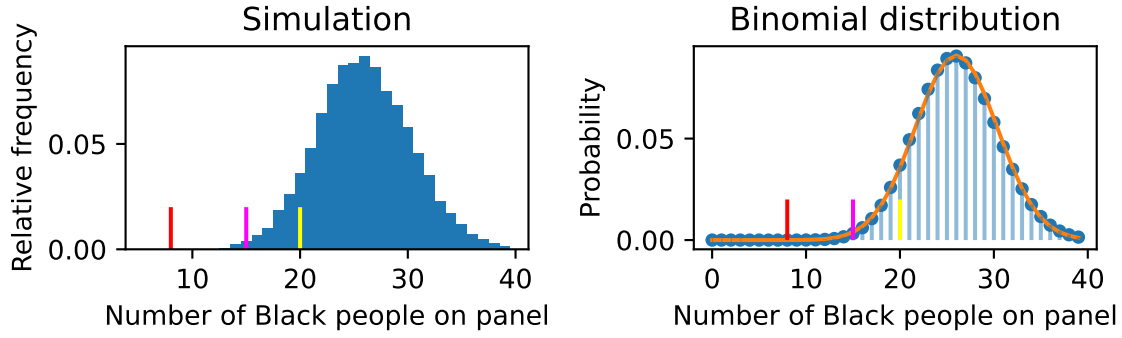


Figure 17.1: Distributions of number of Black people T_0 (test statistic) on a panel of 100 under the null hypothesis that the jury was randomly selected from a population that is 26% Black and 74% non-Black. Left: distribution arising from 10 000 statistical simulations. The red line indicates the number of Black jurors in Swain versus Alabama (1965), the magenta line indicates $t_0 = 15$ and the yellow line indicates $t_0 = 20$. Right: Binomial distribution (blue dots) for $n = 100$ and $p = 0.26$. Normal approximation (orange curve) with $\mu = np$ and $\sigma = \sqrt{np(1-p)}$.

- Deciding on a **rejection region**, i.e. regions of the distribution of the test statistic under H_0 in which we should reject H_0 . Typically, these are the extremities of the distribution. If our test statistic falls into the rejection region, we reject H_0 ; otherwise, we don't reject it.
- Returning a **p-value**, which tells us how compatible the test statistic is with the distribution predicted by chance from H_0 .

Application of test procedure to example In the topic on [Randomness, sampling and simulation](#), we looked at the example of Swain versus Alabama (1965), in which the question was “if 8 Black people were chosen for a jury panel of 100 people, but the fraction of Black people in the population was 26%, does this show bias against Black people?” We found the distribution of the test statistic under the null hypothesis by simulating the null hypothesis model of sampling from a Bernoulli distribution with $P(\text{Black}) = 0.26$. In this case probability theory also tells us that the distribution is a binomial distribution with $n = 100$ and $p = 0.26$. We found that there were no replications in which 8 Black members were chosen (Figure 17.1) – the simulated numbers were always higher.

We did not consider rejection regions or p-values. Since the observed data (8 Black people on the panel; red line in Figure 17.1) were inconsistent with the range of predictions produced by the null hypothesis, it seemed very clear that we should reject the null hypothesis. But what would we have decided if the number of Black people had sat within the distribution of simulated values, e.g. 15 (magenta line) or 20 (yellow line)?

Rejection regions We might want to specify the rejection region as the bottom 5% of the probability mass, i.e. the region that seems unusually low (Figure 17.2, left, region to left of orange boundary). If the observed test statistic falls into that region, we might “reject the hypothesis at the 5% level (one-tailed test)”. We call this a **one-tailed test** because the rejection region occupies only one tail of the distribution. This is justified, as the alternative hypothesis was implicitly “the number of Black people selected is below the number we would have expected by chance”.

If we know the distribution of our null hypothesis model, we can look up statistical tables to determine the boundaries of rejection regions. E.g. in this case, the number n is large enough that we can approximate the binomial distribution with a normal distribution with mean $\mu = np$ and variance $\sigma^2 = np(1-p)$. This means that the standardised statistic

$$Z = \frac{T_0 - \mu}{\sigma} \quad (17.1)$$

is normally distributed. At the edge of the rejection region, this statistic is equal to the z critical value $z_{0.95}$, which has 95% of the probability mass to its right. We can then rearrange Equation 17.1 to find the edge of the rejection region in terms of the original statistic:

$$T_0 = \mu + \sigma z_{0.95} \quad (17.2)$$

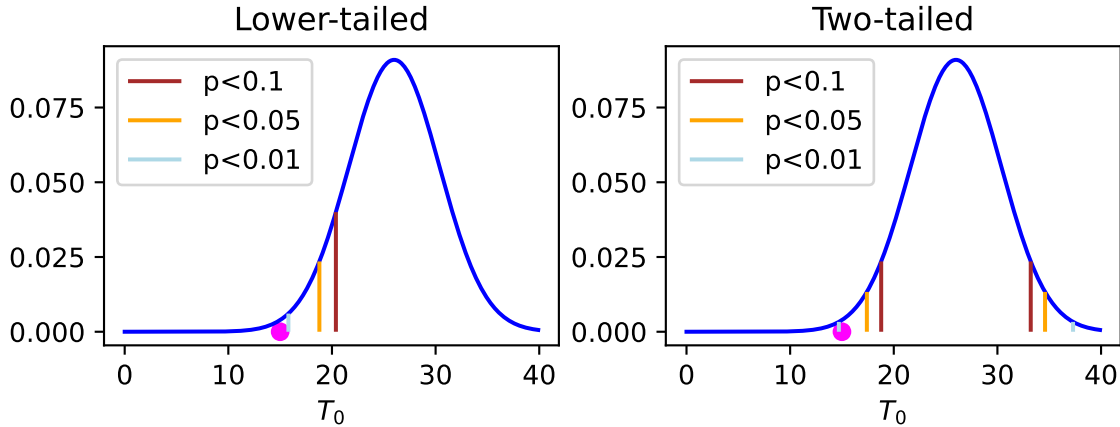


Figure 17.2: Rejection regions. Lower-tailed (left) and upper-tailed (right) rejection regions are shown for the normal approximation to the distribution of the null hypothesis model in the Swain–Alabama example. The observed statistic $t_0 = 15$ is shown with a magenta dot. It lies in the $p < 0.01$ rejection region for a lower-tailed test and in the $p < 0.05$ rejection region for a two-tailed test.

If a test statistic in a hypothesis test is distributed according to a normal distribution, the hypothesis test is sometimes referred to as a “z-test”.

One-tailed and two-tailed tests We could have formulated the alternative hypothesis as “the number of Black people selected is different from (i.e. above or below) the number we would have expected by chance”. In this case we would perform a **two-tailed test** (Figure 17.2, right) by setting the rejection regions to be the bottom 2.5% and the top 2.5% of the probability mass of the distribution. We would “reject the hypothesis at the 5% level (two-tailed test)”.

17.2 p -values

Principle of p -values The principle of p -values is that we set the boundary of the rejection regions to be where the data is, and then report the probability mass in the resulting rejection regions as the p -value .

Determining p -values from statistical simulations Had there been 15 Black people on the panel in Swain versus Alabama (magenta line), a fraction 0.0062 of the 10 000 simulations produced panels with 15 or fewer black members. This would therefore give the p -value $p = 0.0062$, i.e. 0.62%. This certainly calls into question if the observed data is compatible with the null hypotheses.

Suppose that there had been 20 Black people on the jury panel (yellow line in Figure 17.1). The corresponding rejection region is 20 or fewer Black people on the jury. A fraction of 0.101 of the simulations are in this region, so the p -value is $p = 0.101$. We would tend not to reject the null hypothesis at this size of p -value, but this would not mean that the null hypothesis was true.

Sometimes the p -value is reported relative to a round figure rejection region, e.g. in the case with 15 Black people on the jury, $p < 0.01$, indicating that we could “reject the null hypothesis at the 1% level”. However, supplying the actual p -value gives more information than just reporting the rejection region.

Determining p -values from probability distributions Sometimes it is straightforward to compute the probability distribution implied by the null hypothesis. In the Swain versus Alabama example, it is a binomial distribution with $n = 100$ and $p = 0.26$ (Figure 17.1, right). As we are looking at a lower-tailed test, the p -value is the cumulative distribution function of the binomial distribution, cut off at t_0 , the observed number of Black people on the jury panel:

$$P(T_0 \leq t_0) = B(t_0; n, p) = \sum_{t=0}^{t_0} b(t; n, p) \quad (17.3)$$

Table 17.1: P -values computed by various methods for various observed values of t_0 in Swain versus Alabama (1965).

t_0	Simulation	Binomial	Normal
8	0	4.73e-06	2.03e-05
15	0.0067	0.0061	0.0061
20	0.1020	0.1030	0.0857

where $b(t; n, p)$ is the probability of t successes in a binomial distribution with n trials and success probability p ; $B(t; n, p)$ is the corresponding cumulative distribution function (cdf). Stats packages have functions to compute the cdf for various distributions, and the values for the binomial are shown in Table 17.1 along with the simulated values.

Also shown is the normal approximation to the binomial, in which we set $\mu = np$ and $\sigma = \sqrt{np(1-p)}$. The p -values are the values of the normal cumulative distribution function at the standardised value

$$z = \frac{t_0 - \mu}{\sigma} \quad (17.4)$$

Why use rejection regions? The rejection region method works well with printed statistical tables, in which critical values of z and other distributions are available only for particular cut-off values, e.g. 0.01, 0.05. With computer packages it is now possible to define the rejection region relative to the observed data rather than a pre-set cut-off.

Definition of p -value We can define the p -value as follows:

The p -value is the probability, calculated assuming the null hypothesis is true, of obtaining a value of the test statistic at least as contradictory to H_0 as the value calculated from the available sample. (*Modern Mathematical Statistics with Applications*, p. 456)

The whole topic of the interpretation and use of p -values is complex and highly contested. In fact, it took 20 statisticians 2 days and many subsequent days of drafting to produce the American Statistical Association's statement on p -values: [The statement by the American Statistical Association \(Wasserstein and Lazar, 2016\)](#).

What p -values are We quote 2 of the 6 points in the statement here. Firstly, what p -values are:

P -values can indicate how incompatible the data are with a specified statistical model...

The smaller the p -value, the greater the statistical incompatibility of the data with the null hypothesis, if the underlying assumptions used to calculate the p -value hold. This incompatibility can be interpreted as casting doubt on or providing evidence against the null hypothesis or the underlying assumptions. (*ASA Statement on Statistical Significance and P-values*)

In the Swain versus Alabama example where we imagined there were 15 Black people on the jury, the small p -value ($p = 0.0062$) indicates that the data (here 15 Black people on the panel) are quite incompatible with the null hypothesis statistical model (here that Black and non-Black people were drawn from the population at random). The low p -value casts considerable doubt on the hypothesis. Of course the actual data ($t_0 = 8$) has a vanishingly small p -value (Table 17.1).

What p -values are not Secondly, what they are not:

P -values do not measure the probability that the studied hypothesis is true, or the probability that the data were produced by random chance alone.

Researchers often wish to turn a p -value into a statement about the truth of a null hypothesis, or about the probability that random chance produced the observed data. The p -value is neither. It is a statement about data in relation to a specified hypothetical explanation, and is not a statement about the explanation itself. (*ASA Statement on Statistical Significance and P-values*)

	Caucasian	Black/AA	Hispanic	Asian/PI	Other	Total
Population %	54	18	12	15	1	100
Observed panel numbers	780	117	114	384	58	1453
Expected panel numbers	784.62	261.54	174.36	217.95	14.53	1453.00
$\frac{(\text{Observed}-\text{Expected})^2}{\text{Expected}}$	0.03	79.88	20.90	126.51	130.05	357.36

Table 17.2: Alameda County jury panel data. The top row shows the estimated proportions of 5 ethnic groups (Caucasian, Black/African American, Hispanic, Asian/Pacific Islander and Other) in Alameda County. The second row (Observed panel numbers) shows the total number in each group on 11 jury panels from 2009–2010. There was a total of 1453 on the 11 jury panels (final column). The third row (Expected panel numbers) shows the numbers we would expect from each group if the panels had been selected randomly from the population. The final row $\frac{(\text{Observed}-\text{Expected})^2}{\text{Expected}}$ shows the disparity between the observed and expected using this formula. The total disparity is in the final column.

“Statistical significance” A widespread practice in scientific literature is to take p -values of less than $p = 0.05$ as indicating **statistical significance**, i.e. that the null hypothesis should be rejected. Values of less than 0.05 indicate weak evidence against the null hypothesis. Sometimes higher thresholds are used, e.g. $p = 0.01$ and $p = 0.001$. In scientific papers and the output from stats packages you will sometimes see these values indicated with asterisks:

- * means significant at least at the $p < 0.05$ level
- ** means significant at least at the $p < 0.01$ level
- *** means significant at least at the $p < 0.001$ level

There is no “correct” answer about what the right level of significance is. The $p < 0.05$ value was suggested in a paper by the statistician Ronald Fisher¹, who invented the hypothesis test, but it simply seemed “convenient” to him for his purposes. As in the discussion on confidence intervals ([How big should a confidence interval be?](#)), the value we choose to use may depend on the application. For example, we would demand a very low p -value when testing the null hypothesis that a new drug has no effect on the death rate of patients. We might accept a slightly higher p -value for the hypothesis that it has no positive effect on symptoms. In less mission-critical scientific applications, a higher p -value will be acceptable.

17.3 Testing for goodness of fit to a model

Multiple categories In the example so far, there have been just two categories: Black and non-Black. In 2010 the North California branch of the American Civil Liberties Union (ACLU) **investigated** the numbers of Caucasian, Black/African American, Hispanic, Asian/Pacific Islander and Other people on jury panels in Alameda County. They found the data shown in the first two rows of Table 17.2.

We want to test the following null and alternative hypotheses, which are essentially the same as for the case with two categories:

The null hypothesis H_0 : The jury panels were chosen by random selection from the population in a district.

The alternative hypothesis H_a : The jury panels were chosen by some other, unspecified, method.

With two categories, it’s easy to see that the number of Black people could be a test statistic. But in this case, there are 4 numbers that describe the outcome of any simulation (we can always compute the number in the 5th category if we know the total number and the numbers in 4 categories). We can’t have 4 test statistics, so we need to create a statistic that indicates the disparity between the observed and expected outcomes.

Suppose we call the population proportions of each of k groups p_i and the observed numbers in each group n_i . The total number sitting on jury panels is $n = \sum_i n_i$. We can compute the numbers we would expect to be on jury panels as np_i (third row of table). One measure of disparity would be the sum of the

¹Fisher studied under Pearson, and developed a huge body of modern statistics. He also edited the *Annals of Eugenics* and had controversial views on race.

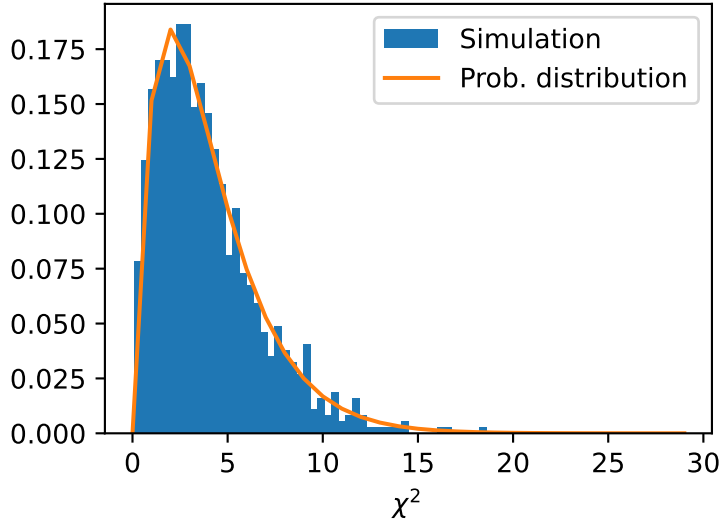


Figure 17.3: Distribution of χ^2 for jury panel selection in Alameda County. Simulations shown in blue and theoretical χ^2 distribution with 4 degrees of freedom shown in orange.

squared differences:

$$\sum_{i=1}^k (n_i - np_i)^2$$

This looks at the *absolute* squared differences between the expected and observed values for each category. If we expected $np_1 = 100$ in one category and observed $N_1 = 95$, this expected-observed pair would contribute 25 to the sum. A difference of $np_2 = 10$ (expected) and $N_2 = 5$ (observed) would also contribute 25 to the sum. However, in *relative* terms, the difference between the first expected-observed pair is 5%, whereas in the second pair it is 50%.

This motivates us to look at the scale the disparity measure by dividing by the expected number in each category, to create a statistic that we call **chi-squared**, written using the Greek symbol χ^2 :

$$\chi^2 = \sum_{i=1}^k \frac{(n_i - np_i)^2}{np_i} \quad (17.5)$$

The components of the χ^2 statistic are seen in the final row of Table 17.2, as is the value of $\chi^2 = 357.36$ for the observed values (in the “Total” column).

Statistical simulation We can now run a statistical simulation to generate the expected distribution of χ^2 . For each repetition we simulate the numbers in each category by drawing from a multinomial distribution with parameters n and p_i . We then compute and store χ^2 for that simulation, which gives us the simulated distribution shown in blue in Figure 17.3. We can immediately see that the observed value of $\chi^2 = 357.36$ is off the scale of the graph, indicating that it has a much bigger value than is compatible with the null hypothesis, so we reject the null hypothesis.

Chi-squared distribution It turns out that, providing every expected value np_i is greater than 5, the χ^2 statistic is distributed approximately according to a χ^2 probability distribution with $k - 1$ degrees of freedom, shown in orange in Figure 17.3. The fit between the probability distribution the simulated distribution is clear.

Goodness-of-fit Large values of χ^2 statistic (i.e. the upper tails of the distribution) indicate a poor **goodness-of-fit** between the model and the data. χ^2 tests therefore tend to be **upper-tailed**. However, the statistic had a very low χ^2 , we might be suspicious that the data had been fiddled with.

²Sometimes you may see the letter X used instead of χ .

	Female	Male	Total		Population 1	Population 2	Total
Depressed	30	12	42	Category 1	n_{11}	n_{12}	$n_{1\bullet}$
Not depressed	2048	1663	3711	Category 2	n_{21}	n_{22}	$n_{2\bullet}$
Total	2078	1675	3753	Total	$n_{\bullet 1}$	$n_{\bullet 2}$	$n_{\bullet\bullet}$

Table 17.3: Left: Contingency table of the number of depressed and not depressed people in a population of females and males; data based on a prospective study [Bornioli et al. \(2020\)](#). Right: General symbolic version of the two-way contingency table. There are I rows and J columns. The number of items falling into a cell in the i th row and j th column is denoted n_{ij} . The total in the i th row is denoted $n_{i\bullet} = \sum_{j=1}^J n_{ij}$, the total in the j th column is $n_{\bullet j} = \sum_{i=1}^I n_{ij}$. The grand total is $n_{\bullet\bullet} = \sum_i n_{i\bullet} = \sum_j n_{\bullet j}$.

	Female	Male		Population 1	Population 2
Depressed	23.25	18.75	Category 1	$\hat{e}_{11}$	$\hat{e}_{12}$
Not depressed	2054.75	1656.25	Category 2	$\hat{e}_{21}$	$\hat{e}_{22}$

Table 17.4: Expected numbers in contingency table, in example (left) and in symbols from Equation 17.7 (right).

The χ^2 statistic can be used to assess the goodness-of-fit of many types of model and data, not just this proportion example. If we find a χ^2 with a p -value greater than desired cut-off, this suggests that we should not reject the model.

Testing for independence with two-way contingency tables We may have multiple populations (e.g. males and females) and multiple categories (e.g. depressed or not depressed). We can arrange these in a **two-way contingency table** (Table 17.3).

We want to test the null hypothesis that being depressed is independent of if you are male or female. In other words $P(X = x, Y = y) = P(X = x)P(Y = y)$. Using a notation similar to Table 17.3 (right), we can write this probability as $p_{ij} = p_{i\bullet}p_{\bullet j}$, where $p_{i\bullet}$ is the marginal probability of an item being in category i and $p_{\bullet j}$ is the marginal probability of an item being in category j . Our best estimates of the marginal probabilities are

$$p_{i\bullet} = \frac{n_{i\bullet}}{n_{\bullet\bullet}} \text{ and } p_{\bullet j} = \frac{n_{\bullet j}}{n_{\bullet\bullet}} \quad (17.6)$$

Therefore the best estimates of the number of in each cell are

$$\hat{e}_{ij} = n_{\bullet\bullet}p_{ij} = \frac{n_{i\bullet}n_{\bullet j}}{n_{\bullet\bullet}} \quad (17.7)$$

The χ^2 statistic is computed as

$$\chi^2 = \frac{(\text{Observed} - \text{Expected})^2}{\text{Expected}} = \sum_i \sum_j \frac{(n_{ij} - \hat{e}_{ij})^2}{\hat{e}_{ij}} \quad (17.8)$$

In this case it is 4.433.

We assume that the numbers of depressed and non-depressed, and males and females are fixed. In general, if there are I rows and J columns in the table there are $(I - 1)(J - 1)$ degrees of freedom. In this case there is therefore only 1 degree of freedom; specifying n_{11} (or any other cell) allows us to compute the values of all the other cells. We therefore look up the cumulative distribution function of χ^2 with 1 degree of freedom, to find that $p < 0.035$, so this is significantly different from independence at the 5% level.

17.4 Issues in hypothesis testing

Type I and Type II errors Regardless of whether we use a one-tailed test or a two-tailed test, if the null hypothesis were true, there is 5% chance that an observed test statistic in the rejection region might really have arisen by chance. By rejecting H_0 , we would have made a **Type I error**: rejecting the null hypothesis when it is true. To reduce the risk of making a Type I error, we could make the rejection region smaller,

e.g. the bottom 1%. However, we would also have increased the chance of making a **Type II error**: not rejecting the null hypothesis when it is false.

There is no right answer about what size of rejection region to use – it depends on what the consequences of Type I versus Type II errors are.

Decisions based on confidence intervals

Scientific conclusions and business or policy decisions should not be based only on whether a p -value passes a specific threshold. (ASA Statement, point 5). (Wasserstein and Lazar, 2016)

$p \leq 0.05$ does not mean that the false; it is one point in a spectrum. However, it is often seen as the “holy grail” of scientific research.

Data dredging, data snooping and p -hacking It is very tempting to try out many experiments in order to get a p -value of less than 0.05. However, the more experiments are run, the more chance there is of Type I errors – i.e. rejecting the null hypothesis when it is true.

Data dredging, data snooping or p -hacking is the practice of rerunning experiments or selecting subsets of datasets until a statistically significant result is achieved. It is a form of cherry-picking, in which data is effectively selected with a bias towards interesting results. It is harder to publish negative results than positive results in academic journals, so there is an incentive to data dredge. Some statistically significant results in the literature will be Type I errors, which makes it important to replicate experimental results.

The ASA statement says:

Proper inference requires full reporting and transparency.

P -values and related analyses should not be reported selectively. Conducting multiple analyses of the data and reporting only those with certain p -values (typically those passing a significance threshold) renders the reported p -values essentially uninterpretable. Cherry-picking promising findings, also known by such terms as data dredging, significance chasing, significance questing, selective inference, and “ p -hacking,” leads to a spurious excess of statistically significant results in the published literature and should be vigorously avoided...(ASA Statement on Statistical Significance and P -values)

Multiple testing Suppose we undertake multiple tests on the same dataset is problematic, and find that one of the tests is significant. As we increase the number of tests, the probability of a Type I error increases (XCKD comic in reading). If we undertake 20 tests, there’s a 0.95^{20} chance of not having a Type I error, and therefore a $1 - 0.95^{20} = 0.64$ chance of a type I error. There are ways to compute more stringent cut-offs in these cases, for example the Bonferroni correction.

Chapter 18

A/B testing

Recommended reading

- *Modern Mathematical Statistics with Applications*, Chapter 10

18.1 The principle of A/B Testing

A/B testing A/B testing is a method for assessing how changes to design of a system affect user behaviour. Figure 18.1 shows a hypothetical example, in which two versions of a website are presented to users selected at random. Group A gets the version with the blue button and group B gets the version with the green button with the inviting arrow. The numbers of users clicking-through is then measured. A/B testing is a form of randomised control trial (See Table 6.2 in [Data collection and statistical relationships](#)). There are a number of commercial systems to implement A/B testing.

Statistical Questions in A/B testing

1. Is A *significantly* better than B?
2. How much better is A than B?

The first question corresponds to the hypothesis testing task introduced in the chapter on [Hypothesis testing and p-values](#). The second question corresponds to the estimation task, introduced in the chapter on [Confidence intervals](#). In this case both tasks can be carried out either using statistical simulations applied to the data, or using theoretical sampling distributions. In this section we will use statistical simulations, and later on (in [Large sample theory of A/B testing](#)), we will apply large sample theory.

Generating confidence intervals for A/B learning using statistical simulations Let's imagine that we present the two versions of the page to group A and to group B the same number of times, n . We find that group A clicks through on 70% of occasions and group B on 72%. We'll call the underlying proportions of users that click through that we are trying to estimate p_A and p_B , and we will define the difference that we are trying to estimate:

$$d = p_A - p_B \quad (18.1)$$

The difference d is positive when A is better than B. We can address the question of how much better than A is than B by finding a point estimate of d – the larger d the better A is than B. We can address the question of if A is significantly better than B by finding a confidence interval.

The natural point estimators for p_A , p_B and d are:

$$\hat{p}_A = \frac{n_A}{n} \quad , \quad \hat{p}_B = \frac{n_B}{n} \quad \text{and} \quad \hat{d} = \hat{p}_A - \hat{p}_B \quad (18.2)$$

where n_A and n_B are the actual numbers clicking through from A and B.

To find the confidence interval, we can use a statistical simulation to generate the sampling distribution of d , assuming the underlying proportions in populations A and B are given by the point estimates $\hat{p}_A$ and $\hat{p}_B$. The routine to generate the sampling distribution of d looks like:

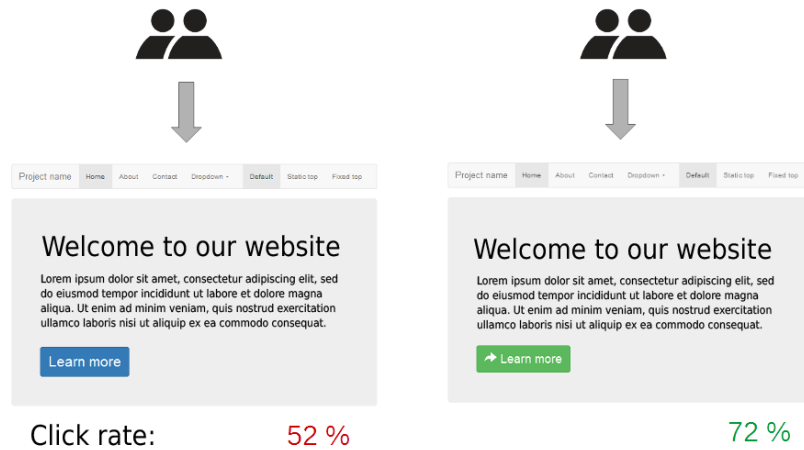


Figure 18.1: A/B Testing. Group A is shown the web page on the left; group B the one on the right. Image credit: [Maxime Lorient, Wikimedia, CC SA 4.0](#).

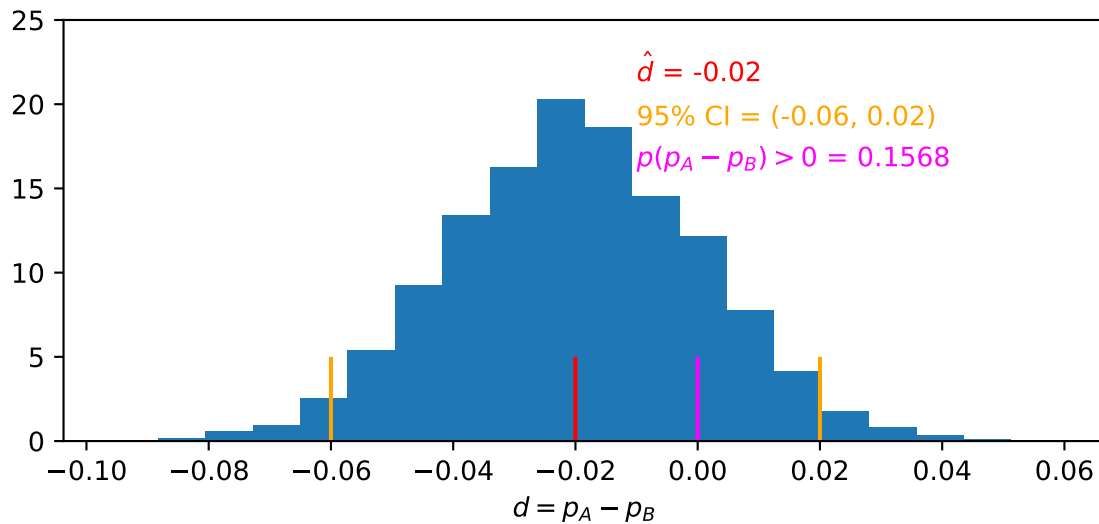


Figure 18.2: Statistical simulation of A/B test with $n = 1000$, $p_A = 0.70$ and $p_B = 0.72$.

- For j in $1, \dots, k$
 - Sample n_A^* from binomial distribution with parameters n and $\hat{p}_A$
 - Sample n_B^* from binomial distribution with parameters n and $\hat{p}_B$
 - Compute and store difference in proportions

$$d_j^* = n_A^*/n - n_B^*/n$$

- Plot the distribution of d^* and compute the desired quantities

The result is shown in Figure 18.2. The point estimate $\hat{d} = -0.02$, suggesting that B is better than A. However, the 95% confidence interval is $(-0.06, 0.02)$, which contains the value $d = 0$, suggesting that A and B could be equally effective.

Note that the area to the left of $\hat{d} = 0$ in the sampling distribution is about 85% and the area to the right is about 15%. We interpret this as meaning that there is an 85% chance that version B is better than version A – but there is still a 15% chance that it isn't.

Reframing the statistical simulation of A/B testing as a bootstrap simulation

This statistical simulation is equivalent to a bootstrap simulation. Suppose we represent the data in Group A as a set of binary numbers $x_1, \dots, x_n$, where 1 represents clicking through and 0 represents not clicking through – the total number of 1s will be n_A . To generate each sample x^* of a bootstrap simulation we resample n times from the data $x_1, \dots, x_n$ and compute the test statistic n_A^* , the number of 1s in x^* . This resampling is equivalent to drawing n_A^* from a binomial distribution, as carried out in the statistical simulation described in this section.

Exercise: a hypothesis test for A/B learning using statistical simulations

It's also possible to approach this A/B problem as a hypothesis test. We leave it as an exercise to formulate the problem in this way and write a statistical simulation.

18.2 Increasing certainty in A/B testing

Getting a more certain result To be more certain, we could keep the test running. But, assuming that the population proportions are $p_A = 0.70$ and $p_B = 0.72$ how many presentations n of both versions of the page would we need to have a chance of (say) only 1% that A is better than B?

The brute force approach to find out how big n should be is to run the statistical simulation again, with different values of n (Figure 18.3). As n increases, the distribution, and the 95% confidence interval gets narrower. By $n = 10000$, we can see that the upper end of the 99% confidence interval is now less than 0. The chance of the underlying proportion p_A being higher than p_B is around 0.0012. We can therefore say that a 99.999% confidence interval is $(-\infty, 0)$.

When to stop sampling Suppose we had collected our first $n = 1000$ A and B responses in 2 hours on a Monday afternoon. We're quite excited by the result, and reckon that we need to keep it running up to $n = 10000$ in order to be 99.999% certain. This will probably take us to Tuesday afternoon, we'll then write a report for the boss, and be done by Wednesday. What could possibly go wrong?

We've made a hidden assumption that every period of the week is like a Monday afternoon. What if people prefer blue to green in the evening? What if the Monday afternoon demographic is older, but the weekend demographic is younger? To avoid selection bias (see [Data collection and statistical relationships](#)), we may wish to collect at least a full week of data to check that our result really is robust – a week's worth of data should mean that any day- or time-specific effects are eliminated, or at least greatly reduced.

18.3 Large sample theory of A/B testing

As with the confidence intervals and hypothesis testing for the sample mean, we can use a theoretical approach to determine confidence intervals or undertake hypothesis testing when doing A/B testing. We have already determined that the estimators for the population proportions are:

$$\hat{p}_A = \frac{n_A}{n} \quad ; \quad \hat{p}_B = \frac{n_B}{n} \quad (18.3)$$

Now we are interested in estimating the difference $d = p_A - p_B$ between our population proportions. An unbiased estimator of d is:

$$\hat{d} = \hat{p}_A - \hat{p}_B \quad (18.4)$$

Supposing the population proportions are p_A and p_B , we expect the number of successes in n trials to be binomially distributed, with the standard deviations of n_A and n_B being:

$$\sigma_{n_A} = \sqrt{np_A(1 - p_A)} \quad ; \quad \sigma_{n_B} = \sqrt{np_B(1 - p_B)} \quad (18.5)$$

Dividing through by n and replacing p_A and p_B by their estimates, we get the estimated standard errors of the estimators $\hat{p}_A$ and $\hat{p}_B$:

$$\hat{\sigma}_{\hat{p}_A} = \sqrt{\frac{\hat{p}_A(1 - \hat{p}_A)}{n}} \quad ; \quad \hat{\sigma}_{\hat{p}_B} = \sqrt{\frac{\hat{p}_B(1 - \hat{p}_B)}{n}} \quad (18.6)$$

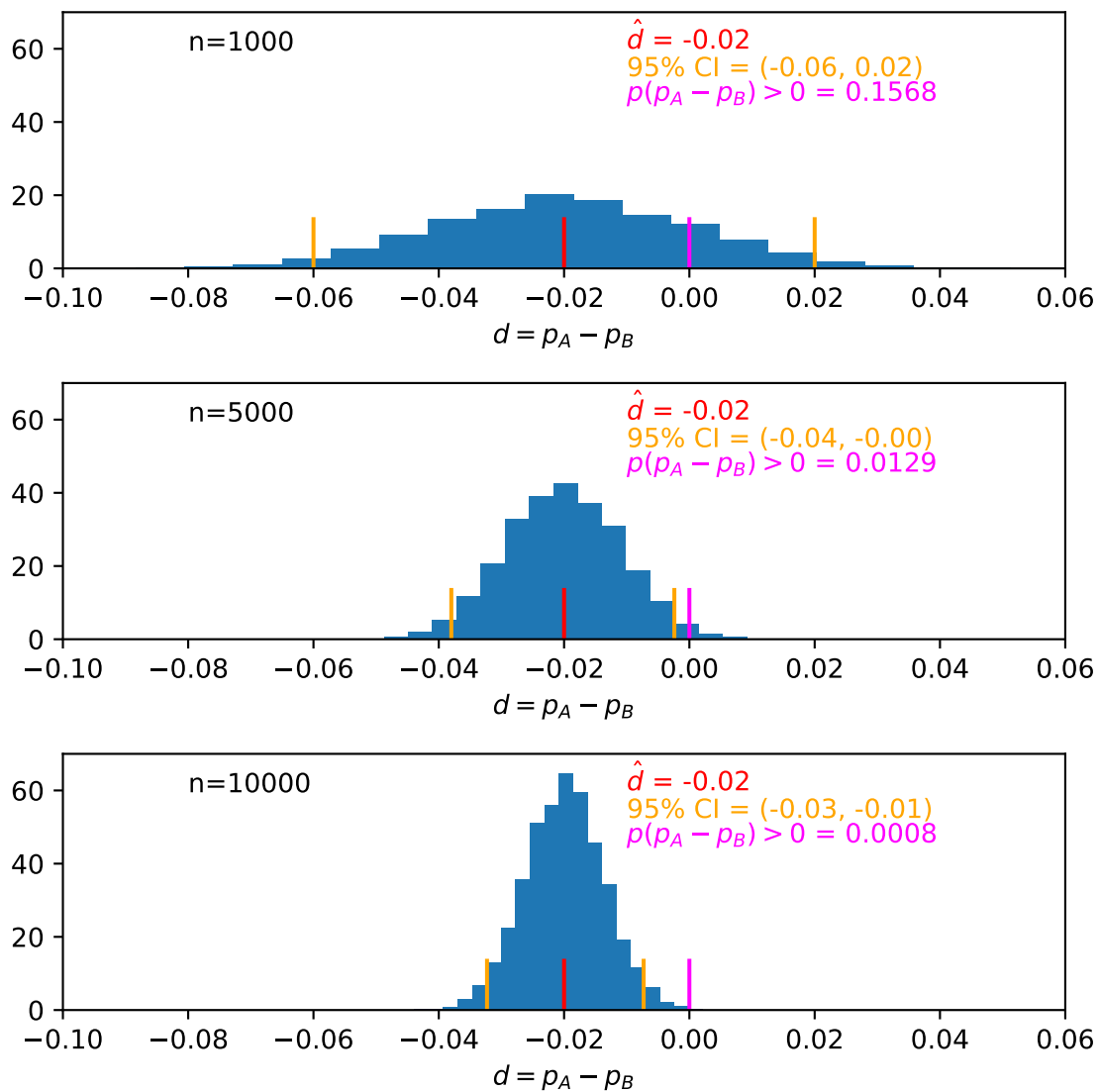


Figure 18.3: Statistical simulation of A/B test with $p_A = 0.70$ and $p_B = 0.72$, and varying numbers of n .

Since the samples from A and B are independent, the variance of the estimator of the difference in proportions $\hat{d}$ is equal to the sum of the variances of $\hat{p}_A$ and $\hat{p}_B$. We take the square root to get the standard error of the estimator $\hat{d}$:

$$\hat{\sigma}_{\hat{d}} = \sqrt{\hat{\sigma}_{\hat{p}_A}^2 + \hat{\sigma}_{\hat{p}_B}^2} = \frac{\sqrt{\hat{p}_A(1 - \hat{p}_A) + \hat{p}_B(1 - \hat{p}_B)}}{\sqrt{n}} \quad (18.7)$$

We'll assume that n is large, in which case the Central Limit Theorem applies, and we can assume that there is little variance in the estimated standard error of $\hat{d}$. We can therefore assume that the statistic

$$Z = \frac{\hat{d} - d}{\hat{\sigma}_{\hat{d}}} = \frac{(\hat{p}_A - \hat{p}_B) - (p_A - p_B)}{\sqrt{(\hat{p}_A(1 - \hat{p}_A) + \hat{p}_B(1 - \hat{p}_B))/n}} \quad (18.8)$$

is normally distributed. We can then use the z -distribution to calculate confidence intervals.

Worked example We'll use figures we used for the statistical simulation to find the 95% confidence interval theoretically:

$$\hat{d} = \hat{p}_A - \hat{p}_B = 0.70 - 0.72 = -0.02 \quad (18.9)$$

$$\hat{\sigma}_{\hat{d}} = \frac{\sqrt{\hat{p}_A(1 - \hat{p}_A) + \hat{p}_B(1 - \hat{p}_B)}}{\sqrt{n}} = \frac{\sqrt{0.70(1 - 0.70) + 0.72(1 - 0.72)}}{\sqrt{1000}} = 0.020 \quad (18.10)$$

For a 95% confidence interval (which makes sense here), we need to use the z critical value $z_{0.025} = 1.96$. The confidence interval for $\hat{d}$ is therefore

$$\begin{aligned} & (\hat{d} - z_{0.025}\hat{\sigma}_{\hat{d}}, \hat{d} + z_{0.025}\hat{\sigma}_{\hat{d}}) \\ & = (-0.02 - 1.96 \times 0.020, -0.02 + 1.96 \times 0.020) \\ & = (-0.60, 0.20) \end{aligned} \quad (18.11)$$

This is almost exactly the same as the simulation estimates.

18.4 Issues in A/B testing

Statistical versus practical significance in A/B testing There is an important distinction between statistical significance and practical significance. We might test the run time of two versions (A and B) of a webserver program on random hits from users. In the example that we've just seen if we make n large enough, we can show that there B is better than A with a p -value of 0.001. However, is the difference of 2% actually that meaningful? In this case it is still probably worth it, since it requires little or no extra effort or energy to create a green button rather than a blue button. But if the processing required to serve version B used a lot more energy, maybe that 2% improvement wouldn't be worth it.

Quoting from the ASA statement on p -values again:

A p -value, or statistical significance, does not measure the size of an effect or the importance of a result. Statistical significance is not equivalent to scientific, human, or economic significance. Smaller p -values do not necessarily imply the presence of larger or more important effects, and larger p -values do not imply a lack of importance or even lack of effect. Any effect, no matter how tiny, can produce a small p -value if the sample size or measurement precision is high enough, and large effects may produce unimpressive p -values if the sample size is small or measurements are imprecise. Similarly, identical estimated effects will have different p -values if the precision of the estimates differs. (Wasserstein and Lazar, 2016)

Ethical questions in A/B testing There are a number of ethical issues we should consider in A/B testing:

- We are undertaking experiments on people
- In a commercial situation, users do not give informed consent.
- In an academic situation, informed consent is required – but how can we get this informed consent without affecting the experiment?
- What about data protection?

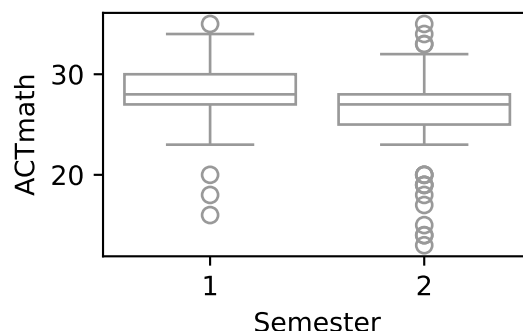


Figure 18.4: Maths scores in Semester 1 (Autumn) and Semester 2 (Spring). Data from [Edge and Friedberg \(1984\)](#), via [Devore and Berk \(2012\)](#).

The experiment in which Facebook manipulated news feeds to explore the effect on users' moods ([Kramer et al., 2014](#)) was an example of A/B testing that was widely seen as problematic ([Verma, 2014](#)) because of a lack of informed consent, no opportunity to opt-out and no institutional review of the experiment, because it was carried out by a private company. Many A/B tests are arguably not having such a significant effect on users – but they may have some effect nonetheless.

It's therefore important to reflect on an A/B test before setting it up. Questions might include:

- Would I feel comfortable if this change was tested on me?
- What potential harms could be caused to users?

18.5 Comparing groups with numeric responses

The problem of two numeric samples In the A/B testing described in the previous sections, the samples A and B comprise n binary response variables, and we estimate the difference in proportions between the groups. A related problem is when the response variables in populations A and B are numeric rather than binary, and we wish to assess whether the distributions are similar or different.

For example, the populations A and B could be students taking a calculus course in Semester 1 (Group A) and students taking the course in Semester 2 (Group B). Figure 18.4 plots these distributions using boxplots. The maths scores of students taking a calculus course in Semester 2 seem to be lower than the grades in Semester 1.

We'll call the grades of the m Semester 1 students $x_1, \dots, x_m$ and the grades of the n Semester 2 students $y_1, \dots, y_n$. We can compute the two means of the grades of each group, and find that difference is $\bar{x} - \bar{y} = 2.37$. But, assuming that the maths scores are representative of performance in Semester 1 and Semester 2 in other years, we'd like to find a 95% confidence interval for the difference between the means – in other words an interval that we would expect to contain the true, underlying difference 95% of the time.

Applying the bootstrap We already know how to use a sample to compute a confidence interval for the mean of one population using the bootstrap estimator ([Bootstrap estimation of confidence intervals](#)). Can we adapt the bootstrap to give us a confidence interval around the difference between two means? To apply the bootstrap, on each bootstrap step we sample with replacement from both groups:

- For j in $1, \dots, B$
 - Take sample x^* of size m from the sample $x_1, \dots, x_m$ with replacement
 - Take sample y^* of size n from the sample $y_1, \dots, y_n$ with replacement
 - Compute the sample mean of the new samples, $\bar{x}_j^*$ and $\bar{y}_j^*$
 - Compute and store the difference in the sample means $\bar{x}_j^* - \bar{y}_j^*$
- Plot the bootstrap distribution of $\bar{x}_j^* - \bar{y}_j^*$

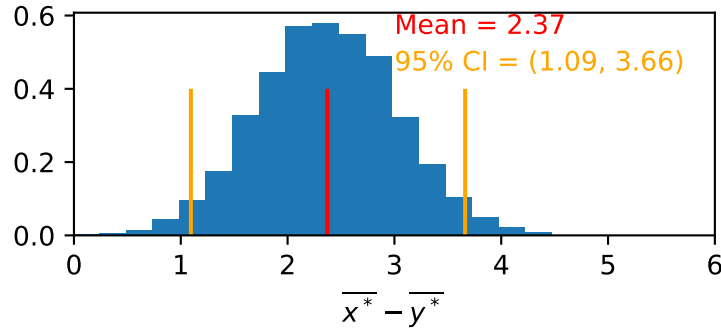


Figure 18.5: Bootstrap distribution of the maths grades example.

- We can also compute the bootstrap estimator of the variance of the difference between the means:

$$s_{\text{boot}}^2 = \frac{\sum_{j=1}^B \left((\bar{x}_j^* - \bar{y}_j^*) - (\bar{x} - \bar{y}) \right)^2}{B - 1}$$

The bootstrap distribution is shown in Figure 18.5. We can see that the 95% confidence interval for the difference in mean grades is 1.11 to 3.66.

Exercise

Write a bootstrap simulation to test the hypothesis that the means of the grades differ between the two semesters.

18.6 The theoretical method of testing for differences between groups

Assumptions We can also compute confidence intervals and undertake hypothesis testing for the difference between the means of two samples theoretically. The assumptions are:

- $X_1, \dots, X_m$ is a random sample from a population with mean μ_1 and variance σ_1^2
- $Y_1, \dots, Y_n$ is a random sample from a population with mean μ_2 and variance σ_2^2
- The samples are independent of each other.

Estimator of the difference The difference between the sample means $\bar{X} - \bar{Y}$ is an unbiased estimator of the difference between the true means $\mu_1 - \mu_2$. This follows from $\bar{X}$ being an unbiased estimator of μ_1 and $\bar{Y}$ being an unbiased estimator of μ_2 . The standard deviation of the estimator is:

$$\sigma_{\bar{X} - \bar{Y}} = \sqrt{\frac{\sigma_1^2}{m} + \frac{\sigma_2^2}{n}} \quad (18.12)$$

This follows from the two samples being independent, so $V(\bar{X}) - V(\bar{Y}) = V(\bar{X}) + V(\bar{Y}) = \sigma_1^2/m + \sigma_2^2/n$.

Theoretical distribution for large samples When both samples are larger than 40 ($m > 40$ and $n > 40$) we can regard the sample as large. As when we estimated confidence intervals for the mean of one population ([Confidence intervals for the mean from small samples](#)), we define a standardised variable, which we expect to be zero in the case of a null hypothesis that the true difference between the population means is $\mu_1 - \mu_2$:

$$Z = \frac{\bar{X} - \bar{Y} - (\mu_1 - \mu_2)}{\sqrt{S_1^2/m + S_2^2/n}} \quad (18.13)$$

The denominator is the sample standard deviation of the estimator, and is a random variable. As when estimating the mean of one population from a small sample ([Confidence intervals for the mean from small](#)

samples), for small m and n this will vary considerably between different samples, and so we do have to consider these random effects. However, for large m and n , it will approximate the true population means, and its variability is low enough to consider it as a fixed parameter as when estimating the mean of one population (Method of estimating confidence interval for the mean of a large sample). In the limit of large n and m , the Central Limit Theorem suggests that the distribution of the statistic should be normal, so we can use a normal distribution with a mean of zero and standard deviation of $\sqrt{s_1^2/m + s_2^2/n}$ as the sampling distribution of the test statistic.

Procedure to calculate a confidence interval for the difference in means from large samples

1. Calculate estimator of the difference in means: $\hat{\mu}_1 - \hat{\mu}_2 = \bar{x} - \bar{y}$
2. Calculate estimated standard error of the difference in means:

$$\hat{\sigma}_{\hat{\mu}_1 - \hat{\mu}_2} = \sqrt{\frac{s_1^2}{m} + \frac{s_2^2}{n}} \quad (18.14)$$

3. The $100(1 - \alpha)\%$ confidence interval is

$$(\bar{x} - \bar{y} - z_{\alpha/2} \hat{\sigma}_{\hat{\mu}_1 - \hat{\mu}_2}, \bar{x} - \bar{y} + z_{\alpha/2} \hat{\sigma}_{\hat{\mu}_1 - \hat{\mu}_2}) \quad \text{which can also be written} \quad \bar{x} - \bar{y} \pm z_{\alpha/2} \hat{\sigma}_{\hat{\mu}_1 - \hat{\mu}_2} \quad (18.15)$$

where $z_{\alpha/2}$ is the z critical value for $\alpha/2$.

In the comparison of grades example (Figure 18.4), we have $m = 74$, $n = 80$, $\bar{x} = 28.2500$, $\bar{y} = 25.8784$, $s_1 = 3.2472$, and $s_2 = 4.5929$. As an exercise compute the 95% confidence interval of the difference in the means $\mu_1 - \mu_2$ and compare this with the estimate from the bootstrap.

Theoretical distribution for small samples When either group has a sample size of less than 40, the variability of the sample standard deviation has to be taken into account, and it turns out that the sampling distribution of the standardised statistic is a t -distribution with a number of degrees of freedom ν that depends on the standard deviations of both distributions:

$$\nu = \frac{(s_1^2/m + s_2^2/n)^2}{\frac{(s_1^2/m)^2}{m-1} + \frac{(s_2^2/n)^2}{n-1}} \quad (18.16)$$

Procedure to calculate a confidence interval for the difference in means from small samples To determine a confidence interval of $100(1 - \alpha)\%$, we follow the first two steps of the large sample procedure above. We then compute ν using Equation 18.16, and find the t critical value $t_{\alpha/2, \nu}$. We can then use the t critical value and the standard error of the estimator to compute the confidence interval in a similar way to the last step of the large sample procedure:

$$\bar{x} - \bar{y} \pm t_{\alpha/2, \nu} \hat{\sigma}_{\hat{\mu}_1 - \hat{\mu}_2} \quad (18.17)$$

18.7 Quantifying the effect size of differences between two numeric samples

Cohen's d The issue of practical significance raised in Issues in A/B testing can be addressed quantitatively with numerical samples using a statistic called **Cohen's d** . Using the same notation for data in two groups $(x_1, \dots, x_m$ and $y_1, \dots, y_n)$, Cohen's d is defined:

$$d = \frac{\bar{x} - \bar{y}}{s} \quad ; \quad s = \sqrt{\frac{(m-1)s_1^2 + (n-1)s_2^2}{m+n-2}} \quad (18.18)$$

Figure 18.6 shows some examples of Cohen's d for different pairs of samples. We can see that in the leftmost and rightmost plots, the value of d is around 0.5, meaning that the difference in the means is about half the size of the pooled standard deviations of the groups – in other words there is a substantial overlap of the two distributions. The middle plot has a much higher value of d , and it is clear here that the distributions overlap little. Although the leftmost and rightmost plots have similar values of d , the rightmost plot, with many more datapoints in each group, has a higher t -statistic and a lower p -value than the one on the leftmost plot.

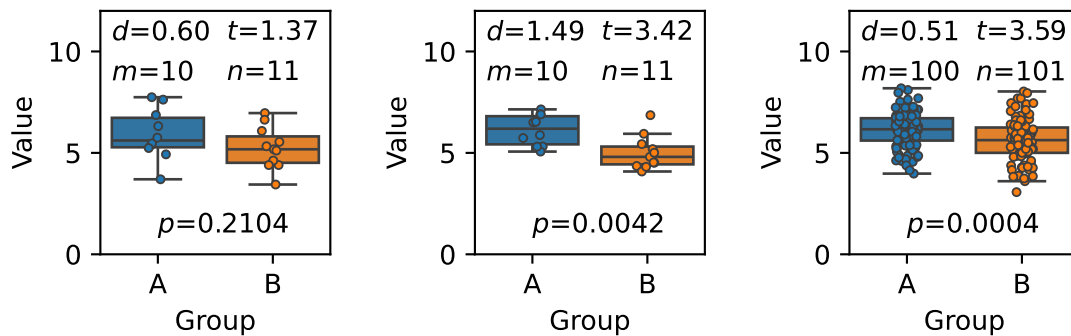


Figure 18.6: Cohen's d for a number of samples, along with the corresponding t statistics and the p -values for testing if the groups are different. Group A is identified with the x s and Group B with the y s.

18.8 Paired data - to appear

18.9 Relationship between hypothesis testing and confidence intervals

Suppose we had wanted to test the hypothesis that average performance in Semester 2 is different to average performance in Semester 1. Our null hypothesis would be

H_0 : The mean performance in semester 1 is the same as the mean performance in semester 2.

The alternative hypothesis would be:

H_a : The mean performance in semester 1 is different from the mean performance in semester 2.

We've previously simulated the null hypothesis (here, no difference in means) to generate a distribution of what the test statistic (here, the difference in the sample means) would be under the null hypothesis, and then compared this distribution with the observed value of our test statistic.

It turns out that there is a duality between confidence intervals and hypothesis testing. Instead of the distribution of the sampling distribution generated under the null hypothesis, we have used the observed data to generate the distribution of the estimator for the parameter corresponding to the test statistic. Instead of the observed value of the test statistic, we have the value the parameter would take under the null hypothesis.

In this example, we used the bootstrap to estimate the distribution of the estimator of the difference of the means $\mu_x - \mu_y$ (Figure 18.5), and we could then ask if the null hypothesis value of the difference in the means (0) lies in either of the rejection regions in the tails of that distribution. If so, we can reject at the level corresponding to the size of the tails. Alternatively, we could compute a p -value, by finding at what quantile the null hypothesis value (here 0) lies on the distribution. In this case we would find $p = 0$, so the null hypothesis would be rejected.

For more on the duality between confidence intervals and hypothesis testing see this [Quora article](#).

Related Workshop: Statistical problems 2

<https://opencourse.inf.ed.ac.uk/inf2-fds/course-materials/semester-2/week-5/task>

The aim of this task and workshop is to apply some of the techniques from inferential statistics, in particular hypothesis testing, logistic regression and A/B testing.

Bayesian inference applied to A/B testing (non-examinable)

So far we've approached statistics from a frequentist or sampling theory perspective. This is predicated on there being a real population parameter that we're trying to estimate, and then using the sampling distribution to model what happens in the process of taking a sample from the population. When we test hypotheses, we find the fraction of sampling simulations of the null hypothesis that could produce a result at least as extreme as what we observe, the test statistic. The two hypothesis

that are competing are the specified null hypothesis and the alternative hypothesis, which is the “not the null hypothesis”.

In Bayesian stats, we do not assume that a single population parameter exists. In the situation we have just described, this seems appropriate: there is a potentially infinite population of users, and so p_A and p_B do not exist in the same way that the number of wildcats in Scotland does.

We can also use Bayesian stats to compare competing hypotheses. In this case our hypotheses might be that A has a higher click through rate than B, or vice versa.

A fundamental quantity in Bayesian statistics is the **posterior probability**, that is the probability distribution of a parameter (e.g. p_A or p_B) *after* we have observed data (e.g. $\hat{p}_A$ and $\hat{p}_B$ from n observations). The posterior probability is derived from the **likelihood** of generating an observation $\hat{p}_A$ given a value of the parameter p_A and the **prior probability** distribution.

The prior distribution allows us to express our beliefs about the likely distribution parameter should be *before* we have seen the data. This sounds like we are biasing the outcome – which we are – but this can make sense, e.g. we probably believe that a coin has a probability of 1/2 of landing on Heads – we would like a lot of data to convince us otherwise. However, if we really don’t have much idea, we can set the prior distribution to be uniform.

Mathematically, Bayes theorem relates the posterior probability (called “the posterior” for short), the likelihood and the prior probability (called “the prior” for short). For our particular example for the A group, we’ll consider the number of observed click-throughs $n_A = n\hat{p}_A$. In this case Bayes’ theorem reads:

$$p(p_A | n_A, n) = \frac{p(n_A | p_A, n) p(p_A)}{\int_{p_A=0}^1 p(n_A | p_A, n) p(p_A) dp_A} \quad (18.19)$$

The denominator is the “probability of the data” $p(n_A)$ and derives from the definition of conditional probability.

Suppose our model of how the data arises is that users decide, independently of each other, to click through with a probability p_A . To start of with, we don’t know what p_A is – we assume it is equally likely to be any number between 0 and 1. In other words, we are assuming a uniform prior: $p(p_A) = 1$ for $0 \leq p_A \leq 1$.

Probability theory tells us that if we have n repeats of a trial in which the probability of “success” on each trial is p_A , then the distribution of the total number of successes is given by a binomial distribution.

$$p(n_A | p_A, n) = \binom{n}{n_A} p_A^{n_A} (1 - p_A)^{n - n_A} \quad (18.20)$$

This is the likelihood of observing n_A out of n users clicking through, given a hypothetical click-through probability of p_A .

Using integration by parts and recursion, we can evaluate the integral in the denominator to get $1/(n + 1)$. We thus find that the posterior distribution is

$$p(p_A | n_A, n) = (n + 1) \binom{n}{n_A} p_A^{n_A} (1 - p_A)^{n - n_A} \quad (18.21)$$

Note that, although this looks like a binomial distribution, the variable is actually p_A , and this is therefore a beta distribution with $a = n_A + 1$ and $b = n - n_A + 1$ ^a

This distribution is now the distribution of the parameter p_A given the data. Using differentiation, we can prove that the distribution has a maximum at $\hat{p}_A = n_A/n$. This makes sense, since we would expect the most likely value to be the observed proportion. The 2.5% and 97.5% centiles will enclose 95% of the distribution and are our Bayesian credibility interval – the analogue of a frequentist confidence interval.

We can write a similar expression for p_B . Since the A and B group are independent, the likelihood is

$$p(n_A, n_B | p_A, p_B, n) = \binom{n}{n_A} p_A^{n_A} (1 - p_A)^{n - n_A} \binom{n}{n_B} p_B^{n_B} (1 - p_B)^{n - n_B} \quad (18.22)$$

It turns out that the posterior $p(p_A, p_B | n_A, n_B, n)$ is also the product of the two posterior distributions.

$$p(p_A, p_B | n_A, n_B, n) = (n + 1) \binom{n}{n_A} p_A^{n_A} (1 - p_A)^{n - n_A} \binom{n}{n_B} p_B^{n_B} (1 - p_B)^{n - n_B} \quad (18.23)$$

To compute the posterior distribution of the difference, we substitute $p_B = p_A - d$ and then integrate over p_A from d to 1 (if d is positive) or from 0 to $1 + d$ (if d is negative).

^aWe can verify the constant is correct:

$$\frac{\Gamma(a+b)}{\Gamma(a)\Gamma(b)} = \frac{\Gamma(n+2)}{\Gamma(n_A+1)\Gamma(n-n_A+1)} = \frac{(n+1)!}{n_A!(n-n_A)!}$$

Chapter 19

Statistical inference of regression coefficients and predictions

Recommended reading

Modern Mathematical Statistics with Applications, Chapter 12

19.1 Inference about linear regression coefficients with the bootstrap

Inference on the slope of the regression line When we first encountered linear regression (in the chapter on [Linear Regression](#)), we hadn't learned about inferential statistics. We can now apply what we've learned about statistical inference to linear regression.

In the linear regression model, the response variable y depends linearly on the predictor variable x (Equation 7.1 in the chapter on [Linear Regression](#)):

$$y = \beta_0 + \beta_1 x \quad (19.1)$$

The variables β_0 and β_1 are called parameters or regression coefficients and are the intercept and slope of the line respectively.

In the chapter on [Linear Regression](#) we used the principle of least squares to obtain point estimates of the intercept and slope $\hat{\beta}_0$ and $\hat{\beta}_1$ of the regression line (Equations 7.6 and 7.7). For any set of data, we would now like to answer the following inferential statistics questions about the regression coefficients:

- What is our confidence in them?
- Do they represent a real effect, and not just a chance correlation in the data?
- Can we quantify uncertainty in our predictions?

Bootstrap estimation confidence intervals Since the value of the regression coefficients can be calculated from sample data, they fit the definition of a statistic (see chapter on [Randomness, sampling and simulation](#)). We can therefore use the bootstrap estimator to estimate the uncertainty in the coefficients. For example, the dataset of female squirrel weight and length (seen in the chapter on [Data collection and statistical relationships](#)) contains $n = 32$ pairs of weight-length pairs. We apply the bootstrap by sampling with replacement $n = 32$ pairs from the original 32 pairs; some pairs will be sampled once, and some not at all. We then apply linear regression to each resampled set of pairs, record the intercept and slope parameter for that sample.

In Figure 19.1, we show examples of regression lines of weight of female squirrels from resampled datasets (top left). The top right plot shows the distribution of the slope generated from 1000 resamples, and the bottom left plot shows the distribution of the intercept. There is considerable variation in the slope – the 95% CI is (2.50, 4.24) g/mm. This is evident in the sample of fits shown in the top left plot. There is also variation in the intercept – however, the scatter plot shows that it is very closely negatively correlated with the slope, since a steeper slope means that the intercept needs to be more negative to ensure that the regression lines pass through the centre of the data.

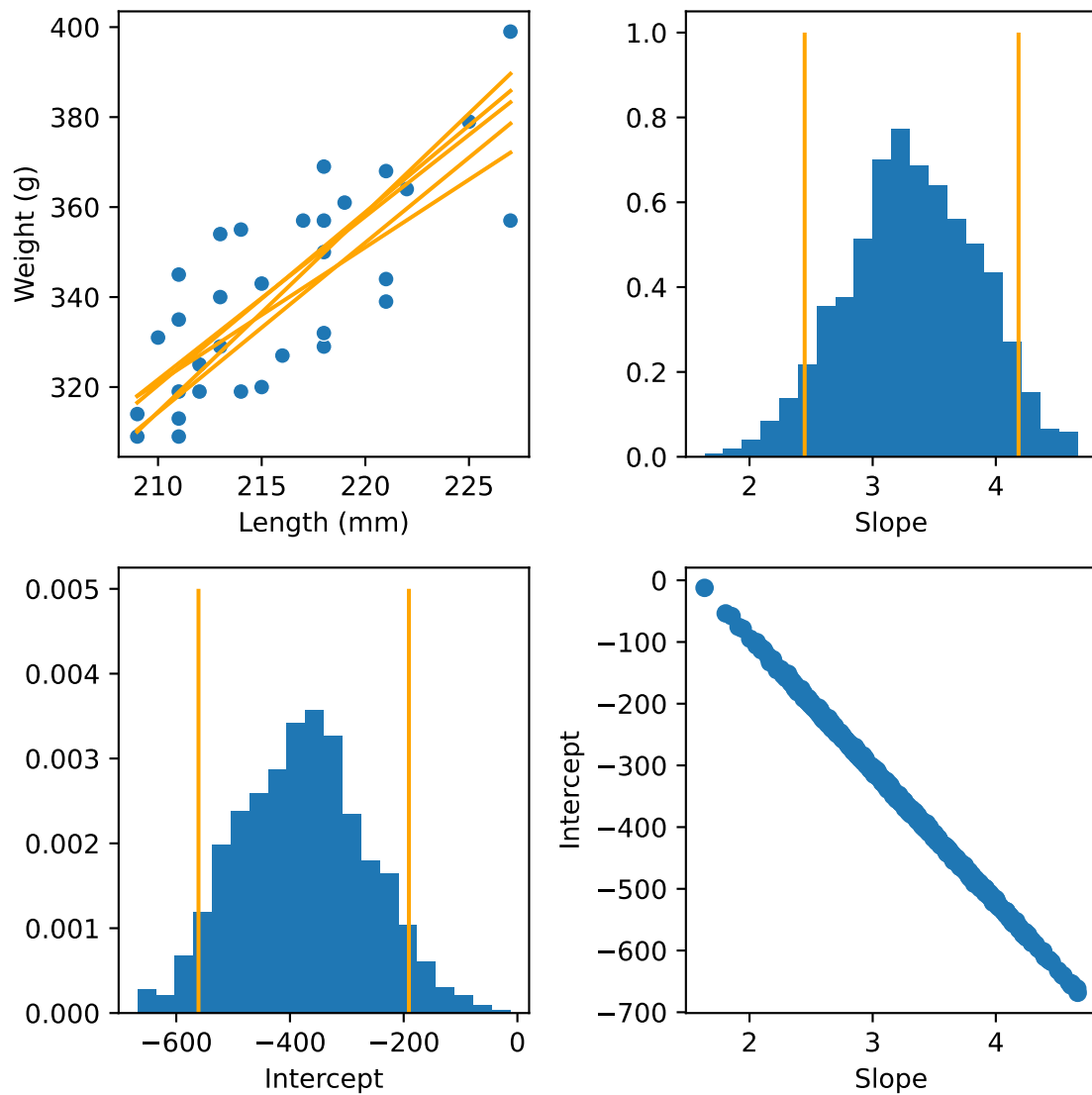


Figure 19.1: Bootstrap applied to squirrel length and weight data. Top left: the Squirrel data of [Wauters and Dhondt \(1989\)](#). A sample of 5 linear regression lines from the bootstrap distribution are shown. Top right: bootstrap distribution of Slope $\hat{\beta}_1$. Bottom left: bootstrap distribution of Intercept $\hat{\beta}_0$. Bottom right: scatter plot of the slope and intercept coefficients from the bootstrap simulations.

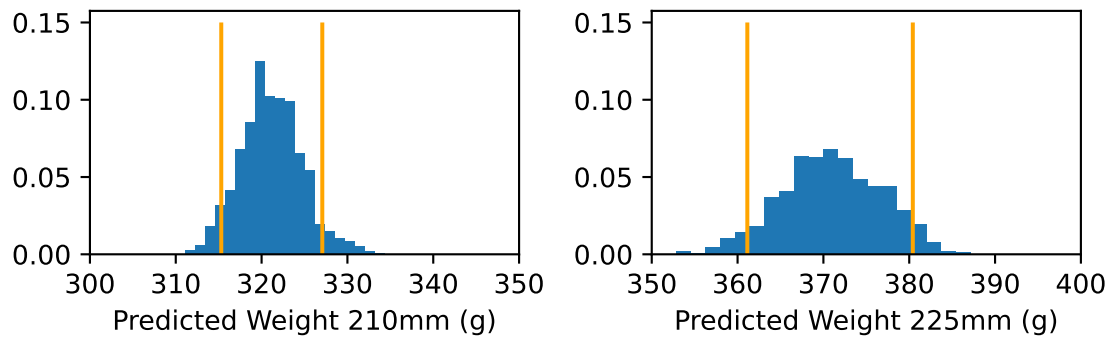


Figure 19.2: Bootstrap applied to squirrel length and weight data. Distribution of weight predictions for squirrel 210mm long (left) and 225mm long (right).

Bootstrap hypothesis testing Often we want to know if the slope of a regression model is significantly different from zero, since a non-zero slope suggests a relationship between the predictor and response variables. In other words, can we reject the null hypothesis that $\hat{\beta}_1 = 0$? The bootstrap distribution of $\hat{\beta}_1$ (Figure 19.1, top right) answers this question immediately, since the range of the bootstrap distribution does not even include 0.

Prediction uncertainty Suppose we want to predict the weight of a squirrel that is a given length, *and* state how uncertain we are about the prediction. We can use the bootstrap distribution of coefficients for this purpose. All we do is substitute the length we’re interested in as the predictor x in the equation $y = \hat{\beta}_0^* + \hat{\beta}_1^*x$ for each pair of the bootstrap estimates $(\hat{\beta}_0^*, \hat{\beta}_1^*)$. Figure 19.2 shows the bootstrap distributions of the predicted weight of squirrels that are 210mm long versus those that are 225mm long. To generate these distributions we have substituted $x = 210$ or $x = 225$ in the equation $y = \hat{\beta}_0^* + \hat{\beta}_1^*x$ for each pair of the bootstrap estimates $(\hat{\beta}_0^*, \hat{\beta}_1^*)$. We can see that the confidence interval for squirrels of length 225mm is wider, reflecting that there are fewer data for longer squirrels than shorter ones (Figure 19.1, top left).

19.2 Understanding stats package linear regression output

Although understanding the principle of regression using the bootstrap is important conceptually, in practice, we use standard linear regression routines such as Python’s `statsmodels` package or the `lm` function in R. As well as outputting the point estimates of the coefficients $\hat{\beta}_0$ and $\hat{\beta}_1$, these packages also supply information about the uncertainty about the point estimates, although using the theoretical methods described in the next sections rather than the bootstrap.

Nevertheless, we’re now in a position to understand a lot more of the output produced by a stats package when it fits a linear regression model, as shown in Figure 19.3.

- Focusing the lower part of the table, the “Intercept” row relates to β_0 and the “Length” row relates to β_1 . In this example, the row is called “Length” because we are regressing on the Length predictor.
- The **coef** column gives the point estimates of the coefficients, $\hat{\beta}_0$ and $\hat{\beta}_1$.
- The **std err** column gives the standard errors in the coefficients $\hat{\sigma}_{\hat{\beta}_0}$ and $\hat{\sigma}_{\hat{\beta}_1}$.
- The **t** column shows the t-statistic for that value, i.e. the value

$$t = \frac{\hat{\beta}}{\hat{\sigma}_{\hat{\beta}}} \quad (19.2)$$

where we can replace β with β_0 and β_1 . The larger the value, the further the coefficient is from 0, measured in multiples of the standard error in the mean, and the less likely it is that the data might have arisen from the null hypothesis that $\beta = 0$.

- The $P > |t|$ column is the p -value, which quantifies how much of the probability mass of the t -distribution is bigger than the magnitude of the t-statistic. The lower the p -value, the less compatible the data is with the null hypothesis that $\beta = 0$.
- The [0.025 and 0.975] columns give the 95% confidence intervals in β .

OLS Regression Results

Dep. Variable:	Weight	R-squared:	0.597			
Model:	OLS	Adj. R-squared:	0.583			
Method:	Least Squares	F-statistic:	44.37			
Date:	Sun, 10 Jan 2021	Prob (F-statistic):	2.24e-07			
Time:	21:08:04	Log-Likelihood:	-129.18			
No. Observations:	32	AIC:	262.4			
Df Residuals:	30	BIC:	265.3			
Df Model:	1					
Covariance Type:	nonrobust					
	coef	std err	t	P> t 	[0.025	0.975]
Intercept	-382.7372	108.680	-3.522	0.001	-604.692	-160.783
Length	3.3515	0.503	6.661	0.000	2.324	4.379
Omnibus:	8.046	Durbin-Watson:	2.337			
Prob(Omnibus):	0.018	Jarque-Bera (JB):	2.231			
Skew:	0.092	Prob(JB):	0.328			
Kurtosis:	1.720	Cond. No.	9.38e+03			

Figure 19.3: Output from Python statsmodels applied to the squirrel dataset.

19.3 Sampling theory inference about linear regression coefficients

Standard error of the gradient coefficient Although the bootstrap estimates work OK, it is possible to derive algebraic expressions for the standard error of the intercept and gradient. We will first present the results. Stats packages will provide estimates of these quantities as a matter of course.

The estimated standard error in the estimator for $\hat{\beta}_1$ is:

$$\hat{\sigma}_{\hat{\beta}_1} = s_{\hat{\beta}_1} = \frac{s}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2}} \quad (19.3)$$

where s is the estimate for σ , the standard deviation of the residuals:

$$\hat{\sigma}^2 = s^2 = \frac{1}{n-2} \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i)^2 = \frac{1}{n-2} \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \frac{\text{SSE}}{n-2} \quad (19.4)$$

The more tightly the points cluster around the regression line (small s in the numerator) the smaller the standard error of the estimate of the gradient we have. Also, the larger the spread of points on the x -axis (large denominator in Equation 19.3), the smaller the standard error of the estimator.

Confidence interval for $\hat{\beta}_1$ As in our estimate of the confidence interval around the mean of small samples ([Confidence intervals for the mean from small samples](#)), we regard both $\hat{\beta}_1$ and the estimated standard error $S_{\hat{\beta}_1}$ as random variables. In order to determine the confidence intervals, we define the quantity:

$$T = \frac{\hat{\beta}_1 - \beta_1}{S_{\hat{\beta}_1}} \quad (19.5)$$

where β_1 is the true but unknown value. It turns out that this quantity has a t -distribution with $n - 2$ degrees of freedom. To determine a $100(1 - \alpha)\%$ confidence interval, we therefore set the probability that the variable T lies between the t critical values $\pm t_{\alpha/2, n-2}$:

$$P \left(-t_{\alpha/2, n-2} < \frac{\hat{\beta}_1 - \beta_1}{S_{\hat{\beta}_1}} < t_{\alpha/2, n-2} \right) < 1 - \alpha \quad (19.6)$$

We rearrange the inside of the probability statement to find the $100(1 - \alpha)\%$ confidence interval for the slope is:

$$\hat{\beta}_1 - s_{\hat{\beta}_1} t_{\alpha/2, n-2} < \beta_1 < \hat{\beta}_1 + s_{\hat{\beta}_1} t_{\alpha/2, n-2} \quad (19.7)$$

Testing hypotheses about $\hat{\beta}_1$ Often we are interested to know if the slope of a regression model is significantly different from zero – in other words, can we reject the null hypothesis that $\hat{\beta}_1 = 0$? To do this we can assume the null hypothesis by setting $\beta_1 = 0$ in Equation 19.5. Under the null hypothesis, the estimator for the slope is expected to be zero. If the estimated value of $\hat{\beta}_1$ is large relative to $S_{\hat{\beta}_1}$ then the T value will be large, and therefore far out in one of the tails of the distribution. By computing the probability mass more extreme than the T value, we get the p -value, and we can then decide if to reject or not reject the null hypothesis that the slope coefficient is different from zero.

If the p -value is large, this suggests that there may be no real relationship between the predictor and response variables, but one appears to be there by chance. Generally when the coefficient of determination R^2 is large and the number n is reasonably large, the p -value will be low, and it is reasonable to start drawing conclusions about the relationship between x and y .

19.4 Derivation of standard error of estimator for slope coefficient

Where does the expression for the standard error of $\hat{\beta}_1$ (Equation 19.3) come from? First we make a subtle change to the linear regression model (Equation 19.1), by thinking of the response variable as being a random variable Y that is a sum of a deterministic linear function of the predictor variable x and a random deviation ε , which we sometimes call an **error term**:

$$Y = \beta_0 + \beta_1 x + \varepsilon \quad (19.8)$$

We will assume that ε is a random variable with an expected value of $E[\varepsilon] = 0$. In principle, we can assume that is drawn from any distribution we would like, but a very common assumption is that it is drawn from a normal distribution with mean 0 and variance σ^2 :

$$\varepsilon \sim N(0, \sigma^2) \quad (19.9)$$

The expression for $\hat{\beta}_1$ is now a random variable:

$$\hat{\beta}_1 = \frac{\sum (x_i - \bar{x})(Y_i - \bar{Y})}{\sum (x_i - \bar{x})^2} \quad (19.10)$$

The term $Y_i - \bar{Y}$ is exactly equal to $\varepsilon_i - \bar{\varepsilon}$, which, from the properties of ε_i has $E[\varepsilon_i - \bar{\varepsilon}] = 0$ and variance $V[\varepsilon_i - \bar{\varepsilon}] = \sigma^2 n / (n - 1)$. Therefore, the expected variance of $\hat{\beta}_1$ is:

$$V[\hat{\beta}_1] = \frac{\sum (x_i - \bar{x})^2 V[\varepsilon_i - \bar{\varepsilon}]}{\sum (x_i - \bar{x})^4} = \frac{\sigma^2 n / (n - 1)}{\sum (x_i - \bar{x})^2} \quad (19.11)$$

Thus the standard error of the estimator is:

$$\sigma_{\hat{\beta}_1} = \sqrt{V[\hat{\beta}_1]} = \frac{\sigma \sqrt{n / (n - 1)}}{\sqrt{\sum (x_i - \bar{x})^2}} \quad (19.12)$$

Since the sample variance s^2 of the residuals is an estimator of $\sigma^2 n / (n - 1)$, we get the estimated standard error:

$$\hat{\sigma}_{\hat{\beta}_1} = \frac{s}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2}} \quad (19.13)$$

For more details see pp. 642–643 of [Devore and Berk \(2012\)](#).

Part V

The Maximum Likelihood Principle and Regression

Chapter 20

Logistic regression

20.1 Principle of logistic regression

What is logistic regression used for? Logistic regression has various uses, including:

- **As a parametric supervised classification algorithm.** For example, a bank has data on previous customers it has considered offering credit cards to, including predictor variables (independent variables) such as their age, income, housing status and employment status. Each of these sets of variables is labelled with the response variable of whether the credit card was approved. The task is to determine if the credit card should be approved for a new customer.
- **As a way of investigating the association between predictor variables and a binary (also called dichotomous) response variable.** For example, suppose we have an observational study of patients of different ages, health levels, ethnicity and gender. Some of the patients have had a dose of vaccine for an illness, and some haven't. We'd like to know how the probability of getting the illness depends on if the vaccine has been administered or not. Like multiple linear regression, we can examine logistic regression coefficients to isolate the effect of the vaccine, controlling for the other variables.

Similarities and differences to k -NN We've already discussed classifiers, when we looked at k -Nearest Neighbours ([Supervised learning: Classification with Nearest neighbours](#)). As a reminder, the problem of classification is to predict the correct label for an unlabelled input item described by a feature vector of variables. As well as acting as a classifier, logistic regression can predict a real-valued number, the *probability* of a data point belonging to a category, on the basis of the predictors/predictor variables. In fact, we convert the logistic regression model into a classifier by choosing at a threshold level of probability at which we make a decision. For example, we might only want to approve credit cards that we think would have a 60% chance of being approved historically.

Association between continuous predictor and binary outcome We will use the example of the credit card approval to illustrate how logistic regression is used as a classifier and as a way of exploring associations between variables. Figure [20.1](#) visualises the relationship between age and approval. Because age is a continuous variable, we can plot individual datapoints on a scatter plot. It looks like older customers were more likely to have their credit approved than younger ones.

Association between binary predictor and binary outcome: Odds and odds ratios Employment is a binary variable ("employed" or "not employed"). If we tried plotting it in the same way as age versus approval, we'd end up with a very uninformative plot, so instead we look at a **contingency table** (Table [20.1](#)), which shows the relative frequency (empirical probability) of having credit approved or not approved based on employment status.

In logistic regression, we will see that it makes sense to describe these probabilities in terms of **odds**¹, which we define as:

$$\text{Odds}(\text{Success}) = \frac{P(\text{Success})}{P(\text{Failure})} = \frac{P(\text{Success})}{1 - P(\text{Success})} \quad (20.1)$$

If success and failure are equally likely, the odds are equal to 1.

¹*Modern Mathematical Statistics with Applications* calls the "odds" the "odds ratio", which is not standard usage.

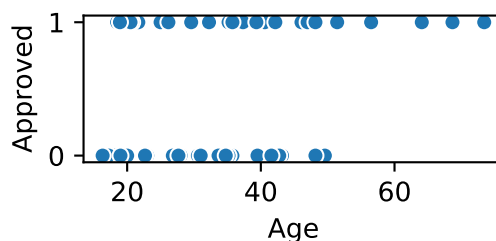


Figure 20.1: Age versus approval. Each datapoint represents the age of a customer and whether their credit was approved (1) or not approved (0). It looks like a greater fraction of older customers had their credit approved than younger ones. A random subsample of data points is plotted to aid visualisation.

Table 20.1: Contingency table showing relative frequencies of approval (“Success”) or not approval (“Failure”) based on employment status (first two columns) and the odds of approval (final column). The odds are the probability of approval divided by the probability of not being approved.

	Approved	Not approved	Approval odds
Employed			
0	0.25	0.75	0.34
1	0.71	0.29	2.42

We call the **odds ratio** (OR) the ratio between the odds of credit approval if employed versus credit approval if not employed.

$$\text{OR}(x) = \frac{\text{Odds}(\text{Success}|x = \text{True})}{\text{Odds}(\text{Success}|x = \text{False})} \quad (20.2)$$

We can find the odds ratio of employment in the credit example by setting “Success” to “Approved” and x to “Employed”, giving an answer of $7.09 = 2.42/0.34$. Thus, the odds of someone who is employed having credit approved are 7.09 times larger than the odds of someone who is not employed having credit approved. The odds ratio is sometimes referred to as an **effect size** and expressed as the percentage change in the odds from x being False to True; this case the effect size is 609%, since the effect of employment increases the odds of approval by this amount.

Principle of logistic regression with one predictor variable As its name suggests, logistic regression is related to linear regression. Suppose that the response variable (or dependent variable) y is a dichotomous variable (i.e. a categorical variable with two categories). We’ll represent the categories by 0 (failure) and 1 (success). We’d like to model the probability $P(Y = 1|X = x)$ that the response variable is 1, given the predictor variable (or predictor) X has a value x . Because we’re predicting a probability, the answer given by logistic regression has to lie between 0 and 1. Therefore, $P(Y = 1|X = x)$ can’t be a linear function of x .

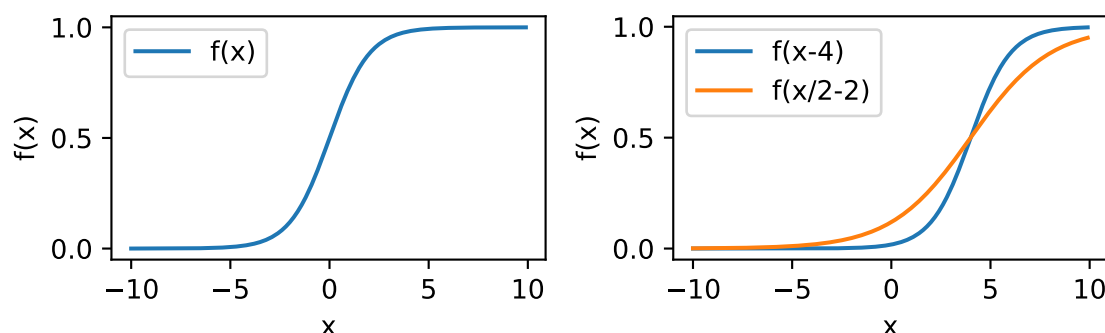


Figure 20.2: The logistic function. Left: the standard logistic function: $f(u) = \exp(u)/(1 + \exp(u))$, which can also be written $f(u) = 1/(1 + \exp(-u))$. Right: Examples of shifted logistic curves.

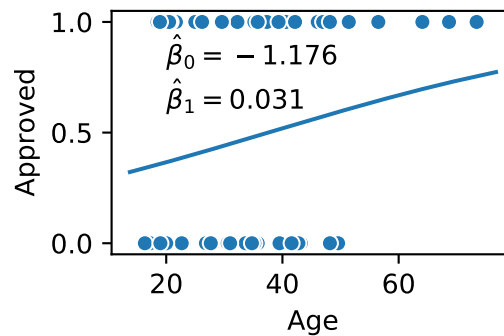


Figure 20.3: Logistic regression of age on credit approval.

We get around this problem by allowing $P(Y = 1|X = x)$ to be a *nonlinear* function of x . A function that works well in many applications is the **logistic function** (Figure 20.2). Using f to denote the logistic function, the probability of a success is:

$$P(Y = 1|X = x) = f(\beta_0 + \beta_1 x) = \frac{e^{\beta_0 + \beta_1 x}}{1 + e^{\beta_0 + \beta_1 x}} = \frac{1}{1 + e^{-\beta_0 - \beta_1 x}} \quad (20.3)$$

i The logistic function

The logistic function is also known as the sigmoid function, and denoted $S(x)$ or $\sigma(x)$, due to its S-shaped curve. However, the term “sigmoid function” can refer to a family of S-shaped functions.

The term *logistique* was first used in 1845 by [Verhulst \(1845\)](#) to describe the solution of a differential equation describing population growth: $\frac{dp}{dt} = p(1 - p)$. However, Verhulst applied the term *logistique* to the expression of time in terms of population, i.e. essentially $t = \ln(p/(1 - p))$. This is in fact the “logit” function: $\text{logit}(p) = \ln(p/(1 - p))$, which is the *inverse* of the logistic function. The name logit was coined much later – see later footnote.

In python `scipy` and some R packages, the logistic function is referred to as “expit”, making “expit” the inverse of “logit”, just as “log” is the inverse of “exp”.

Just as with linear regression, we can adjust the values of the coefficients β_0 and β_1 to fit the data as best as possible by:

1. Defining an error function (also called a loss function) that measures how good the fit between the data and the model is for any pair of β_0 and β_1 .
2. Adjusting β_0 and β_1 to minimise the error function.

The error function for linear regression was the sum of the squared errors. We need a different error function for logistic regression, which we will derive using probability theory in the chapter on [Maximum likelihood and generalised linear regression](#).

Application to credit example with one variable Figure 20.3 shows the logistic regression of age on credit approval – we are ignoring all the other variables for now. The curve doesn’t look very much like a logistic curve, but that’s because it’s got a very shallow slope, since $\hat{\beta}_1 = 0.03$. We can see that the probability ranges between about 0.37 for teenagers and 0.8 for 70-year-olds.

20.2 Interpretation of logistic regression coefficients

Interpretation of $\hat{\beta}_0$ In linear regression $\hat{\beta}_0$ is the intercept: it tells us the predicted value of the response variable when the predictor variable is 0. In logistic regression $f(\hat{\beta}_0) = 1/(1 + \exp(-\hat{\beta}_0))$ tells us the probability the response variable being 1 (“success”) when the predictor variable is 0. In the credit example it suggests the likelihood of a newborn baby receiving credit approval is $f(-1.176) = 0.236$ – which seems rather high!

Log odds Remember the definition of odds (Equation 20.1). To interpret the coefficient $\hat{\beta}_1$ it helps to rewrite the logistic regression model (Equation 20.3) in terms of **log odds**, i.e. the log of the odds:

$$\text{Log Odds}(\text{Success}) = \ln \frac{P(\text{Success})}{P(\text{Failure})} = \ln \frac{P(\text{Success})}{1 - P(\text{Success})} \quad (20.4)$$

Log odds of 0 mean that success and failure are equally likely: $P(\text{Success}) = P(\text{Failure}) = 0.5$. Positive log odds mean that success is more likely than failure, and vice versa for negative log odds. An increase of 1 unit of the log odds means that the odds increase by a factor of e . As the probability tends towards 1, the log odds tend towards infinity; as the probability tends towards 0, the log odds tend towards negative infinity.

When we express probability in terms of log odds, we sometimes say it has units of “logits”, which stands for *logistic units*. Going back to the example, we can say that when the predictor variable is 0, the log odds of approval are $\hat{\beta}_0 = -1.176$ logits.

The logit function converts the probability of success into the log odds of success to failure²:

$$\text{logit}(p) = \ln \frac{p}{1 - p} \quad (20.5)$$

We can now re-express Equation 20.4 as $\text{Log Odds}(\text{Success}) = \text{logit}(P(\text{Success}))$.

Rewriting the logistic regression model in terms of log odds The probability of a failure is:

$$P(Y = 0|X = x) = 1 - f(\beta_0 + \beta_1 x) = 1 - \frac{e^{\beta_0 + \beta_1 x}}{1 + e^{\beta_0 + \beta_1 x}} = \frac{1}{1 + e^{\beta_0 + \beta_1 x}} = f(-\beta_0 - \beta_1 x) \quad (20.6)$$

We can divide Equation 20.3 by Equation 20.6 to obtain³:

$$\frac{P(Y = 1|X = x)}{P(Y = 0|X = x)} = \frac{P(Y = 1|X = x)}{1 - P(Y = 1|X = x)} = e^{\beta_0 + \beta_1 x} \quad (20.7)$$

The ratio on the left is the odds for success. It tells us how many times more likely the “success” ($Y = 1$) is than the “failure” ($Y = 0$) for any value of x (see Equation 20.1). If we take natural logs of both sides of the equation, we see that the log odds is a linear function of the predictor:

$$\text{logit}(P(Y = 1|X = x)) = \ln \frac{P(Y = 1|X = x)}{1 - P(Y = 1|X = x)} = \ln \frac{P(Y = 1|X = x)}{P(Y = 0|X = x)} = \beta_0 + \beta_1 x \quad (20.8)$$

We can now see that $\hat{\beta}_0$ is the log odds when the predictor variable is equal to 0.

Interpretation of $\hat{\beta}_1$ From Equation 20.8, we can see that the parameter β_1 tells us the increase in the log odds when we increase x by 1 unit.

In other words, when we increase x by 1 the odds multiply by a factor $\exp(\beta_1)$. We refer to this factor as the odds ratio (OR) for the variable x . In this example the $\text{OR} = \exp(0.03) = 1.03$. Thus, for every year of age, you’re 1.03 times more likely to have a loan approved, an effect size of 3%.

20.3 Multiple logistic regression and confidence intervals

Principle of multiple logistic regression Just as with multiple regression, we can extend the logistic regression model (Equation 20.3) to include extra predictor variables $x^{(2)} \dots$ by adding these variables multiplied by corresponding coefficients $\beta_2^{(2)} \dots$:

$$P(Y = 1|X^{(1)} = x^{(1)}, X^{(2)} = x^{(2)}, \dots) = f(\beta_0 + \beta_1 x^{(1)} + \beta_2 x^{(2)} + \dots) \quad (20.9)$$

This equation applies regardless of whether the predictor variables are binary (such as employment status) or continuous (such as age).

²If we have a continuous response variable between 0 and 1 (e.g. the proportion p of organisms killed by a toxin), we could transform the response variable into logits using $\text{logit}(p)$. In fact, logistic regression and the term logit were invented to deal with this sort of data (Berkson, 1944).

³This identity should help to see this:

$$\frac{f(u)}{f(-u)} = \frac{e^u}{1 + e^u} \frac{1 + e^{-u}}{e^{-u}} = e^u$$

	Variable	Coefficient	Odds or OR
$\hat{\beta}_0$	Intercept	-1.969	0.140
$\hat{\beta}_1$	Age	0.029	1.030
$\hat{\beta}_2$	Employed	1.881	6.562

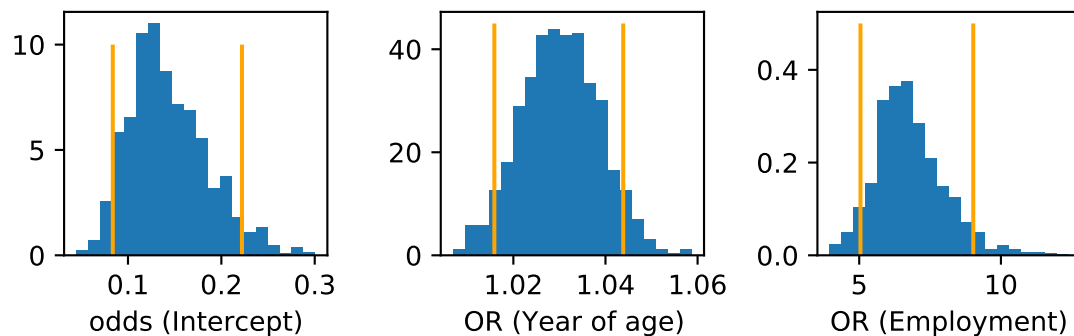
Table 20.2: Coefficients expressed in raw form and as odds ratio $\exp(\beta)$.

Figure 20.4: Bootstrap distributions for the baseline odds and odds ratios for age and employment in the credit scoring example.

Multiple logistic regression applied to credit example If we apply multiple logistic regression to the credit example, we end up with the coefficients and odds ratios shown in Table 20.2. We can see that the effect of being employed increases the odds of being awarded a loan by a factor of 6.56, an effect size of 556%. By contrast each year of age only multiplies the odds by 1.03, an effect size of 3%. To see the effect of increasing age by 10 years, we'd need to raise this OR to the power 10, and would find that the odds are only multiplied by 1.35. The effect of an increase in age from 20 to 70 is about 4.36 – still less than the effect of being in employment.

Bootstrap confidence intervals on coefficients Just as the mean and median are statistics, so are the coefficients $\hat{\beta}_1$ and $\hat{\beta}_2$ in logistic regression. We can therefore use the bootstrap to generate confidence intervals for logistic regression (Figure 20.4). The central estimate and the 95% confidence intervals computed from the 2.5% and 97.5% centiles are:

- Age: OR=1.030, CI=(1.017, 1.044)
- Employment: OR=6.562, CI=(5.110, 8.805)

20.4 Logistic regression as a classifier

Converting linear regression to a classifier Setting a threshold probability p_{thresh} corresponds to setting threshold log odds, which we'll define as c . If we choose log odds $c = 0$, this means odds of 1, i.e. the probability of success (approval) is 1/2. Substituting $c = \ln \frac{P(Y=1|x)}{1-P(Y=1|x)}$ into Equation 20.8, we find:

$$c = \beta_0 + \beta_1 x \quad (20.10)$$

This defines a linear decision boundary – in the region where $\beta_0 + \beta_1 x > c$, the log odds are greater than the threshold, and we classify unseen datapoints in this region as “Success”, and elsewhere, we classify unseen datapoints as “Failure”.

Figure 20.5 shows decision boundaries for various threshold levels when we consider two continuous variables in the credit dataset: age and the log of the income. Note: as with linear regression, it often makes sense with logistic regression to transform variables so that their distribution is as normal as possible.

Transparency of logistic regression The credit agency might want to explain to its customers why their application was or was not approved. Logistic regression makes it very easy to do this, since essentially we have a credit scoring system:

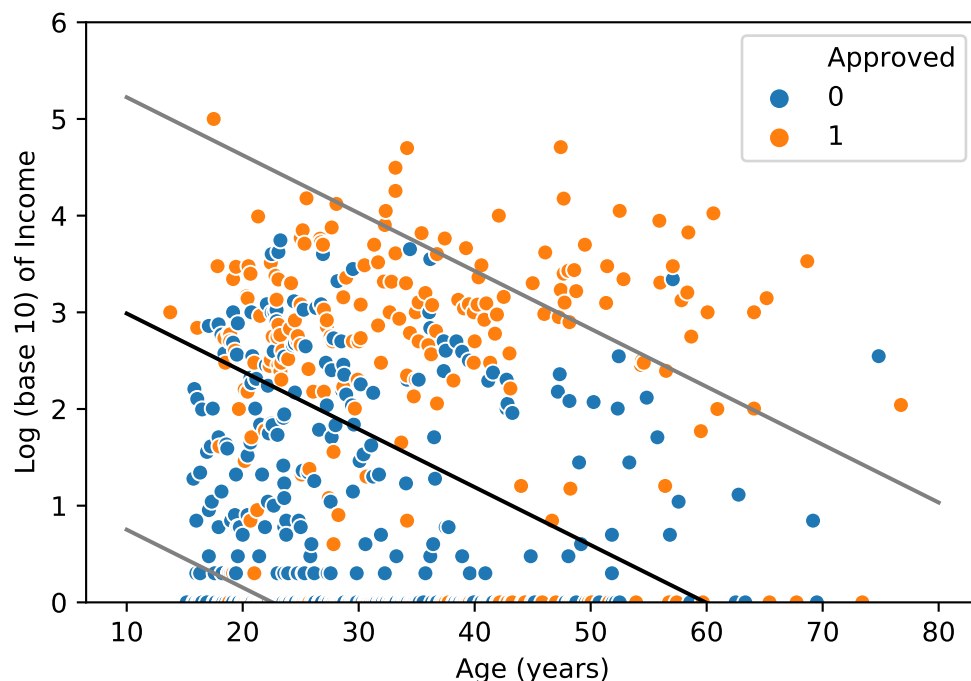


Figure 20.5: Logistic regression applied to credit approval dataset. The age and log income of each application is plotted, along with its approval status. The black line is the decision boundary found by logistic regression with an odds ratio of 1, i.e. log odds $c = 0$. The grey lines are the thresholds corresponding to odds ratios of 3 and 1/3 (i.e. 75%/25% and 25%/75%).

- If you are in employment you score 1.625, if not you score 0
- Multiply your age by 0.029 and add the result to your score
- Round your income to the nearest 1000. Multiply the number of zeros in this figure by 0.320 and add the result to your score⁴
- If you scored more than 2.246, your credit will be approved

Thus, a logistic regression classifier is potentially a very **transparent** classifier. It could help to reduce the ethical harms of data science to individuals by allowing them to understand why their loan was rejected. One of the recommendations of Vallor's *Introduction to Data Ethics* is to "Promote Values of Transparency, Autonomy, and Trustworthiness" (Vallor, 2018).

Logistic regression versus k -nearest neighbour The logistic regression classifier differs in a number of ways from the nearest neighbour classifier:

1. The logistic regression decision boundary (obtained by setting a probability criterion) is a straight line, whereas the nearest neighbour decision boundary is nonlinear (Figure 20.6).
2. The k -NN thus gives more flexibility and the ability to have higher accuracy, but it is also more likely to over-fit, as seen in the chapter on [k-NN](#), [hyperparameters](#), [metrics](#), [cross-validation](#).
3. The logistic regression algorithm is more transparent than k -NN.
4. k -NN classifiers benefit from having standardised predictor variables as inputs; logistic regression doesn't need this, though it can help if we are regularising a logistic regression classifier (which we will not do in this course).

Often it is worth trying logistic regression first in classification problems.

⁴OK, this is an approximation to a log! We could ask people to take logs or provide a table: or make the algorithm itself work in this way.

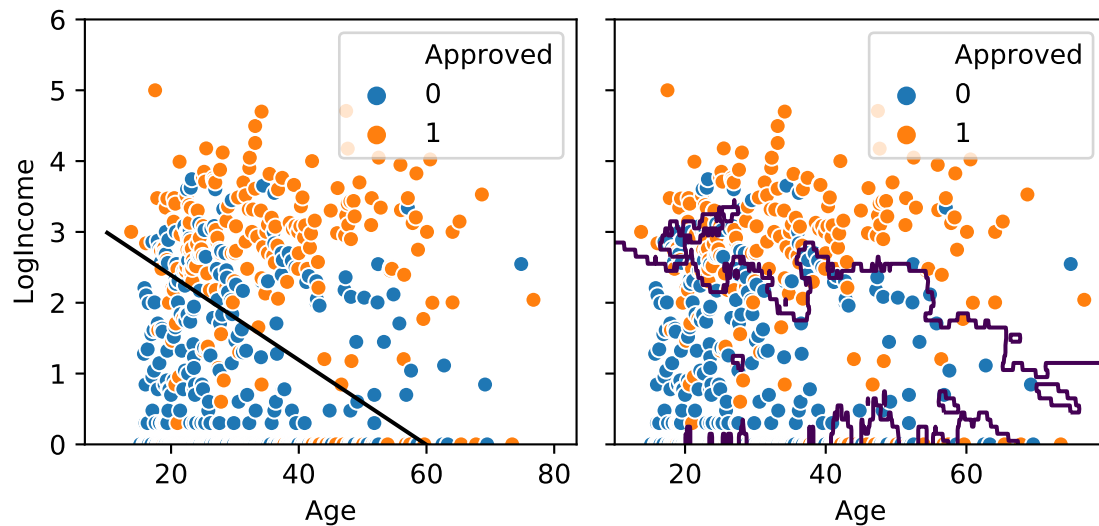


Figure 20.6: Logistic regression (left) versus 11-NN (right) applied to the credit data. The decision boundaries are shown as dark lines.

Related Python Lab: Logistic regression

<https://github.com/Inf2-FDS/FDS-S2-04-logistic-regression>

In this lab you will learn about logistic regression, interpretation of logistic regression coefficients and generating confidence intervals for logistic regression coefficients. By the end of this lab you should be able to:

- identify what transformations would be helpful to variables before applying logistic regression
- apply logistic regression to a dataset
- interpret the coefficients from application of logistic regression
- apply the bootstrap to logistic regression to obtain confidence intervals
- interpret the confidence intervals

Chapter 21

Maximum likelihood and generalised linear regression

21.1 The principle of maximum likelihood

Regression revisited In the chapter on [Linear Regression](#), to find the values of the coefficients $\hat{\beta}_0$ and $\hat{\beta}_1$ we used the principle of least squares, i.e. we adjusted the values of the coefficients to minimise the error function (or loss function), which was the sum of squared errors between the predicted and observed response variables. You might wonder why we chose the error function to be the sum of squared errors rather than another measure, like the sum of the absolute differences between the data points and the regression line. One reason was that we could derive formulae for the values of $\hat{\beta}_0$ and $\hat{\beta}_1$ that minimise the squared error.

In the chapter on [Logistic regression](#), we stated that to obtain the coefficients, we need to minimise an error function, and promised to derive the function later.

In this chapter we:

- introduce an important principle in statistics and machine learning, the principle of maximum likelihood
- apply the principle of maximum likelihood to linear regression leads to the error function for linear regression being a sum of squared errors
- apply the principle of maximum likelihood to derive an error function in the case of logistic regression, where there is a binary outcome
- show how the principle maximum likelihood can be used to derive error functions some types of data are not well fit by either linear regression or logistic regression.

Along the way, to help understand the principle of maximum likelihood, we will consider some simpler examples involving only one variable.

Likelihood We can use a statistical model to generate “fake” data points (the polite way of referring to this “fake” data is “synthetic data”). We’ll start with the model being a univariate distribution, in this example a normal distribution with mean μ and variance σ^2 . We express the idea of the distribution generating data using the following notation:

$$Y \sim \mathcal{N}(\mu, \sigma^2) \quad (21.1)$$

which we can read as “we draw the random variable Y from the normal distribution with mean μ and variance σ^2 ”. We generate n samples $y_1, \dots, y_n$ from the distribution by taking n draws; these samples are the “fake data”.

Now imagine that the data $y_1, \dots, y_n$ is real, and that we ask: out of all the possible values of μ and σ^2 , which of them would have been most likely to generate that data? We define the **likelihood** as the probability of generating the data (here $y_1, \dots, y_n$) given the parameters (here μ and σ^2). In general, if we have a set of m parameters $\vartheta_1, \dots, \vartheta_m$ we can denote the likelihood

$$P(Y = y_1, \dots, y_n | \vartheta_1, \dots, \vartheta_m)$$

The form of the likelihood function will depend on the distribution we’re assuming the data comes from; we will see a number of examples of likelihood functions later.

Principle of maximum likelihood The principle of maximum likelihood states that for a set of observed data and a given statistical model, we adjust the model parameters to maximise the likelihood that the observed data arises from the model. The resulting parameters are referred to as the **maximum likelihood estimates**. The principle of maximum likelihood is also referred to as the **maximum likelihood principle**, or “maximum likelihood”, for short.

21.2 Maximum likelihood principle applied to a simple example

Likelihood of univariate data assuming a normal distribution and independent samples Suppose we have some univariate data $y_1, \dots, y_n$ (for example, as shown in Figure 21.1a) and we make two assumptions about how it was generated:

1. Each sample is drawn independently of the others
2. Each sample is drawn from a normal distribution

In some cases, these assumptions might not be reasonable – for example if we were picking card from a deck one after another without replacement, and recording the rank of the card (with Ace = 1, Jack = 11, Queen = 12 and King = 13). However, we will continue on the basis that they are reasonable assumptions for our data.

From the second assumption (normal distribution), the likelihood of generating *one* data point with a value y is:

$$P(Y = y | \mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y - \mu)^2}{2\sigma^2}\right) \quad (21.2)$$

From the first assumption (independence) the likelihood of generating all the data given the parameters is:

$$P(Y = y_1, \dots, y_n | \mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y_1 - \mu)^2}{2\sigma^2}\right) \times \dots \times \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y_n - \mu)^2}{2\sigma^2}\right) \quad (21.3)$$

Since the data $y_1, \dots, y_n$ are fixed, the likelihood is a function of the parameters μ and σ^2 .

Figure 21.1b shows the likelihood function of generating the data shown in Figure 21.1a. In this case, the highest likelihood arises when the mean μ is close to 0 and when the variance σ^2 is close to 1. Moving away from this point (e.g. to a higher mean or higher variance), the likelihood of having generated the data decreases, which should make sense when we compare at the mean and the standard deviation of the data itself. In fact, we generated the “data” in Figure 21.1a from a normal distribution with a mean of 0 and a variance of 1, so the likelihood function is as expected.

Maximising the likelihood function Although we can see fairly clearly where the maximum is in this example, we want to find a way of determining what values of the parameters maximise the likelihood without having to do a plot every time, or experiment with computing the likelihood for different values of μ and σ^2 . This is another example of an optimisation problem, like finding the minimum of the error function in [Linear Regression](#). Unlike in the case of linear regression, there will be no analytical solution, and we do not have time in this course to go into the many methods of optimisation that could be used to find the maximum likelihood estimates of the parameters. However, we can show some of the steps that lead up to the optimisation.

The first step is to write the likelihood in Equation 21.3 more compactly by using $\prod$ to denote a product, analogously to $\sum$ denoting summation:

$$P(Y = y_1, \dots, y_n | \mu, \sigma^2) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y_i - \mu)^2}{2\sigma^2}\right) \quad (21.4)$$

This rewriting allows what follows to be written more compactly and generally.

The log likelihood It is often helpful to work with the log of the likelihood. In general, whenever each data point is generated independently, the log likelihood of the data is equal to the sum of the log likelihood of generating each data point:

$$\ln P(Y = y_1, \dots, y_n | \vartheta_1, \dots, \vartheta_m) = \ln \prod_{i=1}^n P(Y = y_i | \vartheta_1, \dots, \vartheta_m) = \sum_{i=1}^n \ln P(Y = y_i | \vartheta_1, \dots, \vartheta_m) \quad (21.5)$$

	Data
y_1	1.624345
y_2	-0.611756
y_3	-0.528172
y_4	-1.072969
y_5	0.865408
y_6	-2.301539
y_7	1.744812
y_8	-0.761207
y_9	0.319039
y_{10}	-0.249370

(a) Univariate dataset

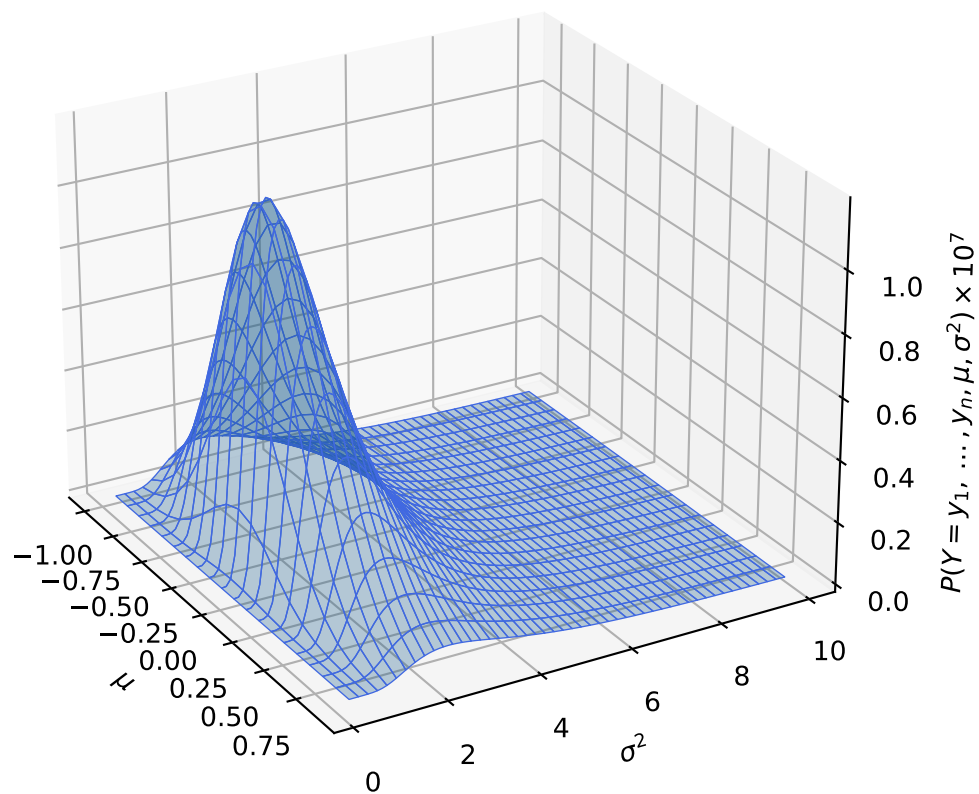
(b) Likelihood that the data in (a) were generated by a normal distribution with a variety of parameters μ and σ^2 .

Figure 21.1: Demonstration of maximum likelihood applied to univariate data.

The equation above holds because the log of a product is the sum of logs of the components of the product ($\ln ab = \ln a + \ln b$).

In the example, when we substitute the probability of one data point (Equation 21.3) into the general Equation (21.5), we obtain the log likelihood as a function ℓ of the parameters (μ and σ^2):

$$\begin{aligned}\ell(\mu, \sigma^2) &= \ln P(Y = y_1, \dots, y_n | \mu, \sigma^2) = \sum_{i=1}^n \ln \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y_i - \mu)^2}{2\sigma^2}\right) \\ &= \sum_{i=1}^n \ln\left(\frac{1}{\sqrt{2\pi\sigma^2}}\right) + \ln\left(\exp\left(-\frac{(y_i - \mu)^2}{2\sigma^2}\right)\right) \\ &= -\frac{n}{2} \ln 2\pi\sigma^2 - \sum_{i=1}^n \left(\frac{(y_i - \mu)^2}{2\sigma^2}\right)\end{aligned}\quad (21.6)$$

We can see that the final is the sum of squared differences between the data points and the mean multiplied by $-1/(2\sigma^2)$. Because of the negative sign, adjusting μ to decrease the sum of the squared differences will increase the likelihood. Increasing σ^2 will make this term smaller (i.e. still negative, but closer to zero). However, increasing σ^2 is penalised by the first term, so we should expect that there is an optimal value of σ^2 .

Figure 21.2a shows the plot of the log likelihood for the univariate data example in Figure 21.1a. We can see that it is much smoother than the plot of the likelihood (Figure 21.1b), which is a helpful property when applying optimisation algorithms. If we plot the log of the variance, the likelihood becomes smoother still (Figure 21.2b). Furthermore, the likelihood is defined for all values of $\log \sigma^2$, compared to only positive values of σ^2 . Having no hard boundary at 0 is also helpful for optimisation.

Why do we use the log likelihood? There are three advantages of using the log likelihood rather than using the likelihood:

1. The sum of logs is easy to represent within the limits of floating point arithmetic.
2. The log likelihood function is smoother than the likelihood function.
3. Sums are easy to differentiate; products are not. Differentiation is helpful for the optimisation required to find the maximum likelihood parameters.

Maximum likelihood estimates of the mean and the variance Because the log function is a monotonically increasing function, the values of β_1 , β_2 and σ^2 that maximise the log likelihood will be exactly the same values that maximise the likelihood. To find the parameter values that maximise the log likelihood, we first partially differentiate ¹ Equation 21.6 with respect to μ and σ^2 . We can then find the values of μ and σ^2 at which the derivatives $\partial\ell/\partial\mu$ and $\partial\ell/\partial\sigma^2$ are zero, which leads to the maximum likelihood estimates:

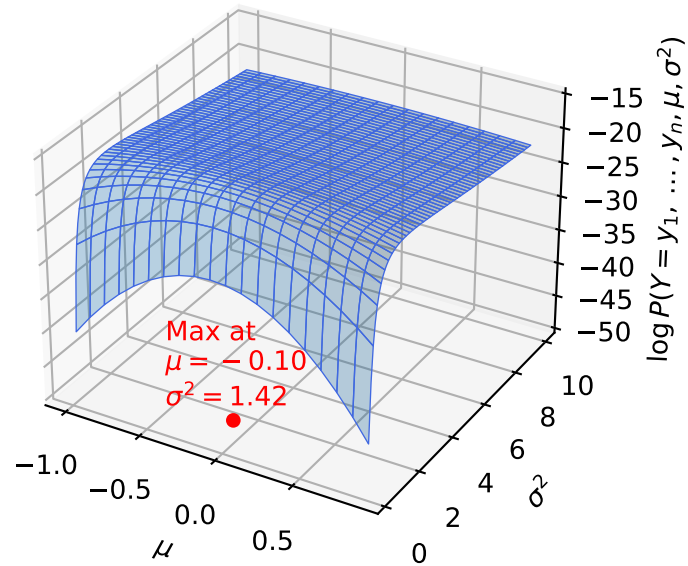
$$\hat{\mu}_{\text{MLE}} = \frac{1}{n} \sum_{i=1}^n y_i \quad \text{and} \quad \hat{\sigma}_{\text{MLE}}^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{\mu})^2 \quad (21.7)$$

The estimate of the parameter μ is the familiar definition of the sample mean seen in the chapters on [Descriptive statistics](#) and [Estimation](#). However, the estimate of the variance has the divisor n , rather than $n-1$ as used previously. In the chapter on [Estimation](#) we show that this would be a biased estimator, though when n is large there is very little difference between the biased and unbiased estimates. In Figure 21.2 we have indicated the maximum likelihood estimates with the red dot; it corresponds clearly to the peak of the log likelihood function.

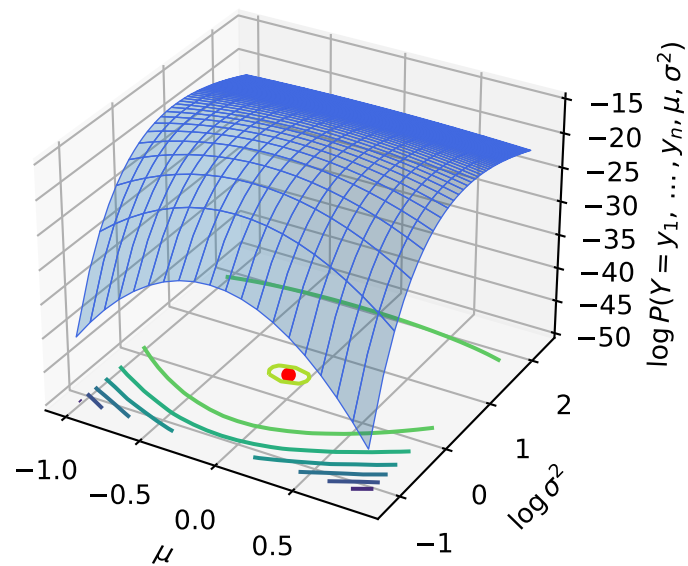
21.3 Maximum likelihood estimation of linear regression coefficients

A probabilistic model for linear regression We can repeat our thinking about the single variable in the context of linear regression. We think of the values x_i of the predictors (Equation 7.1 in the topic on [Linear Regression](#)), as being known, and the response variable as being a random variable Y_i that is a sum of a

¹Most students will not have covered partial differentiation; knowledge of this derivation is not examinable.



(a) Log likelihood of data in Figure 21.1a under a normal distribution with parameters μ and σ^2 .



(b) Log likelihood of data in Figure 21.1a under a normal distribution with parameters μ and σ^2 , but plotted against $\log \sigma^2$ instead of σ^2 .

Figure 21.2: The log likelihood function applied to univariate data.

deterministic linear function of the predictor variable x_i and a random deviation ε_i , which we sometimes call an error term:

$$Y_i = \beta_0 + \beta_1 x_i + \varepsilon_i \quad ; \quad \varepsilon_i \sim \mathcal{N}(0, \sigma^2) \quad (21.8)$$

This equation is equivalent to the variable Y_i being drawn from a normal distribution with mean $\beta_0 + \beta_1 x_i$ and variance σ^2 :

$$Y_i \sim \mathcal{N}(\beta_0 + \beta_1 x_i, \sigma^2) \quad (21.9)$$

This equation is almost the same as Equation 21.1 for the generation of univariate data from a normal distribution, except that we have replaced μ with $\beta_0 + \beta_1 x_i$.

Application of maximum likelihood The likelihood of this probabilistic linear regression model depends on the parameters β_0 , β_1 and σ^2 , as well as the data points. Since we are assuming a normal distribution, we can adapt Equation 21.6 for the log-likelihood of the univariate data by substituting $\beta_0 + \beta_1 x_i$ in place of μ :

$$\begin{aligned} \ell(\beta_0, \beta_1, \sigma^2) &= \ln P(Y = y_1, \dots, y_n | x_1, \dots, x_n | \beta_0, \beta_1, \sigma^2) \\ &= -\frac{n}{2} \ln 2\pi\sigma^2 - \sum_{i=1}^n \left(\frac{(y_i - \beta_0 - \beta_1 x_i)^2}{2\sigma^2} \right) \end{aligned} \quad (21.10)$$

Before trying to visualise the log likelihood, it's worth looking at its structure. From the chapter on [Linear Regression](#), remember that we defined the sum of squared errors (SSE) as:

$$\text{SSE} = \sum (y_i - \hat{y}_i)^2 \quad (21.11)$$

where the predicted value $\hat{y}_i = \beta_0 + \beta_1 x_i$. The second term in Equation 21.10 is the negative of the SSE, divided by $2\sigma^2$. Thus, to maximise the log likelihood, we need to minimise the SSE with respect to β_0 and β_1 , which is exactly what we did when we derived the linear regression coefficients. We thus have a probabilistic motivation for the principle of least squares.

The values of the coefficients We've already calculated $\hat{\beta}_0$ and $\hat{\beta}_1$ using the principle of least squares in the chapter on [Linear Regression](#):

$$\hat{\beta}_0 = \bar{y} - \hat{\beta}_1 \bar{x} \quad ; \quad \hat{\beta}_1 = \frac{\sum x_i y_i - n \bar{x} \bar{y}}{\sum x_i^2 - n \bar{x}^2} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2} \quad (21.12)$$

However, we now have one more parameter to estimate: σ^2 . To do this we maximise Equation 21.10 by differentiating it with respect to σ^2 , and arrive at the maximum likelihood estimator for σ^2 :

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i)^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \frac{\text{SSE}}{n} \quad (21.13)$$

Note that this is a biased estimator. An unbiased estimator, which can be obtained via estimation theory, is:

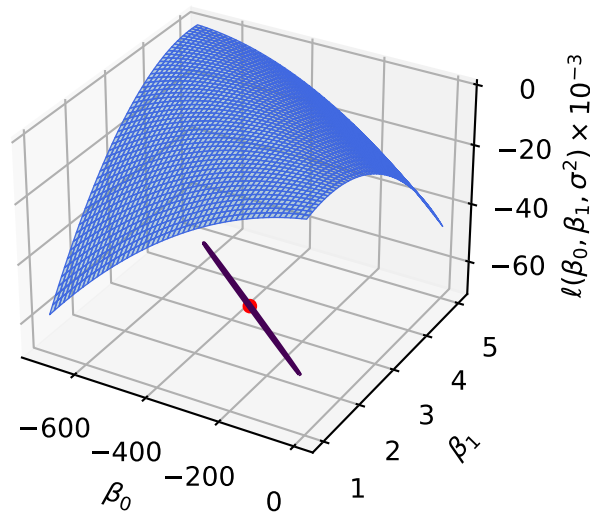
$$\hat{\sigma}^2 = s^2 = \frac{1}{n-2} \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i)^2 = \frac{1}{n-2} \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \frac{\text{SSE}}{n-2} \quad (21.14)$$

because are $n - 2$ degrees of freedom: knowing $\hat{\beta}_0$ and $\hat{\beta}_1$ means that we only need to know $n - 2$ of the y_i to work out the values of the remaining two. In practice, n is large enough that the difference between the two estimates of $\hat{\sigma}^2$ is negligible.

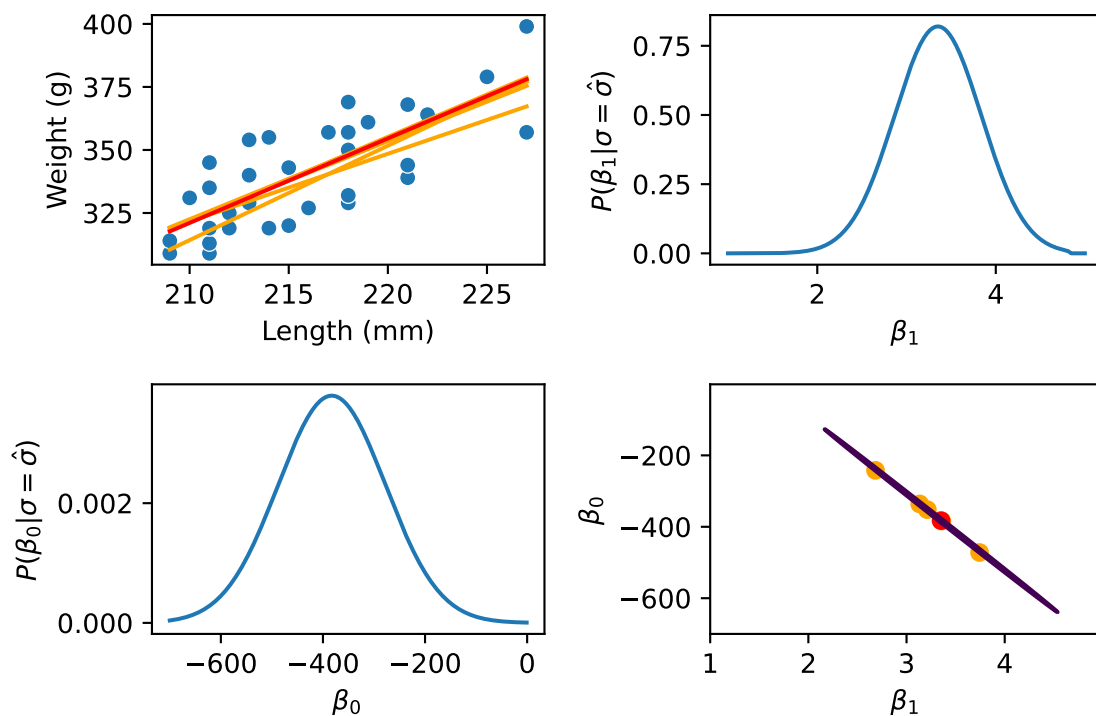
It is also worth noting that, since $(y_i - \hat{y}_i)$ are the residuals, $\hat{\sigma}^2$ is the variance of the residuals.

Uncertainty estimates Because there are three parameters, we cannot visualise the likelihood for regression easily. However, we can examine the likelihood as a function of the intercept and slope when the variance is fixed at its maximum value (Figure 21.3). There is a very ellipse that encloses the most likely values for pairs of (β_0, β_1) . The larger the slope, the lower the gradient, which ensures that the regression line passes close to the centre of the data.

This relationship is the same as when we tried to estimate confidence intervals using the bootstrap ([Statistical inference and regression](#)). Examining the structure of the likelihood function (which depends on



(a) The log likelihood for linear regression applied to the squirrel data (below). The purple contour encloses 95% of the probability in the likelihood function, and the red point shows the maximum likelihood estimate.



(b) Top left: the squirrel data weight and height data, with the regression line in red from the maximum likelihood parameters, and others in orange sampled from the distribution of the values of β_0 and β_1 . Top right: the marginal distribution of the likelihood of the slope β_1 . Bottom left: the marginal distribution of the likelihood of the intercept β_0 . Bottom right: contour enclosing 95% of the probability mass of the likelihood function. Points in red and orange correspond to the maximum likelihood estimates and the parameters sampled to generate alternative regression lines.

Figure 21.3: Demonstration of maximum likelihood applied to regression of squirrel weight on height.

all the data) is quite a different process to the bootstrap, in which we are effectively reporting the maximum likelihood estimate for each bootstrap set of data. Yet despite these differences, the two methods give very similar results.

We can also compute the marginal distributions of β_0 and β_1 . The standard deviations of these distributions correspond to the standard errors of the estimators for β_0 and β_1 that are reported in the output of linear regression from stats packages ([Statistical inference and regression](#)).

21.4 Maximum likelihood estimation of logistic regression coefficients

A probabilistic model for logistic regression To find the likelihood of the data given the parameters in logistic regression, we can follow a similar process to the one for linear regression. In order to apply maximum likelihood, we need a statistical model that can generate the response variable Y_i . Since Y_i is only 1 or 0 (or “Success” or “Fail”) it could be generated by a Bernoulli distribution, which is parameterised by the probability of success p :

$$Y_i \sim \text{Bernoulli}(p) \quad (21.15)$$

In the chapter on [Logistic regression](#), we saw that in logistic regression, the probability of success is related to the predictors by putting a linear model through the logistic function (Equation 20.3):

$$p = f(\beta_0 + \beta_1 x_i) \quad (21.16)$$

where f is the logistic function.

Reversing the generation process, the likelihood of generating a particular datapoint of $Y = y$ (where $y \in \{0, 1\}$) is:

$$P(Y = y|p) = yp + (1 - y)(1 - p) \quad (21.17)$$

This equation is equivalent to saying that if $y = 1$, there was a probability p that it was generated by the Bernoulli process; if $y = 0$ there was a probability of $1 - p$ that it was generated. When we substitute in the dependence of p on the predictors (Equation 21.16), we obtain:

$$P(Y = y|\beta_0, \beta_1) = yf(\beta_0 + \beta_1 x_i) + (1 - y)f(\beta_0 + \beta_1 x_i) \quad (21.18)$$

This equation expresses the fact that if $y = 1$, there was a probability $p = f(\beta_0 + \beta_1 x_i)$ of having generated it, whereas if $y = 0$, there was a probability of $1 - p$ of having generated it.

With this expression for the likelihood of *one* data point, under the assumption of independence of each data point, we can apply the general definition of the log likelihood (Equation 21.5) to obtain the log-likelihood all the data given the parameters β_0 and β_1 :

$$\begin{aligned} \ell(\beta_0, \beta_1) &= \ln P(Y = y_1, \dots, y_n; x_1, \dots, x_n | \beta_0, \beta_1, \sigma^2) \\ &= \sum_{i=1}^n \ln (y_i f(\beta_0 + \beta_1 x_i) + (1 - y_i) f(\beta_0 + \beta_1 x_i)) \end{aligned} \quad (21.19)$$

We now have an expression for the likelihood of an observed data given the model. The likelihood is a function of β_0 and β_1 . However, unlike in the case of linear regression, we can't derive formulae to give the best estimates of the coefficients $\hat{\beta}_0$ and $\hat{\beta}_1$ that maximise the likelihood. We have instead to use a numerical optimisation procedure that gets us to the best estimates.

We now turn to how to estimate the logistic regression coefficients to give the best estimates $\hat{\beta}_0$, $\hat{\beta}_1$ etc. In *linear* regression we minimised the sum of squared errors between the predicted and observed response variables. We could do this analytically, ending up with a formula for the regression coefficients. In logistic regression, it doesn't make sense to minimise the sum of squared errors, since our response variable is only 0 or +1 whereas the predictor variables can have an infinite range.

Intuition of maximum likelihood applied to logistic regression The maximum likelihood principle and derivation may look a bit complicated. We can imagine that the logistic function $P(Y = +1|X = x_i)$ is like a piece of elastic. Datapoints that are “successes” ($y_i = 1$) pull $P(Y = 1|X = x_i)$ upwards towards 1 at the location x_i , since this will make success more likely. Datapoints that are “failures” ($y_i = 0$) pull $P(Y = 1|X = x_i)$ downwards towards 0 at the location x_i . Of course, the successes and failures may be mixed up, in which case they will be competing with each other to pull the logistic function up or down.

Name	Distribution	Link function
Linear regression	Normal	$g(\mu) = \mu$
Logistic regression	Bernoulli	$g(p) = \ln \frac{p}{1-p}$
Poisson regression	Poisson	$g(\lambda) = \ln \lambda$

Table 21.1: Generalised linear models.

21.5 Generalised linear models

We have now used the maximum likelihood principle to derive regression models for two types of response variables: continuous ones in the case of linear regression, and binary variables in the case of logistic regression.

It may be that the response variable can't be described well by either of those patterns: for example suppose we are modelling how the number of arrests per year in a geographic area depends on the Scottish Index of Multiple Deprivation (see chapter on [Dealing with high dimensions – PCA](#)). The number of arrests is zero or a positive integer, so the Gaussian distribution is not a good model for it. The domain of the Poisson distribution is positive integers, so it could be a good choice. We then have the problem of linking the parameter of the Poisson distribution to the predictor variables.

Generalised linear models (GLMs) provide a general framework for dealing with response variables modelled using different distributions. A generalised linear model comprises:

A distribution This is the distribution of the response variable given a parameter μ . The distribution is parameterised so that the expectation distribution is equal to the parameter itself. For example the Bernoulli distribution is parameterised with p and $E[Y] = p$. In this case the domain of p is $[0, 1]$.

A link function The link function g links the parameter of the distribution (or one of the parameters of the distribution) to a linear model of the predictors. Regardless of the domain of the parameter, the link function maps a value of the parameter onto the real number line. For example in the case of logistic regression, the logit function maps p onto the real number line.

A linear model The value of the link function $g(\mu)$ is equal to the prediction from the linear model:

$$g(\mu) = \beta_0 + \beta_1 x_1 + \beta_2 x_2 \dots \quad (21.20)$$

Once the distribution and link function are specified, a maximum likelihood function can be written, and via the magic of optimisation and stats packages, the maximum likelihood estimates of the parameters can be found. Then, to make a prediction from the model, we can invert the link function:

$$\mu = g^{-1}(\beta_0 + \beta_1 x_1 + \beta_2 x_2 \dots) \quad (21.21)$$

Table 21.1 summarises some types of GLM. There are many more, for example in the Python `statsmodels` package.

21.6 From maximum likelihood to Bayesian inference

In this course we have focused on frequentist statistics, the idea that when we estimate there is a true answer, and in the long run our estimate will encompass that true answer a specified fraction of the time. In contrast, in Bayesian statistics, we focus on the probability of the parameters. It's a short hop from maximum likelihood to Bayesian inference, so we'll make the link.

Bayes theorem states that the probability of the parameters given the data depend on the likelihood (which we have been computing in this chapter), and the prior probability of the parameters (how likely particular values are):

$$P(\vartheta|Y = y) = \frac{P(Y = y|\vartheta)P(\vartheta)}{P(Y = y)} \quad (21.22)$$

The denominator is the "probability of the data" also known as the evidence, and is equal to:

$$\int P(Y = y|\vartheta)P(\vartheta)d\vartheta \quad (21.23)$$

In the case where we make no assumptions about the distribution of the parameters ($P(\vartheta)$ is a constant), then the posterior is proportional to the likelihood.

Maximum likelihood elsewhere Maximum likelihood is a widely-used method for deriving estimators for parameters, ranging from simple models like linear regression to more complicated probabilistic machine learning ones. If it is possible to write a likelihood function for a model, it is generally possible to find its maximum.

Chapter 22

Ethical issues with supervised learning



Recommended reading

Equality law can disadvantage women in algorithmic credit decisions

22.1 Fairness in machine learning

Algorithms and fairness Algorithms have an effect on our day-to-day lives in many different domains, for example:

- The COMPAS algorithm, used to predict recidivism in California ([Introduction to data ethics](#)), which gives recommendations to judges who decide whether to award individuals parole.
- When you apply for credit or a mortgage, to buy a home, your data goes into an algorithm to produce a recommendation.
- In education, intelligent tutoring systems employ machine learning to assess how fluent the students are in different types of skills. Based on their inferred skill level, the intelligent tutor sets problems for the students that best match their own personal learning style. The algorithms that infer the students' skills use vast amounts of data collected from interactions with students from all over the world, from different backgrounds.

Supposing that humans were making these decisions about us, we might believe that we were being treated “unfairly”, perhaps because they did not share our interests, because they were distracted or tired, or because they were biased against people from our background, or discriminating against us on the basis of our race or gender.

Algorithms do not suffer from tiredness, and are consistent, but they may still demonstrate biases and be regarded as “unfair”, especially if they are trained on previous decisions made by humans. “Fairness” has multiple definitions, which we don’t have time to cover in this course. However, fairness in machine learning is very important, and this chapter should introduce us to ...

Example: Classification of names by Google Google tries to classify real versus fake names, by training on massive amounts of data collected from the Internet. Most of these names are common U.S. names, and ethnic minority names are rarer. This results in the database being biased toward common U.S. names versus ethnic names. As a result, the algorithm is always going to be better at predicting outcomes that relate to the majority class in the dataset than it does to a minority class. Therefore, someone who has an unusual name is likely to have their name classified as fake by this system.

If someone at Google is using these predictions to make a decision about whether to authorise a new Google account to a user, and they don’t want to authorise the account to someone with a likelihood that has picked a false name, they’re going to be discriminating against a person of ethnic minority who might have a non-standard name.

Questions when training ML algorithms This example raises a number of questions:

- Has the machine learning algorithm been trained on enough representative data in order to treat the population at large with the treatment that we deserve as human beings?
- Has the machine learning algorithm acquired the biases of any human decisions it was trained on?

Prediction and judgement In all these domains the algorithms make predictions. If humans are not “in-the-loop”, the predictions are effectively judgement calls in the sense that they make a decision that affects the lives of those it relates to.

In some systems, there is a human in-the-loop. For example, in the COMPAS algorithm itself does not make the decision of whether to release that person, but it is making a recommendation to a judge, someone who has been educated and trained over many years. If humans are in the loop, it may still be difficult to justify making a different call to the machine.

22.2 Overview of equality legislation

Discrimination We deserve to be treated fairly by other people and institutions. This should include the algorithms that manipulate our data and provide recommendations to people who act upon them. Being “fair” implies that how we are treated should not depend on attributes such as our race or gender, whether we’re applying for a mortgage, applying for parole, or learning mathematics. **Discrimination** arises if one person is treated better than another because of one or more of their attributes that are intrinsic to them.

Equality legislation Legislation such as the EU Equal Treatment Directive and the UK Equality Act of 2010 prohibits discrimination based on **protected characteristics**. In the UK Equality Act, there are nine protected characteristics:

- age
- disability
- gender reassignment
- marriage and civil partnership
- pregnancy and maternity
- race
- religion or belief
- sex
- sexual orientation.

Age has special status and can be used in some cases (credit scoring is one). Prohibiting discrimination can also be seen as requiring **equal treatment** of individuals, or **equal opportunity** for individuals.

Fairness and discrimination Note that a lack of discrimination on its own does not mean that someone is treated fairly; there are ways of being unfair other than discrimination. For example, suppose that an interviewer asks each candidate what their favourite band is, and hires the candidate with preferences most similar to theirs. Discrimination is one manifestation of unfairness.

Just as human decision-makers can exhibit bias, so can algorithms, especially if we have trained them on previous human decisions.

22.3 Case Study: The effect of gender on credit scoring

Credit scoring Dr. Galina Andreeva, from the University of Edinburgh Business School, with Anna Matuszyk wrote a paper called “The Law of Equal Opportunities or Unintended Consequences: the impact of unisex risk assessment in consumer credit” (Andreeva and Matuszyk, 2019); the paper has a [helpful summary](#). This paper is about credit scoring, which is how banks assess the risk of individuals applying for loans. The aim of credit scoring is to compute the probability of default, i.e. the probability that a potential or existing borrower will not pay back a loan. The bank chooses a cut-off probability of default. Those applications with a score above that cut-off will receive the loan and those with a score below that cut-off will not receive the loan. As discussed in the chapter on [Logistic regression](#), the probability of default can be computed for each application based on the attributes such as income, time in current address, occupation, etc.

Protected characteristics and credit scoring A natural question to ask is what attributes (features) should we allow these algorithms to use in order to avoid bias. The equality legislation implies that banks should not use any protected characteristics as features in credit scoring algorithms, such as logistic regression. So, for example, protected attributes in the data should be dropped before training. Thus, any two individuals who only differ on protected attributes should not differ in the likelihood to obtain a loan, and there should be equal treatment.

Age is a special case because on the one hand, it would be wrong to refuse someone a job because they're 55 versus 35. On the other hand, banks should approve a car loan for an individual who is 11 years old. Therefore, age may need to be considered at some point by the decision maker or algorithm.

What is the effect on outcomes of removing protected characteristics? We might hope that by banks practising equal opportunity in credit scoring, there would be **equal outcome**, meaning that as the fractions of loan requests approved would be the same for men and women, or groups corresponding to other protected characteristics.

This study focused on gender, and if there is equal outcome for men and women as a group. The authors had access to a dataset containing a portfolio of car loans from a major bank in an EU country from 2003–2010. The dataset contained the protected attribute of gender, although the bank did not use it in their algorithms. It is often difficult to obtain such datasets, given that using the protected characteristics is not allowed for decision-making.

The dataset contained 78,052 records, 26.7% from women and 73.3% from men. The definition of “default” is defaulting (i.e. not keeping up with repayments) on the loan for 2 months (65 days). Only small fractions of loans had defaults: 1.30% of loans to women and 1.82% of loans to men. Thus, men are more likely to default on their payments than women.

The data was split into a training dataset (80%) and a test set (20%). Because of the very low numbers of “positive” instances (i.e. defaults), stratified sampling was used to generate exactly the same proportions of males and females with and without default in the training and test sets. Four logistic regression models were trained, and their test results reported:

1. Model with Gender, i.e. Gender is retained as a predictor variable (training sample comprising both men and women) – this would be illegal for a bank to use, but is possible in a research setting.
2. Model without Gender, i.e. Gender is removed from the data (training sample comprising both men and women)
3. Model trained and tested only on men
4. Model trained and tested only on women

The models were compared in two aspects:

- how they affect the chances of men/women being offered credit
- predictive accuracy

Training the model Predictor variables to include in the model were selected by significance and predictive accuracy, as described in the paper. In the final full models there were 11 predictor variables, including gender. Some variables related to the applicant (e.g. number of children), and others to aspects of the loan (e.g. the down payment, i.e. how much of the cost was being paid up front). Categorical variables were encoded using indicator variables (see chapter on [Data](#)). Numeric variables (e.g. “car age”) were converted to categorical variables by splitting into ranges (e.g. up to 2 years old, 2 years old, 3–4 years old, etc.) – think about why! The regression coefficients were assessed for significance, with p -values reported (as in the chapters on [Logistic regression](#) and [Statistical inference and regression](#)).

Effect of model on chances of men and women being offered credit in Model 1 (with Gender) For a given cut-off level, men in the sample were more likely to have the loan rejected and women were less likely. For example, if the cut-off is set so that 60% of applicants are rejected, in the model with gender 45% of women in the test set would have their application rejected but 71% of men in the test set would have. This result indicates that in the training set, with all other factors being equal, women had lower default rates in the past. The result doesn't tell us *why* this might be the case – we will return to this question.

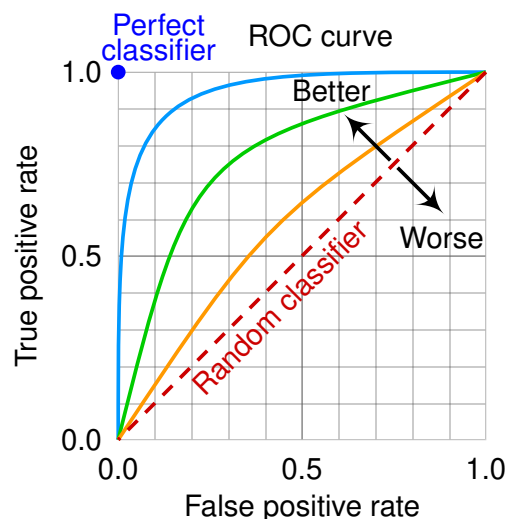


Figure 22.1: The receiver-operator characteristic. The diagonal shows the performance of a random classifier; the area under this curve (AUC) is 0.5. The orange, green and blue lines show the performance of three successively better classifiers; the AUC increases from 0.5 towards 1, which would be the AUC of a perfect classifier, shown by the blue point at the top left. Credit: CMG Lee, Wikimedia Commons, CC BY-SA 4.0.

Effect of model on chances of men and women being offered credit in Model 2 (without Gender) In the model without gender, the chances of being accepted for credit decreased for women, and increased for men, but women were still more likely to get credit for a given cut-off level. For example, if the cut-off is set so that 60% of applicants are rejected, in the model with gender 54% of women in the test set would have their application rejected but 65% of men in the test set would have.

Thus, we can conclude that equal treatment of individuals by ignoring a protected characteristic, as per the equality legislation, does not lead to equal outcome at the group level.

Proxy variables Why is there still an effect of gender, when gender is not one of the predictor variables? The reason is that other variables are correlated with gender, and so the model learns the association between these variables and the outcome of whether there was a default. Such variables are called **proxy variables**. Most coefficients vary little between the two models, but the coefficients relating to typically female or male professions do differ, and will mean that someone with a typically female profession is less likely to default and someone with a typically male profession is more likely to default. Thus, Profession is acting as a proxy variable for gender. We can also imagine that income (women still make less money than men) and employment (women take more parental leave, so they often have less time in employment than do men).

Predictive accuracy We can ask if the model without gender is as accurate as the model with gender. To measure accuracy, we will use a measure called the **area under the curve** (AUC). To understand AUC, we first need to understand the **receiver operator characteristic** (ROC). In the section on [Metrics](#), we encountered selectivity and sensitivity. With threshold-based method like logistic regression, we can control the number of instances classified as positive (default on loan) or negative (no default) by changing the threshold – the lower the threshold, the more instances classified as positive. Since some of these positives will be true positives, the sensitivity metric of the classifier will also increase. However, the classifier will detect fewer negatives, and therefore fewer true negatives, so its selectivity metric goes down. The receiver-operator characteristic (Figure 22.1) plots the Sensitivity (also known as the true positive rate) against the Selectivity subtracted from 1 (also known as the false positive rate). [This interactive demonstration online](#) is also helpful in understanding the ROC.

We hope that there is an optimal threshold at which both the sensitivity and selectivity are high. However, we can assess the performance of the classifier, independent of the threshold chosen, by finding the area under the ROC curve. An AUC of 0.5 corresponds to the diagonal line in the ROC, which corresponds to a random classifier. A perfect classifier has an AUC of 1.

In the credit classification example the AUC of the models with and without gender are both about 0.89 – there is no loss predictive accuracy in the model without gender. However, some individuals will be treated

differently in the models with and without gender. If we measure the AUC by testing only men or only women, we find that the AUC for women is lower (0.80, model 1 or 0.79, model 2) than for men (0.91 for both models 1 and 2). Thus, prediction accuracy is lower for the group less represented in the training data – women made up only about 27% of the sample.

Making equal outcomes By training separate models for men and for women (Models 3 and Models 4) and setting the cut-off for each model so that a given fraction of loan applications are rejected, it is possible to have equal outcomes at the group level. However, there are then different models for men and women, and the coefficients are different for men and for women.

Conclusions The authors conclude that the existing law is not effective in promoting equality in the context of algorithm decision-making. Equal treatment of individuals does not translate into equal outcomes at the group level. The paper adds to the existing evidence that it is not possible to completely remove the effect of a protected characteristic without deleting all correlated characteristics. It also demonstrates that minority segments are dominated by majority ones, due to the lack of relative lack of data.

The authors conclude existing law is not effective in promoting equality in algorithmic decision-making. However, to make the algorithms fair at the group level, separate models for men and women are needed, which then means different treatment, in violation of the equality legislation.

Returning to the question of why women appear to be more reliable at paying loans back, it's perhaps possible that Gender or its proxies are acting as proxies for other variables more directly linked with paying back loans. Alternatively, it may be that women are simply less likely to default on loans than men, with all other factors being equal.

Part VI

Project skills

Chapter 23

Software engineering for data science



Recommended reading

- [Good enough practices in scientific computing Wilson et al. \(2017\)](#)
- [Reproducible Analysis Through Automated Jupyter Notebook Pipelines, Amanda Birmingham](#)
- [The scientific paper is obsolete – historical overview of notebooks](#)

23.1 Reproducible research

What is reproducible research?

- If I have carried out an experiment or done some analysis, I can give you the instructions to re-run the experiment and analysis, and you can reproduce my results
- I can also come back to my files and run the analysis again in 3 months' time – yourself from 3 months ago doesn't answer email!

Requirements for reproducible data analysis

- Data
- Code
- Documentation

Barriers to reproducible research – data

- Data not shared
- Data lost or in incompatible format. Here are some emails I have received when requesting data the authors of scientific papers:
 - 2004 (data collected in 1990s)

I got these data 10 years ago, two operative systems ago, when mass storage was based on magneto-optic disks, thay I do not use any more. I will look for the potassium raw data, but I cannot promiss you to get it easily. Actually, I have to look at the Institute magazine to see if I kept the disks in a box...
 - 2021 (data collected in 2005)

...I think that I do have the data somewhere but I may not have access to it at the moment. ...
 - 2021 (data collected in 1995)

...Sadly the data (if it still exists) is on a dusty floppy disc somewhere – so I wouldn't be able to find or access – otherwise you would have been most welcome. ...

Facilitators of reproducible research – data

- Journal policies
- Data repositories
 - Institutional (e.g. Edinburgh Data Share)
 - Subject-specific (e.g. EBRAINS)
 - General (e.g. Zenodo)
 - Governmental (e.g. data.gov.scot)

Barriers to reproducible research – code

- Code not shared
- Code doesn't run or takes days to get running
- Code runs but gives different results

Although some journals now ask peer reviewers to rerun and verify code, sharing it publicly is still far from an academic norm. The amount of time researchers have to spend either helping people use their software or refuting claims stemming from its misuse is a “big worry” among many academics, says Neil Chue Hong, founding director of the Software Sustainability Institute in Edinburgh. “There are ways you can run the code that mean you won't get sensible results, but the researchers who use the code know what those ways are,” he says. “It's like you or me being given a Formula One racing car, and being surprised when it crashes when it goes around the first corner.” (Chawla, 2020)

- Why?
 - Missing code – not tested on clean system
 - Buggy code, maintained within a group
 - Changes in programming languages, e.g. libraries

Facilitators to reproducible research – code

- Version control and code repositories: git and github
- Unit testing: continuous integration (e.g. Travis-CI) helps to identify problems if underlying libraries get uploaded
- Conda environments, including specification of library versions used
- Virtual environments/containers
- Notebooks?

Barriers and facilitators to reproducible research – documentation Barriers:

- What do the columns in the tables mean?
- How were they collected?
- How do I run the code?

Basic solution: a README file explaining all the above!

23.2 Notebooks versus programs**Why notebooks?**

Notebooks

- Mathematica (first released in 1988) was first package to have notebook interface
- Idea of combining text with computation or maths to produce a living paper, which is also reproducible
- Related to literate programming (invented by Donald Knuth author of T_EX, which is written as a literate program)¹
- The open-source Jupyter (first released under the name iPython in 2011) system² has implemented the notebook principle for Python, R and Julia and many other languages³

Advantages with notebooks

- Good for storing a one-off analysis of a few datasets, and producing figures and visualisation
- Helps reproducibility by recording list of steps in data processing
- Many packages have python interfaces – e.g. computational neuroscience packages
- Can be used on a remote server – can be helpful when data has to remain isolated (e.g. health records)

Problems with notebooks

- What if I use the same set of steps on a new dataset
 - e.g. I've run some analysis on one set of gene sequencing data – now I want to repeat it on another set?
 - Do I edit my notebook with the new dataset?
 - Or create a function?
 - We want **automation** rather than **interactivity**
- Inconsistent state [Demo]
- Collaboration on large projects
- Object-oriented code
- <http://compbio.ucsd.edu/reproducible-analysis-automated-jupyter-notebook-pipelines/>
- Version control

Best of both worlds

- If you're wanting to repeat analyses, or you're finding that you are repeating snippets between in notebooks, it might be time to create a Python module to contain the functions

23.3 Data and code management

Data versus code Many of the suggestions in this section are based on the good advice in *Good enough practices in scientific computing* (Wilson et al., 2017).

- Version control for code (e.g. using Github) is a **good thing** but does it work for data?
- VC systems deal well with “small” text files – kilobytes rather than megabytes, and definitely not gigabytes
- Thus they are good for code but not so good for data, especially large datasets
- Thus we need different solutions for storing and keeping track of changes to data and code

¹Donald Knuth (the inventor of T_EX, the basis for L^AT_EX) introduced **literate programming** in 1984. The essential idea is that rather than writing in computer code, programmers embed the code in a description of the logic of the program, in human readable form. There is thus one source file (or set of source files) containing human both human-readable documentation and computer code. A preprocessor is used for two operation on these source files: (1) “tangle” – extract computer-readable source code that can be run; (2) “weave”: extract human-readable documentation, with code embedded. Potential advantages of literate programming are a higher-quality program, and that the author can recall thought-processes when coming back to the program. T_EX, upon which L^AT_EX is based, is written as a literate program.

²Jupyter notebook is considered to be a highly influential piece of scientific software; see, e.g., *Ten computer codes that changed science* (Perkel, 2021).

³Although Mathematica is a closed source package, Wolfram Research have made the Wolfram engine available as a **kernel for Jupyter notebooks**. Thanks to 2020/21 FDS students for alerting us to this.

Data

- Save the raw data
 - Don't overwrite with a better, cleaned-up version
 - Protect, e.g. with file permissions
 - If downloading from a database, record details used to obtain data (exact query, date of retrieval, version of db on that date)
- Backup the data
- Save and share a clean version of the data
 - Cleaned
 - Open data format, e.g. CSV, JSON, YAML, XML
 - Meaningful variable names and file names
 - Metadata about meaning of columns (if not clear from original dataset)
- Make sure you the cleaned dataset is "Tidy"
 - One observation per row
 - One variable per column
 - Unit not stored with numbers – store in metadata or separate column
 - Ideally unique ID for each observation
- If generating data, share in a repository

Code

- Version control – git is currently dominant
- Documentation
- Specifying versions – this can be done using Conda environments and .req files

```
jupyter=1.0.0
matplotlib=2.2.3
numpy=1.15.0
pandas=0.23.4
scikit-learn=0.19.1
scipy=1.1.0
seaborn=0.9.0
python-graphviz=0.8.4
```

Jupyter notebooks and version control

- Problem: Jupyter notebooks are written in JSON, so the diffs stored in version control systems are not very intelligible
- (R markdown doesn't suffer from this problem)
- Workaround: packages such as **nbdime**

Chapter 24

Writing skills

List of datasets

p.	§	Description
13	2.6.0	Titanic
13	2.6.0	Life expectancy (WHO). https://www.kaggle.com/datasets/kumarajarshi/life-expectancy-who
13	2.6.0	Drinks by country
14	2.6.0	Bad form entry data
14	2.6.0	University of Edinburgh timetables
15	3.2.0	Squirrels (Wauters and Dhondt, 1989). Data scraped from paper
26	5.2.0	Informatics Forum Electricity consumption. Collected by David Sterratt
36	5.6.0	CO ₂ and other Greenhouse Gas Emissions (Ritchie et al., 2023)
36	5.6.0	Diabetes (Kahn, 1994). Categorical dep variable, mostly numeric indep variables.
36	5.6.0	Drinks by country
36	5.6.0	Diabetes (Kahn, 1994). Categorical dep variable, mostly numeric indep variables.
38	6.1.0	Online learning trial (Alpert et al., 2016). Randomised trial data with numeric outcome and numeric and categorical covariates. From Replication data for: A Randomized Assessment of Online Learning
41	6.3.0	Wages data (Wooldridge, 2020). Observational data with numeric outcome and mostly numeric explanatory variables. From wooldridge: 115 Data Sets from "Introductory Econometrics: A Modern Approach, 7e" by Jeffrey M. Wooldridge
49	7.1.0	Galton's heights. From CSV in Data 8
55	7.3.0	World population 1940–2000 (Klein Goldewijk et al., 2017). Source: HYDE 3.2 database https://dataportaal.pbl.nl/downloads/HYDE
57	7.3.0	Synthetic dataset showing heteroscedasticity
58	7.4.0	Diabetes (Kahn, 1994)
60	8.2.0	Galton's heights
62	8.4.0	Grades (Edge and Friedberg, 1984). downloaded from website accompanying Devore and Berk (2012), where it is Example 12.25 https://extras.springer.com/2012/978-1-4614-0390-6.zip
69	9.2.0	Synthetic dataset showing correlation
71	10.1.0	Scottish Index of Multiple Deprivation, 2016 edition. https://simd.scot
77	10.3.0	Grades (Edge and Friedberg, 1984)
83	10.4.0	Breast cancer (Wolberg et al., 1995)
87	11.1.0	Oranges, lemons and apples (Murray, 2006)
103	12.4.0	Wine quality
105	13.1.0	Oranges, lemons and apples (Murray, 2006)
113	13.4.0	Breast cancer (Wolberg et al., 1995)
117	14.1.0	Squirrels (Wauters and Dhondt, 1989). Data scraped from paper
122	14.2.0	Swain versus Alabama
126	14.4.0	Basketball players. From http://www.basketball-reference.com and http://www.sportrac.com
126	14.4.0	Swain versus Alabama
139	16.3.0	Japanese restaurant data
145	16.6.0	Basketball players
148	17.1.0	Swain versus Alabama
151	17.3.0	Alameda County Jury pools. The North California branch of the American Civil Liberties Union (ACLU) investigated the numbers of Caucasian, Black/African American, Hispanic, Asian/Pacific Islander and Other people on jury panels in Alameda County.
160	18.5.0	Grades (Edge and Friedberg, 1984)
168	19.1.0	Squirrels (Wauters and Dhondt, 1989)
175	20.1.0	Credit approval dataset, from UCI.

Bibliography

- Adhikari, A., DeNero, J. et al. (2020). 'Computational and inferential thinking: The foundations of data science'. Online, Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International (CC BY-NC-ND 4.0), retrieved at various times between June 2020 and March 2021. URL <https://www.inferentialthinking.com>
- Alpert, W. T., Couch, K. A. et al. (2016). 'A randomized assessment of online learning'. *American Economic Review* 106:378–82. URL <https://www.aeaweb.org/articles?id=10.1257/aer.p20161057>
- Anaconda (2020). 'The state of data science 2020: Moving from hype toward maturity'. URL <https://www.anaconda.com/state-of-data-science-2020>
- Andreeva, G. and Matuszyk, A. (2019). 'The law of equal opportunities or unintended consequences?: The effect of unisex risk assessment in consumer credit'. *Journal of the Royal Statistical Society: Series A (Statistics in Society)* 182:1287–1311. URL <https://rss.onlinelibrary.wiley.com/doi/abs/10.1111/rssa.12494>. <https://rss.onlinelibrary.wiley.com/doi/pdf/10.1111/rssa.12494>
- Anscombe, F. J. (1973). 'Graphs in statistical analysis'. *The American Statistician* 27:17–21. URL <http://www.jstor.org/stable/2682899>
- BBC (2022). 'German murderer wins 'right to be forgotten''. URL <https://www.bbc.co.uk/news/world-europe-50579297>. Retrieved 4 February 2024
- Berkson, J. (1944). 'Application of the logistic function to bio-assay'. *Journal of the American Statistical Association* 39:357–365
- Bornioli, A., Lewis-Smith, H. et al. (2020). 'Body dissatisfaction predicts the onset of depression among adolescent females and males: a prospective study.' *J Epidemiol Community Health* 75:343–348. URL <https://dx.doi.org/10.1136/jech-2019-213033>
- Chawla, D. S. (2020). 'Critiqued coronavirus simulation gets thumbs up from code-checking efforts'. *Nature* 582:323–324. URL <https://doi.org/10.1038/d41586-020-01685-y>
- Dasu, T. and Johnson, T. (2003). *Exploratory data mining and data cleaning*. Wiley
- Davies, B. (2020). 'Is web scraping legal in 2020?' URL <https://scrapediary.com/is-web-scraping-legal/>. Retrieved on 15 October 2022
- Densmore, J. (2017). 'Ethics in web scraping'. URL <https://towardsdatascience.com/ethics-in-web-scraping-b96b18136f01>. Retrieved on 15 October 2022
- Devore, J. and Berk, K. (2012). *Modern Mathematical Statistics with Applications*. Springer Texts in Statistics. Springer, New York, second ed.
- Edge, O. P. and Friedberg, S. H. (1984). 'Factors affecting achievement in the first course in calculus'. *Journal of Experimental Education* 52:136–140
- Gelman, A., Carlin, J. B. et al. (2004). *Bayesian Data Analysis*. Chapman & Hall/CRC, Boca Raton, second ed.
- Gelman, A. and Nolan, D. (2017). *Teaching statistics – a bag of tricks*. Oxford University Press
- Harford, T. (2014). 'Big data: Are we making a big mistake?' URL <https://timharford.com/2014/04/big-data-are-we-making-a-big-mistake/>. Retrieved on 11 October 2025

- Harriss, L., Fearn, J. et al. (2023). 'Data science skills in the UK workforce'. POSTnote 697, UK Parliament. URL <https://doi.org/10.58248/PN697>
- Hastie, T., Tibshirani, R. et al. (2009). *The elements of statistical learning*. Springer, New York, second ed. URL <http://dx.doi.org/10.1007/b94608>
- Hill, K. (2022). 'A dad took photos of his naked toddler for the doctor. Google flagged him as a criminal'. *The New York Times* URL <https://www.nytimes.com/2022/08/21/technology/google-surveillance-toddler-photo.html>. Retrieved 4 February 2024
- Ivory, S. B. (2021). *Becoming a critical thinker: for your university studies and beyond*. Oxford University Press
- Kahn, M. (1994). 'Diabetes'. UCI Machine Learning Repository. URL <https://doi.org/10.24432/C5T59G>. Data was made available for the 1994 AAAI Spring Symposium on Artificial Intelligence in Medicine (AIM-94), Palo Alto, CA
- Klein Goldewijk, K., A., Beusen, J. D. et al. (2017). 'Anthropogenic land use estimates for the holocene; HYDE 3.2'. *Earth System Science Data* 9:927–953
- Kohavi, R., Henne, R. M. et al. (2007). 'Practical guide to controlled experiments on the web: Listen to your customers not to the HiPPO'. In P. Berkhin, R. Caruana, X. Wu and S. Gaffney, eds., *KDD-2007 Proceedings of the thirteenth ACM SIGKDD international conference on knowledge discovery and data mining*, pp. 959–967. Association of Computing Machinery, New York, USA
- Kramer, A. D., Guillory, J. E. et al. (2014). 'Experimental evidence of massive-scale emotional contagion through social networks.' *Proc Natl Acad Sci U S A* 111:8788–90. URL <https://doi.org/10.1073/pnas.1320040111>
- Lane-Clayton, J. E. (1926). *A Further Report on Cancer of the Breast With Special Reference to its Associated Antecedent Conditions*. No. 32 in Reports on Public Health and Medical Subjects. HMSO, London
- Lord, C. G., Ross, L. et al. (1979). 'Biased assimilation and attitude polarization: The effects of prior theories on subsequently considered evidence'. *Journal of Personality and Social Psychology* 37:2098–2109
- MacKay, D. J. C. (2003). *Information theory, inference and learning algorithms*. Cambridge University Press, Cambridge, UK, sixth ed. URL <http://www.inference.phy.cam.ac.uk/itprnn/book.pdf>
- Murray, I. (2006). 'Oranges, lemons and apples dataset'. URL http://homepages.inf.ed.ac.uk/imurray2/teaching/oranges_and_lemons/
- Pearl, J. and Mackenzie, D. (2018). *The Book of Why*. Allen Lane, London
- Perkel, J. M. (2021). 'Ten computer codes that transformed science'. *Nature* 589:344–348. URL <https://doi.org/10.1038/d41586-021-00075-2>
- Press, David J; Pharoah, P. (2010). 'Risk factors for breast cancer: A reanalysis of two case-control studies from 1926 and 1931'. *Epidemiology* 21:566–572. URL <https://doi.org/10.1097/EDE.0b013e3181e08eb3>
- Ritchie, H., Rosado, P. et al. (2023). 'CO₂ and greenhouse gas emissions'. Our World in Data. URL <https://ourworldindata.org/co2-and-greenhouse-gas-emissions>
- Rutherford, A. (2022). *Control: the Dark History and Troubling Present of Eugenics*. Weidenfeld & Nicolson, London
- Schwabish, J. (2021). *Better data visualizations: a guide for scholars, researchers and wonks*. Colombia University Press
- Scottish Government (2016). 'Scottish index of multiple deprivation (SIMD) 2016'. URL <https://www.webarchive.org.uk/wayback/archive/20200117165925/https://www2.gov.scot/SIMD>
- Shea, J. M. (2024). *wooldridge: 115 Data Sets from "Introductory Econometrics: A Modern Approach, 7e" by Jeffrey M. Wooldridge*. URL <https://CRAN.R-project.org/package=wooldridge>. R package version 1.4-4

- Spencer, E. A. and Heneghan, C. (2018). 'Catalogue of bias collaboration. confirmation bias. in: Catalogue of bias 2018'. URL <https://www.catalogueofbiases.org/biases/confirmationbias>
- The Turing Way Community (2022). *The Turing Way: A Handbook for Reproducible Data Science (Version v1.0.3)*. Zenodo. URL <http://doi.org/10.5281/zenodo.6909298>
- Tufte, E. (1982). *The visual display of quantitative information*. Graphics Press, Cheshire, Connecticut
- Tufte, E. (2001). *The visual display of quantitative information*. Graphics Press, Cheshire, Connecticut, second ed.
- Tufte, E. R. (2006). *Beautiful evidence*. Graphics Press, Cheshire, Connecticut
- Tukey, J. W. (1962). 'The future of data analysis'. *The Annals of Mathematical Statistics* 33:1–67
- Tukey, J. W. J. W. (1977). *Exploratory data analysis*. Addison-Wesley Pub. Co., Reading, Mass
- Vallor, S. (2018). 'An introduction to data ethics'. Online. URL <https://www.scu.edu/media/ethics-center/technology-ethics/IntroToDataEthics.pdf>
- Verhulst, P.-F. (1845). 'Recherches mathématiques sur la loi d'accroissement de la population'. *Nouveaux mémoires de l'Académie Royale des Sciences et Belles-Lettres de Bruxelles* 18:4–55. URL [https://gdz.sub.uni-goettingen.de/id/PPN129323640_0018?tify={%22pages%22:\[14\],%22panX%22:0.459,%22panY%22:0.815,%22view%22:%22info%22,%22zoom%22:0.721}](https://gdz.sub.uni-goettingen.de/id/PPN129323640_0018?tify={%22pages%22:[14],%22panX%22:0.459,%22panY%22:0.815,%22view%22:%22info%22,%22zoom%22:0.721})
- Verma, I. M. (2014). 'Editorial expression of concern: Experimental evidence of massive-scale emotional contagion through social networks.' *Proc Natl Acad Sci U S A* 111:10779. URL <https://dx.doi.org/10.1073/pnas.1412469111>
- Vopson, M. M. (2021). 'The world's data explained: how much we're producing and where it's stored'. *The Conversation* URL <https://theconversation.com/the-worlds-data-explained-how-much-were-producing-and-where-its-all-stored-159964>
- Wainer, H. (1999). 'The most dangerous profession: A note on nonsampling error'. *Psychological Methods* 4:250–256
- Wasserstein, R. L. and Lazar, N. A. (2016). 'The ASA statement on p -values: Context, process, and purpose'. *The American Statistician* 70:129–133. URL <https://doi.org/10.1080/00031305.2016.1154108>
- Wauters, L. A. and Dhondt, A. A. (1989). 'Variation in length and body weight of the red squirrel (*Sciurus vulgaris*) in two different habitats'. *Journal of Zoology* 217:93–106
- Wexler, S., Shaffer, J. et al. (2017). *The Big Book of Dashboards: Visualizing Your Data Using Real-World Business Scenarios*. Wiley
- Wickham, H. (2014). 'Tidy data'. *Journal of Statistical Software* 59. URL <https://www.jstatsoft.org/index.php/jss/article/view/v059i10/v59i10.pdf>
- Wilkinson, L. (2005). *The Grammar of Graphics*. Springer, New York, second ed.
- Willsher, K. (2022). 'French tax officials use AI to spot 20,000 undeclared pools'. *The Guardian* URL <https://www.theguardian.com/world/2022/aug/29/french-tax-officials-use-ai-to-spot-20000-undeclared-pools>. Retrieved 4 February 2024
- Wilson, G., Bryan, J. et al. (2017). 'Good enough practices in scientific computing'. *PLOS Computational Biology* 13:1–20. URL <https://doi.org/10.1371/journal.pcbi.1005510>
- Wolberg, W., Mangasarian, O. et al. (1995). 'Breast cancer Wisconsin (diagnostic)'. UCI Machine Learning Repository. URL <https://doi.org/10.24432/C5DW2B>
- Wolfe, J. M. and Horowitz, T. S. (2004). 'What attributes guide the deployment of visual attention and how do they do it?' *Nat Rev Neurosci* 5:495–501
- Wolfe, J. M. and Horowitz, T. S. (2017). 'Five factors that guide attention in visual search.' *Nat Hum Behav* 1:0058

- Wooldridge, J. M. (2020). *Introductory Econometrics: A Modern Approach*. Cengage, seventh ed.
- Xu, D. and Tian, Y. (2015). 'A comprehensive survey of clustering algorithms'. *Annals of Data Science* 2:165–193. URL <https://doi.org/10.1007/s40745-015-0040-1>
- Yanai, I. and Lercher, M. (2020). 'A hypothesis is a liability'. *Genome Biol* 21:231. URL <https://doi.org/10.1186/s13059-020-02133-w>

Index

- 1-hot encoding, *see* indicator variable
- A/B testing, 155
- accuracy, *see* classification accuracy
- adjusted coefficient of determination, 62
- agglomerative clustering, *see* clustering
- algorithmic bias, *see* bias
- alternative hypothesis, 147
- API, 41
- area under the curve, 196
- association, 39
- asymptotically normal, 125
- attribute, 8
- bandwidth, 31
- bar plot, 31
- bias, 23
 - algorithmic, 24
 - confirmation, 23
 - gender, 5
 - in machine learning, 24
 - measurement, 24
 - of an estimator, 129
 - racial, 5
 - sampling, 23, 40, 121
 - selection, 23, 157
- biased estimator, *see* estimator, biased
- bimodal histogram, 16
- bin, 15
- bivariate, 9
- bootstrap estimator, 139, 141
 - applied to linear regression, 167
 - applied to logistic regression, 179
 - for confidence interval of the mean, 141
 - for difference between means, 160
 - in A/B test, 157
- boxplot, 32, 160
- capture-recapture, 128
- case-control study, 39
- categorical data, 20
- causal link, 45
- Central Limit Theorem, 131
- Central Limit Theorem, 125, 135, 159, 162
- Chebyshev's inequality, 126
- cherry-picking, 23, 154
- chi-squared, 152
- classification, 87, 105, 175
- classification accuracy, 98
- classification error, 87
- classification error rate, 91, 96, 98
- classifier, 87
 - baseline, 98
- clustering, 105
 - hierarchical, 108
 - agglomerative, 108
 - top-down, 108
 - K-means, 108
 - batch, 113
 - online, 113
 - minimum variance, 112
 - partitional, 108
- coefficient of determination, 58
- Cohen's d , 162
- cohort study, 39
- collinear variables
 - in multiple regression, 65, 70
- component score, 72, 81
- conditioning, 46
- confidence interval, 127, 135, 138
- confidence interval estimation, 118
- confounding variable, 45
- contingency table, 175
 - two-way, 153
- control group, 38
- correlation coefficient, 58
- covariance matrix, 68
- covariate, *see* predictor variable
- credibility interval, 164
- credit scoring, 194
- cross-validation, 98
 - holdout, 98
 - K-fold, 101
- cumulative distribution function, 139
- cumulative scree plot, 77
- curse of dimensionality, 71, 113
- data, 7
 - structured, 7
 - unstructured, 8
- data compression, 105
- data dredging, 154
- data interpretation, 105
- data matrix, 8
- data science, 3
- data snooping, *see* data dredging
- data stories
 - Gallup and the *Literary Review*, 40
 - John Snow and the Broad Street pump, 39

- data wrangling, [10](#)
- datasets
 - Alameda County Jury pools, [151](#)
 - Bad form entry data, [14](#)
 - Basketball players, [126](#), [145](#)
 - Breast cancer ([Wolberg et al., 1995](#)), [83](#), [113](#)
 - CO₂ and other Greenhouse Gas Emissions ([Ritchie et al., 2023](#)), [36](#)
 - Credit approval dataset, from UCI., [175](#)
 - Diabetes ([Kahn, 1994](#)), [36](#), [58](#)
 - Drinks by country, [13](#), [36](#)
 - Galton's heights, [49](#), [60](#)
 - Grades ([Edge and Friedberg, 1984](#)), [62](#), [77](#), [160](#)
 - Informatics Forum Electricity consumption, [26](#)
 - Japanese restaurant data, [139](#)
 - Life expectancy (WHO), [13](#)
 - Online learning trial ([Alpert et al., 2016](#)), [38](#)
 - Oranges, lemons and apples ([Murray, 2006](#)), [87](#), [105](#)
 - Scottish Index of Multiple Deprivation, 2016 edition. <https://simd.scot>, [71](#)
 - Squirrels ([Wauters and Dhondt, 1989](#)), [15](#), [117](#), [168](#)
 - Swain versus Alabama, [122](#), [126](#), [148](#)
 - Synthetic dataset showing correlation, [69](#)
 - Synthetic dataset showing heteroscedasticity, [57](#)
 - Titanic, [13](#)
 - University of Edinburgh timetables, [14](#)
 - Wages data ([Wooldridge, 2020](#)), [41](#)
 - Wine quality, [103](#)
 - World population 1940–2000 ([Klein Goldewijk et al., 2017](#)), [55](#)
- decision boundary, [87](#), [88](#)
- degrees of freedom, [19](#)
- density, *see* relative frequency density
- density plot, [31](#)
- dependent variable, *see* response variable
- descriptive statistics, [15](#)
- design matrix, [68](#), [68](#)
- deviation from the mean, [18](#)
- dichotomous, [176](#)
- dimensionality reduction, [71](#), [76](#)
- discrimination, [194](#)
- distance
 - between two data points, [108](#)
 - Euclidean, [108](#)
- dummy variable, *see* indicator variable
- effect size, [34](#), [176](#)
- elbow, *see* scree plot, elbow
- empirical distribution, [16](#), [140](#)
- empirical probability, *see* relative frequency
- endogenous variable, *see* response variable
- equal opportunity, [194](#)
- equal outcome, [195](#), [196](#)
- equal treatment, [194](#), [195](#), [196](#)
- error function, [52](#), [67](#), [177](#), [183](#)
- error rate, *see* classification error rate
- error term, [171](#), [188](#)
- estimated standard error, [171](#)
- estimated standard error of an estimator, [131](#)
- estimated standard errors, [157](#)
- estimation, [117](#), [127](#)
- estimator
 - biased, [129](#), [188](#)
 - unbiased, [19](#), [129](#), [188](#)
- ethical case studies
 - AI-assisted tax collection in France, [21](#)
 - Facebook emotional contagion experiment, [24](#)
 - Incorrect classification of images of children by Google, [21](#)
 - murderer wins right to be forgotten, [23](#)
 - OK Cupid web scraping, [24](#)
- ethical reasoning, [22](#)
- ethical sensitivity, [22](#)
- ethics, [21](#)
- exogenous variable, *see* predictor variable
- experiment
 - scientific, [38](#)
- explanatory variable, *see* predictor variable
- Exploratory Data Analysis, [25](#)
- feature vector, [87](#), [87](#), [175](#)
- features, *see also* predictor variable, [87](#)
- filtering, [10](#)
- first principal component, [72](#), [82](#)
- fold, [101](#)
- frequency, [15](#), [31](#)
- frequency density, [31](#)
- General Data Protection Regulation Rights, [22](#)
- generalisation, [93](#), [96](#)
- generalised linear model, [191](#)
- goodness-of-fit, [152](#)
- grammar of graphics, [29](#)
- heteroscedasticity, [57](#)
- hierarchical clustering, *see* clustering
- histogram, [15](#), [30](#)
 - frequency, [15](#), [16](#)
 - frequency density, [15](#)
 - relative frequency, [15](#)
- homoscedasticity, [57](#)
- hyperparameter, [88](#), [96](#)
- hypothesis, [147](#)
- independent variable, *see* predictor variable
- independently and identically distributed (i.i.d.) random sample, *see* random sample
- indicator variable, [10](#), [14](#), [59](#), [195](#)
- inertia, [112](#)
- inferential statistics, *see* statistical inference
- instance, [8](#), [17](#), [18](#)
- interaction terms, [62](#)
- interquartile range, [20](#)
- inverse survival function, [138](#)

- join, 11
 - inner join, 11
 - left join, 12
 - Pandas merge function, 11
 - outer join, 12
- key, 11
- knee, *see* scree plot, elbow
- label, 87
- Law of Large Numbers, 125
- legislation
 - EU Equal Treatment Directive, 194
 - UK Equality Act, 2010, 194
- likelihood, 164, 183
- linear regression, 183
- linear regression model, 51, 167
- literate programming, 203
- loadings, 72
- local minima, 112
- log odds, 178, 178
- logistic function, 177
- logit function, 178
- long form data, 8
- longitudinal study, *see* cohort study
- loss function, *see* error function, 183
- lower quartile, *see* quartile
- lower-tailed test, 149
- lurking variable, 62
- machine learning, 5
- maximum likelihood principle, *see* principle of maximum likelihood
- maximum likelihood estimate, 184
- maximum likelihood estimator, 188
- mean squared error, 57, 102
 - within-cluster, *see* within-cluster mean squared error
- mean squared error of an estimator, 129
- measurement bias, *see* bias
- median, 32
- metadata, 8, 10
- metric, 102
- missing data, 10
 - as a source of bias, 23
- mode, 20
- model, 51, 127, 147
 - probabilistic, 121
 - statistical, 52, 121
- moment matrix, 68
- multimodal histogram, 16
- multiple regression, 59
- multivariate, 9
- non-parametric method, 51
- normal equations, 53, 68
- normal matrix, 68
- Not a Number, 11
- Not Applicable, 11
- null hypothesis, 147
- numeric data, 20
- numeric variable, 17, 18
- observation, 8, 17, 18
- observational data, 39
- odds, 175, 178
- odds ratio, 176, 178
- one-hot encoding, *see* indicator variable
- one-tailed test, 148
- optimisation, 52, 184
- outcome, 38, 39, *see* response variable
- outlier, 18, 44
 - in box plot, 32
- over-fitting, 62, 96
- over-generalisation, 96
- p -hacking, *see* data dredging
- p -value, 148, 149, 150, 154, 159
 - in output from linear regression, 169
- parameter, 49, 127, 167
- parametric, 175
- parametric model, 51, 59
- partitional clustering, *see* clustering
- partitional clustering, 108
- piecewise linear, 89
- point estimate, 167
- point estimation, 117
- point estimator, 127
- population, 16, 117, 127
- population mean, 17, 18
- population median, 18
- population standard deviation, 19
- population variance, 18
- posterior probability, 164
- practical significance, 159
- preattentive attributes, 28
- predicted values, 55
- predictor, 49
- predictor variable, *see* predictor
- principal components, 74
- principal components analysis, 71
- principle of least squares, 52, 167, 183
 - probabilistic motivation for, 188
- principle of maximum likelihood, 184
- prior probability, 164
- probability density function, 31
- prospective study, 39
- protected characteristics, 194
- prototype, 108
- proxy variables, 196
- pseudorandom numbers, 120
- quantiles, 20
- quartile, 20, 32
 - lower, 20
 - upper, 20
- random number generator, 120

- random sample, 120
- randomised control trial, 38, 39, 155
- receiver operator characteristic, 196
- regression
 - non-parametric, 99
- regression coefficient, 49, 60, 167
- regressor, *see* predictor variable
- regressor matrix, 68
- regular expressions, 13
- regularisation, 96
- regularisation parameter, 96
- rejection region, 148
- relative frequency, 122, 175
- relative frequency density, 16, 31
- residual, 55, 188
- response variable, 49
- retrospective study, 39
- right to be forgotten, *see* right to erasure
- right to erasure, 23
- root mean squared error, 57
- rotation matrix, 82
- sample, 16, 117, 127
 - of convenience, 121
 - random, 119
 - with replacement, 120
 - without replacement, 120
 - stratified, 121
- sample correlation coefficient, 42
- sample covariance, 41
- sample mean, 17, 17
- sample median, 17
- sample variance, *see* variance
- sampling
 - without replacement, 184
- sampling distribution, 119
- scatter, 112
- scree plot, 77, 113
 - elbow, 77, 113
- selection bias, *see* bias
- selectivity, 99, 196
- sensitivity, 99, 196
- skewed distribution, 124
 - negative or left, 18
 - positive or right, 18, 119, 132
- specificity, *see* selectivity
- split-apply-combine, 12
- spurious correlations, 45
- standard deviation
 - sample standard deviation, 18
- standard error of an estimator, 131
- standard error of the mean, 123, 125, 129, 131, 135, 140
- standardised variable, 19, 54, 72, 125
- statistic, 121, 138, 167
- statistical approaches
 - Bayesian, 126
 - Frequentist, 126
- statistical inference, 117
- statistical significance, 151, 159
- statistical simulation, 119, 121, 155
- statistical theory, 119
- statistics, 5
- structured data, *see* data
- study design, 38
- sum of squared errors
 - in regression, 58, 188
 - within-cluster, *see* within-cluster sum of squared errors
- supervised learning, 87, 88, 105
- symmetric distribution, 18
- t* critical value, 144
- t*-distribution, 144
 - in output from linear regression, 169
- tabular data, 8
- target variable, *see* response variable
- test point, 87
- test procedure, 147
- test set error rate, 93
- test statistic, 147
- testing set, 93
- tidy data, 8
- total sum of squares, 58
- training set, 87, 93, 93
- training set error rate, 93
- transparent, 180
- treatment, 39
- treatment group, 38
- two-tailed test, 149
- Type I error, 153
- Type II error, 154
- unbalanced classes, 98
- unbiased estimator, *see* estimator, unbiased
- under-fitting, 96
- under-generalisation, 96
- unimodal histogram, 16
- unit vector, 81
- univariate, 9
- univariate data, 15
- unstructured data, *see* data
- unsupervised learning, 87, 105
- upper quartile, *see* quartile
- upper-tailed, 152
- validation data, 98
- variable, 8
 - categorical, 9
 - numerical, 9
 - ordinal, 9
 - string, 9
- variables
 - associated, 44
 - causally related, 44
 - correlated, 44
- variance

- of an estimator, [129](#), [129](#)
 - sample variance, [18](#)
- vector quantisation, [108](#)
- visualisation, [25](#)
- Voronoi diagram, *see* Voronoi tessellation
- Voronoi tessellation, [89](#)

- web scraping, [14](#), [41](#)
- weights, [72](#)
- within-cluster mean squared error, [112](#)
- within-cluster sum of squared errors, [112](#)

- z critical value, [135](#), [159](#), [162](#)
- z-distribution, [135](#), [159](#)
- z-score, *see* standardised variable