

Mathematics Level 9

Chapter 6 — Statistics · 5 subtopics

This work is licensed under the Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License. To view a copy of this license, visit <http://creativecommons.org/licenses/by-nc-nd/4.0/> or send a letter to Creative Commons, PO Box 1866, Mountain View, CA 94042, USA.

Attribution: Classbasics Pty Ltd, vwx.au

Aboriginal and Torres Strait Islander content has been excluded as accuracy cannot be guaranteed, and it is better to be respectful than cover the entire curriculum inaccurately.

Significant effort has been made to ensure this content is legal. Please send any areas of concern to admin@vwx.au with appropriate document/legal references so errors can be verified and changes be made.

Contents

Section	Page
Comparing distributions: back-to-back stem-and-leaf plots, histograms, centre, spread, shape and outliers	2
Choosing and justifying a data display	24
Reading survey reports: how the data was obtained, and estimating population means and medians	46
Sampling methods, sample-to-sample variation, and representation that promotes a point of view	61
Planning and conducting a statistical investigation	79
Test	99

Comparing distributions: back-to-back stem-and-leaf plots, histograms, centre, spread, shape and outliers

6A · Chapter 6 — Statistics · Level 9

Content description VC2M9ST03 (word for word): *represent the distribution of multiple data sets for numerical variables using comparative representations such as back-to-back stem-and-leaf plots and histograms; describe data, using terms including ‘skewed’, ‘symmetric’ and ‘bi-modal’; compare data distributions using mean, median and range to describe and interpret numerical data sets with consideration of centre, spread and shape, and the effect of outliers on these measures*

Achievement standard anchor (AS 16): *Students compare and analyse the distributions of multiple numerical data sets, choose representations, describe features of these data sets using summary statistics and the shape of distributions, and consider the effect of outliers.*

Elaboration ids for linkage: VC2M9ST03.E01, VC2M9ST03.E02, VC2M9ST03.E03, VC2M9ST03.E05

Prerequisite pointer: Level 8 Statistics is assumed — population vs sample and collection techniques (VC2M8ST01), analysing and reporting on the distribution of data from primary and secondary sources (VC2M8ST02), sample variation and the effect of sample size (VC2M8ST03), and investigations with samples and fair inference (VC2M8ST04). The stem-and-leaf plot, the histogram, and the mean, median and range are rehearsed below as tools, not re-taught. Level 8 described one distribution; 6A compares several — centre against centre, spread against spread, shape against shape, and what one unusual value does to the story.

Note on VC2M9ST03.E04 (exclusion recorded): E04’s ‘Closing the Gap’ material is excluded from this book under CONTENT_POLICY Rule 1 (D10; recorded as ACCEPTED here, since TOC.md must not be edited from inside a section). The content description is delivered in full via VC2M9ST03.E01–E03 and E05 contexts — comparing real and scenario distributions by centre, spread and shape — per D5: elaborations are illustrations of a CD, not the CD itself.

Note on sources: the weather contexts below use real data — Source: Australian Bureau of Meteorology, “Daily Weather Observations for Melbourne (Olympic Park), Victoria, July 2026” (station 086338), prepared 14 August 2026. Retrieved 20 August 2026; and Source: Australian Bureau of Meteorology, “Daily Weather Observations for Ballarat, Victoria, July 2026” (station 089002), prepared 14 August 2026. Retrieved 20 August 2026. (The CSVs are archived in `_build9/_audit/bom_dwo_202607/`.) Every other context is a clearly-labelled SCENARIO or a constructed sample drawn to show a shape (R2, R11).

Theory

Where Level 8 leaves off

At Level 8 you described single data sets: you drew dot plots, stem-and-leaf plots and histograms, and you summarised a set with its **mean**, its **median** and its **range**. Level 9 keeps those tools and changes the job. The content description’s verb is *compare*: put two or more **distributions** side by side and say, in full sentences, how their centre, spread and shape differ — and what happens to the measures when one value sits well away from the rest.

This subtopic adds three things to your toolkit:

1. a comparative display — the **back-to-back stem-and-leaf plot**, which keeps every value of two sets visible on one shared scale;
2. the **histogram** used comparatively — two large sets drawn with the same classes on the same axis;
3. a comparison routine — mean, median and range for centre and spread, the shape words **symmetric**, **skewed** and **bi-modal**, and a rule for the **outlier**.

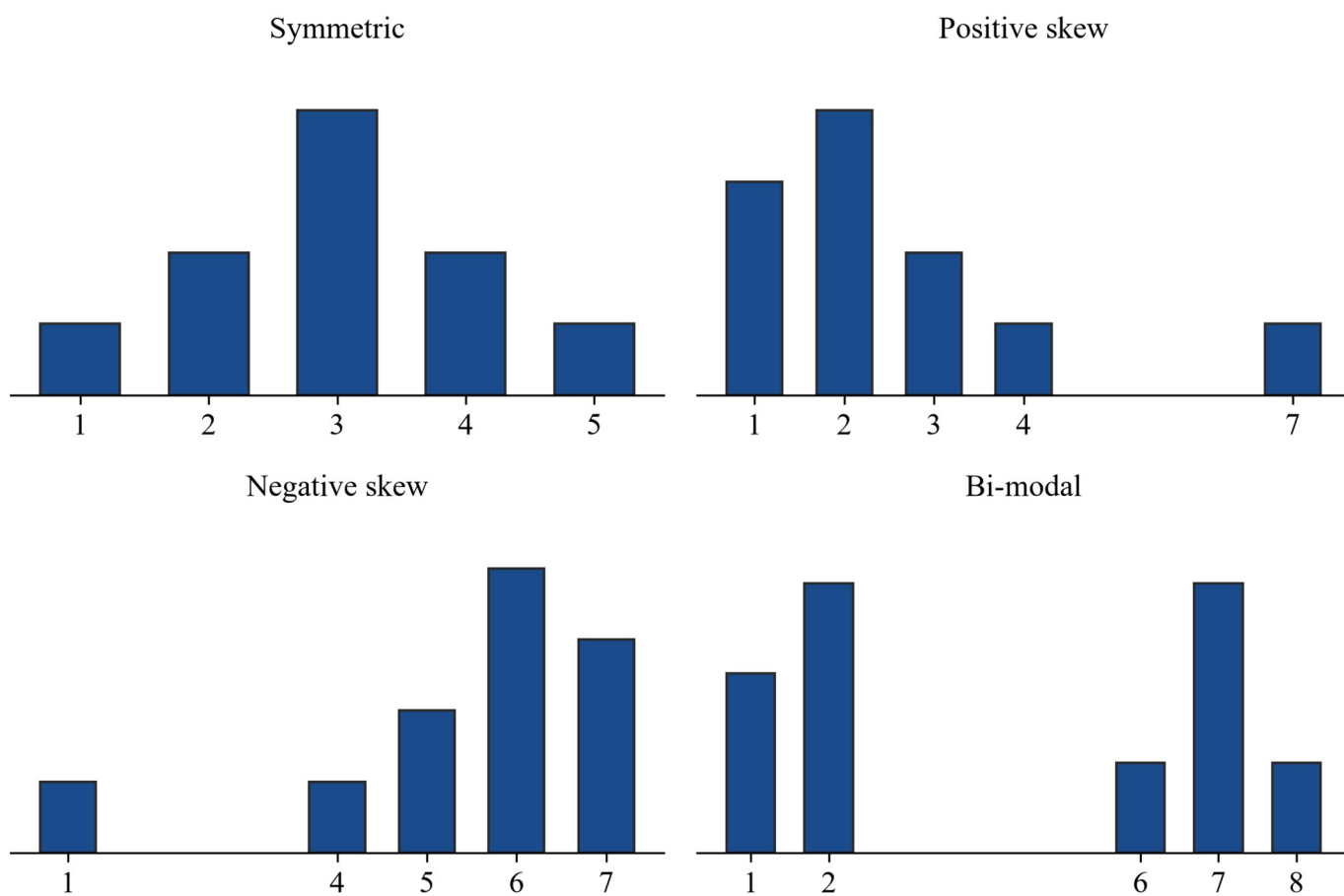
Centre, spread and shape — the comparison tools

The three measures are Level 8 tools, restated here for comparing:

- The **mean** is the sum of the values divided by how many there are. It uses every value — which is its strength, and its weakness, because one unusual value pulls it.
- The **median** is the middle value once the data is ordered, or the mean of the middle two values when the count is even. One unusual value moves it very little.
- The **range** is the largest value minus the smallest. It is a one-number summary of spread, and one unusual value at either end changes it completely.

At Level 9, shape has examinable names:

- **symmetric** — the two halves of the display are close to mirror images, and the mean and median sit close together;
- **positive skew** — a tail of high values runs off to the right; the mean is pulled *above* the median;
- **negative skew** — a tail of low values runs off to the left; the mean is pulled *below* the median;
- **bi-modal** — the display has two peaks, usually because two different groups are mixed into one data set; a mean or median sitting in the valley between the peaks can describe almost nobody.



Constructed samples drawn to show shape only — the values are not real data.

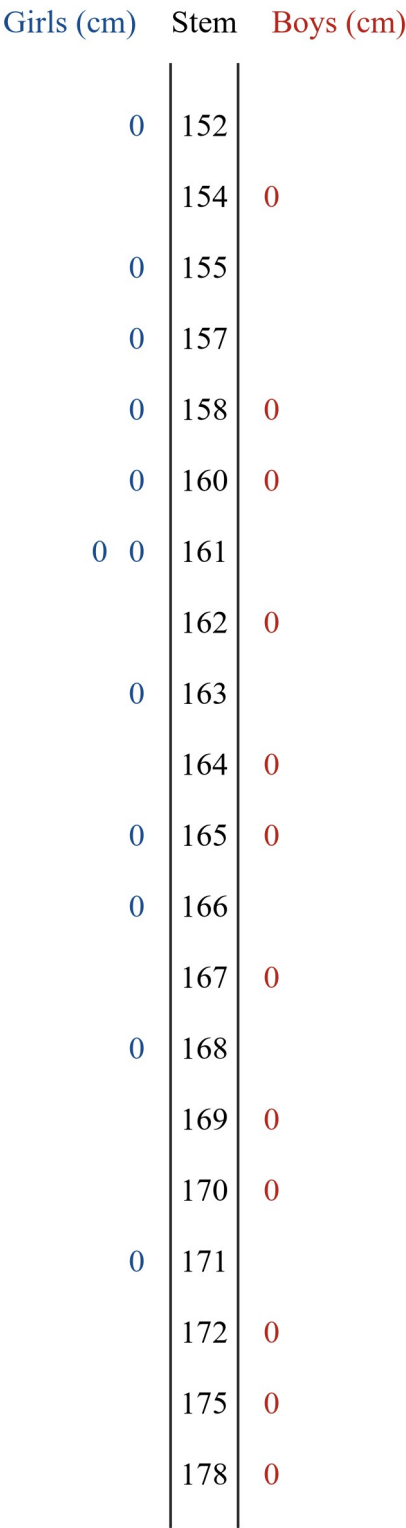
Four constructed samples drawn as histograms: a symmetric shape, a positive skew with its tail to the right, a negative skew with its tail to the left, and a bi-modal shape with two peaks

The figure shows the four shapes as constructed samples — the values exist only to show the shapes. Read the skew from the **tail**, not from the tall bars: the positive-skew panel has its tall bars at the left and its long tail at the right.

Back-to-back stem-and-leaf plots

A stem-and-leaf plot keeps every value visible: the stem carries the leading digits and each leaf is one value's final digit. A **back-to-back** plot puts two data sets on one diagram — the sets share the stem column, and their leaves run outwards on opposite sides. By convention the smallest leaf on each side sits next to the stem, and a **key** fixes the place value. Because both sets sit on the same stems, differences in centre and spread are visible at a glance.

Heights of 12 girls and 12 boys in a Year 9 class (SCENARIO)



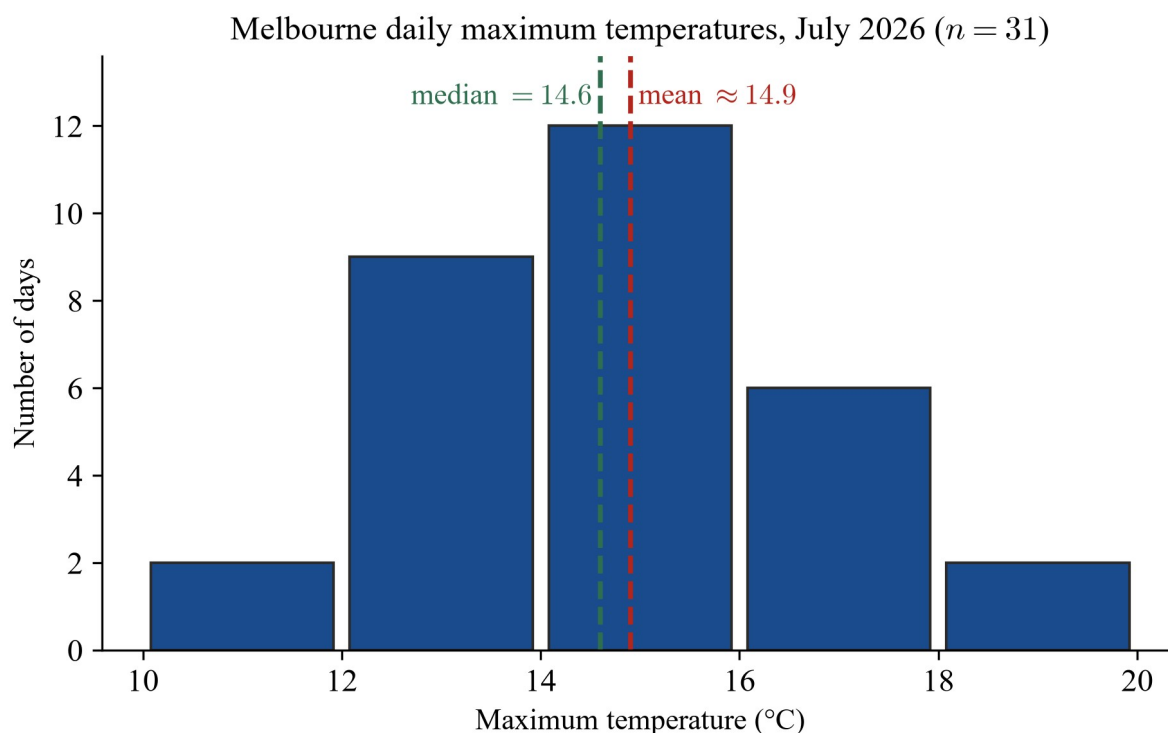
Key: 7 | 16 | 0 means 167 cm for a girl and 160 cm for a boy.

Back-to-back stem-and-leaf plot of the heights in centimetres of 12 girls and 12 boys in a Year 9 class, with the girls' leaves on the left and the boys' leaves on the right

In this plot (a SCENARIO), the boys' leaves sit on higher stems and stretch over more rows: the boys are typically taller and more spread out. Nothing is lost in the picture — all 24 heights can still be read off, leaf by leaf, which is the point of choosing this display for two smaller sets.

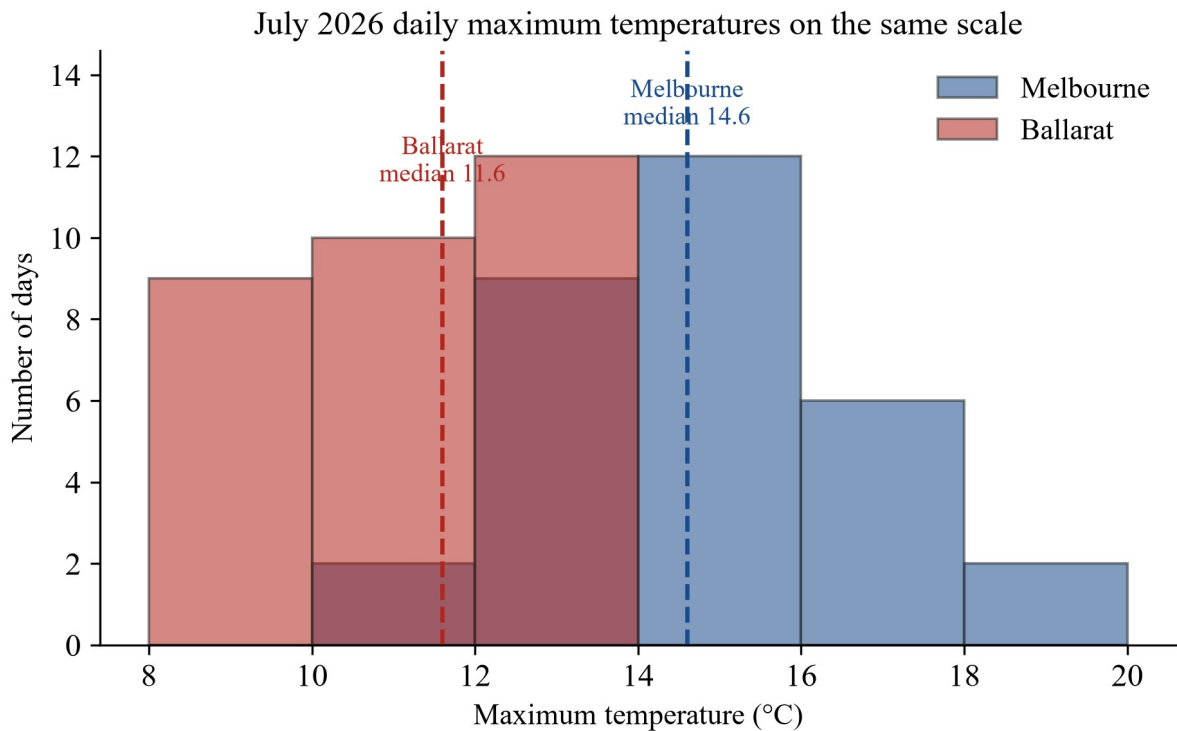
Histograms, and comparing two on one scale

When a set is large, listing every value gets crowded, so the values are grouped into equal-width **classes** and a **histogram** is drawn: the height of each bar is the **frequency** (the number of values in that class), and the bars touch, because the horizontal axis is a number line. (That is what separates a histogram from a bar chart: a bar chart shows categories, with gaps.)



Histogram of the 31 daily maximum temperatures at Melbourne (Olympic Park) in July 2026, in two-degree classes, with the mean of about 14.9 degrees Celsius and the median of 14.6 degrees Celsius marked

The histogram of Melbourne's July 2026 daily maxima (real data — source line above) shows the **modal class**, the class with the tallest bar, at 14–16 °C, and a slight tail towards the warm days. To compare two large sets, draw both histograms with the same classes on the same axis; then centre, spread and shape can be read straight off the picture.

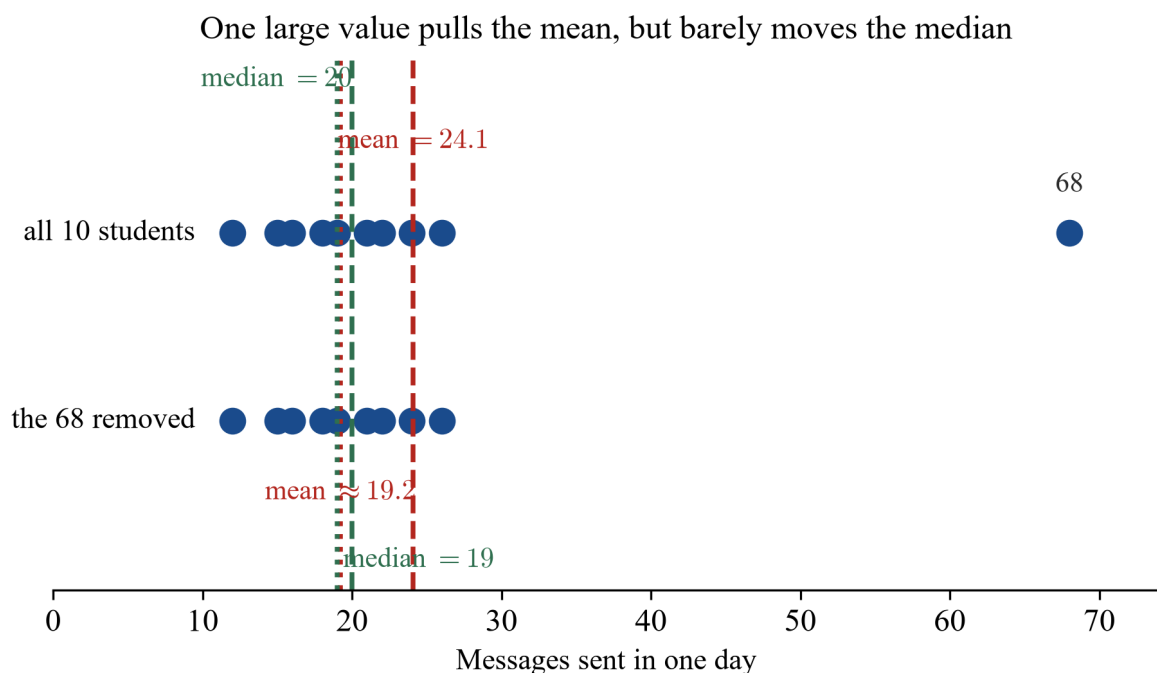


Two histograms on the same axis and the same two-degree classes: Melbourne's July 2026 daily maximum temperatures in blue and Ballarat's in red, with the median of each city marked

On one scale the story is immediate: Ballarat's whole distribution sits about three degrees to the left of Melbourne's and inside a narrower band, while both months keep a single peak. Same shape family, different centre and spread — exactly the comparison the content description asks for.

Outliers: what one value can do

An **outlier** is a value that sits well away from the bulk of the data. The curriculum uses the word without defining it, so this is the definition this book works with; at Level 9 the judgement is visual — look at the display and ask whether one value sits apart from the crowd. Its importance is mechanical: the mean uses every value, so an outlier pulls it; the median only cares about order, so it moves very little; the range stretches to include it.



Dot plots of the messages sent in one day by ten students, drawn twice: once with all ten values, showing mean 24.1 and median 20, and once with the outlier of 68 removed, showing mean about 19.2 and median 19

In the SCENARIO of the figure, one student's 68 messages pulls the mean from about 19.2 up to 24.1 and stretches the range from 14 to 56, while the median moves from 19 to 20. So a fair description names the outlier and says what it does — and where a typical value is wanted, the median is the honest number.

The comparison checklist

This checklist is the working tool of the subtopic. Every worked example and every question below runs some part of it.

1. **Read the display** — key, stems and leaves, classes, scales. Are both sets drawn on the same scale?
 2. **Centre** — the mean and the median of each set. Which set is typically higher, and do the mean and median agree?
 3. **Spread** — the range of each set. Which set is more variable, and which more consistent?
 4. **Shape** — symmetric, skewed (and which way), or bi-modal?
 5. **Outliers** — does any value sit apart from the bulk, and which measures would it pull?
 6. **Interpret in context** — write full sentences that compare: “typically higher”, “more variable”, “a similar shape”.
-

Worked examples

Example 1 — scaffolded

SCENARIO. The heights, in centimetres, of the 12 girls and 12 boys in a Year 9 class are recorded. Girls: 152, 155, 157, 158, 160, 161, 161, 163, 165, 166, 168, 171. Boys: 154, 158, 160, 162, 164, 165, 167, 169, 170, 172, 175, 178. Build a back-to-back stem-and-leaf plot and compare the two distributions.

Heights of 12 girls and 12 boys in a Year 9 class (SCENARIO)

Girls (cm) Stem Boys (cm)

0	152	
	154	0
0	155	
0	157	
0	158	0
0	160	0
0 0	161	
	162	0
0	163	
	164	0
0	165	0
0	166	
	167	0
0	168	
	169	0
	170	0
0	171	
	172	0
	175	0
	178	0

Key: 7 | 16 | 0 means 167 cm for a girl and 160 cm for a boy.

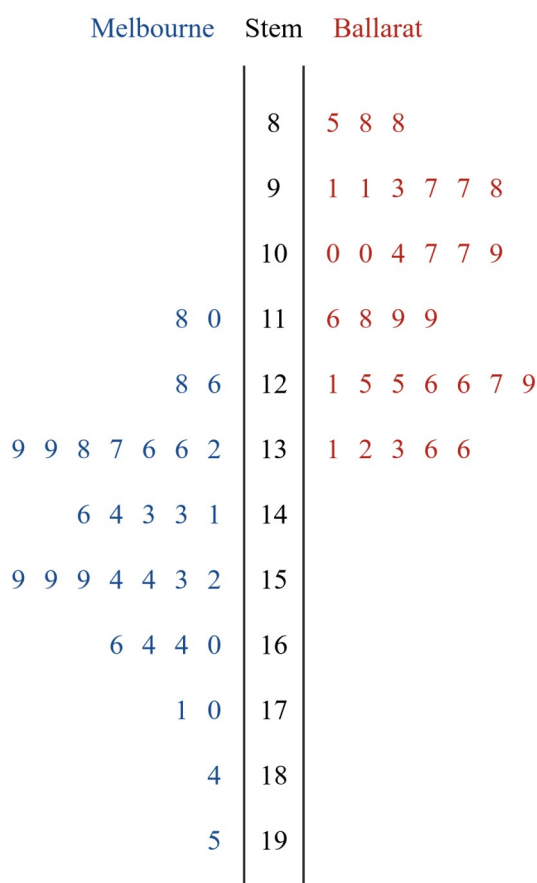
Back-to-back stem-and-leaf plot of the girls' and boys' heights, stems 15 to 17

Working	Comment
Stems 15, 16, 17; each leaf is one height; key: 7 / 16 / 0 means 167 cm for a girl and 160 cm for a boy	One shared stem column, leaves running out each side, smallest leaf next to the stem.
Girls: mean $\frac{1937}{12} \approx 161.4$ cm; median 161 cm; range 171 – 152 = 19 cm	The middle two of the 12 ordered girls' heights are 161 and 161.
Boys: mean $\frac{1994}{12} \approx 166.2$ cm; median $\frac{165 + 167}{2} = 166$ cm; range 178 – 154 = 24 cm	Every summary sits higher for the boys.
Shape: one block of stems on each side, no gaps, no leaf sitting apart	Neither set shows an outlier.
$\therefore$ The boys are typically about 5 cm taller (medians 166 against 161) and somewhat more spread (ranges 24 against 19); both distributions have one peak and no outliers	Step 6 of the checklist: the comparison sentences are the answer, not the numbers alone.

Example 2 — scaffolded

The back-to-back stem-and-leaf plot below shows the daily maximum temperature, in °C, recorded at Melbourne (Olympic Park) and at Ballarat through July 2026 — 31 days at each station. Compare the two distributions. Source: Australian Bureau of Meteorology, “Daily Weather Observations for Melbourne (Olympic Park), Victoria, July 2026” (station 086338), prepared 14 August 2026. Retrieved 20 August 2026; and the matching Ballarat (station 089002) file.

Daily maximum temperature, July 2026: Melbourne (left) and Ballarat (right)



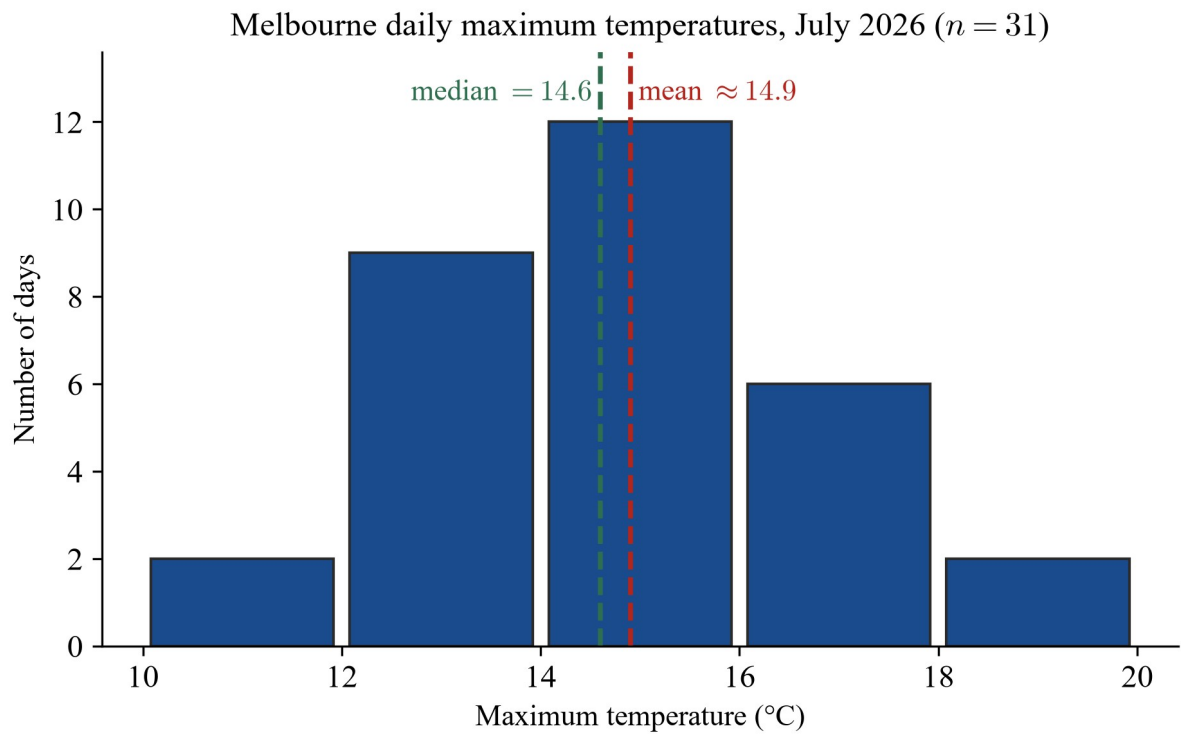
Key: 3 | 14 | 1 means 14.3 °C at Melbourne and 14.1 °C at Ballarat. Leaves are tenths of a degree. Source: Bureau of Meteorology.

Back-to-back stem-and-leaf plot of July 2026 daily maximum temperatures, Melbourne on the left and Ballarat on the right, stems 8 to 19 with leaves as tenths of a degree

Working	Comment
Melbourne: mean ≈ 14.9 °C; median 14.6 °C (the 16th of 31 leaves); range $19.5 - 11.0 = 8.5$ °C	Step 1: the key says 3 / 14 / 1 means 14.3 °C at Melbourne and 14.1 °C at Ballarat.
Ballarat: mean ≈ 11.2 °C; median 11.6 °C; range $13.6 - 8.5 = 5.1$ °C	Same routine on the right-hand side.
Centre: $14.9 - 11.2 = 3.7$ °C by the means, $14.6 - 11.6 = 3.0$ °C by the medians	Melbourne was typically about three degrees warmer.
Spread: 8.5 °C against 5.1 °C	Melbourne's month varied more; Ballarat stayed inside a five-degree band.
Shape: Melbourne's tail runs to the warm side (leaves at 17, 18, 19) — slight positive skew; Ballarat is one block, roughly symmetric	Mean above median at Melbourne, just below at Ballarat, agreeing with the shapes.
$\therefore$ July 2026 was typically about three degrees warmer at Melbourne than at Ballarat and clearly more variable, with a slight positive skew at Melbourne; neither month shows an outlier, so mean and median tell the same story at each station	A full comparison: centre, spread and shape, in sentences.

Example 3 — scaffolded

The histogram shows the 31 daily maximum temperatures recorded at Melbourne (Olympic Park) in July 2026, in 2-degree classes. Describe the distribution: size, modal class, centre and shape. Source: Australian Bureau of Meteorology, “Daily Weather Observations for Melbourne (Olympic Park), Victoria, July 2026” (station 086338), prepared 14 August 2026. Retrieved 20 August 2026.

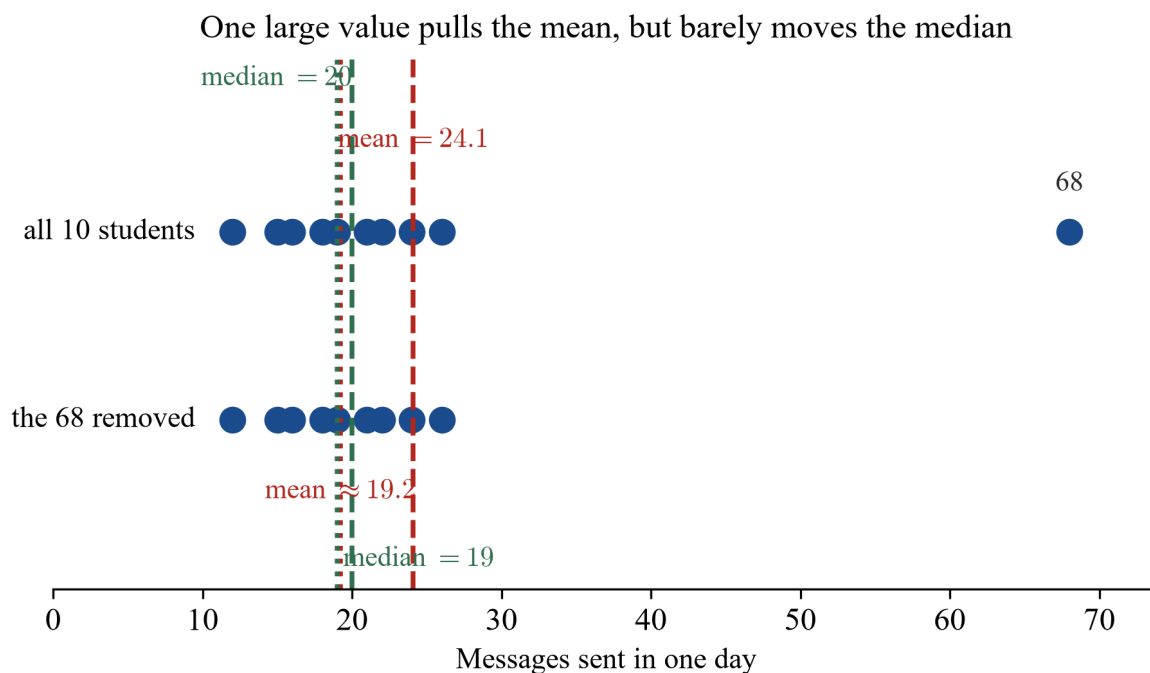


Histogram of the 31 Melbourne daily maxima in two-degree classes with the mean and median marked

Working	Comment
$n=31$ days; classes 10–12, 12–14, 14–16, 16–18, 18–20 °C	A value on a class boundary is counted in the higher class.
Modal class 14–16 °C, holding 12 of the 31 days: $\frac{12}{31} \approx 38.7\%$	The tallest bar.
Mean ≈ 14.9 °C; median = 14.6 °C	Both are marked on the figure.
Mean just above the median, with a small tail of warm days at the right → slight positive skew	The two days at 18–20 °C form the tail.
∴ Melbourne’s July 2026 maxima clustered in the 14–16 °C class, sat between 11.0 °C and 19.5 °C (range 8.5 °C), and were slightly positively skewed	One sentence each for centre, spread and shape.

Example 4 — scaffolded

SCENARIO. Ten students count the messages they send in one day: 12, 15, 16, 18, 19, 21, 22, 24, 26, 68. Investigate the effect of the value 68 on the mean, the median and the range.



Dot plots of the message counts with and without the outlier of 68, with the mean and median of each version marked

Working

Comment

With all ten: mean $\frac{241}{10} = 24.1$; median $\frac{19+21}{2} = 20$;
range $68 - 12 = 56$

On the dot plot, 68 sits far from the crowd — an outlier, judged by eye.

Without the 68: mean $\frac{173}{9} \approx 19.2$; median 19; range
 $26 - 12 = 14$

The other nine values are untouched.

Mean falls by $24.1 - 19.2 = 4.9$; median falls by 1;
range falls by $56 - 14 = 42$

The mean and the range feel the outlier; the median shrugs it off.

$\therefore$ The single value 68 pulls the mean up by about five messages and stretches the range from 14 to 56, while the median moves by one — so the median describes the typical student here, and a fair report names the outlier instead of hiding it

This is the “effect of outliers on these measures” of the content description.

Example 5 — unscaffolded

SCENARIO. Two swimmers each race ten times over 50 m. Times in seconds — Swimmer A: 36.2, 36.5, 36.5, 36.8, 37.0, 37.1, 37.4, 37.6, 37.9, 38.3. Swimmer B: 35.1, 35.9, 36.4, 36.8, 37.2, 37.5, 37.8, 38.6, 39.4, 40.7. Compare the two distributions and say which swimmer a relay selector might prefer.

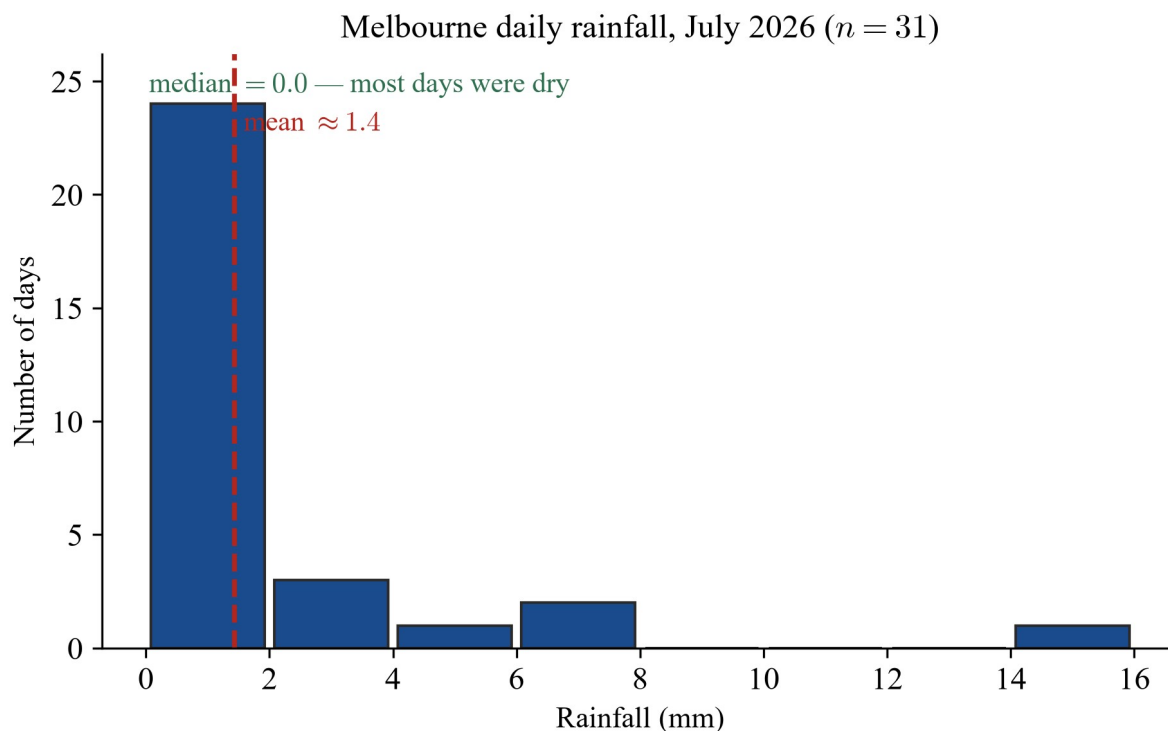
Swimmer A: mean $\frac{371.3}{10} = 37.13$ s; median $\frac{37.0+37.1}{2} = 37.05$ s; range $38.3 - 36.2 = 2.1$ s. Swimmer B: mean

$\frac{375.4}{10} = 37.54$ s; median $\frac{37.2+37.5}{2} = 37.35$ s; range $40.7 - 35.1 = 5.6$ s. Centre: A is faster on both measures —

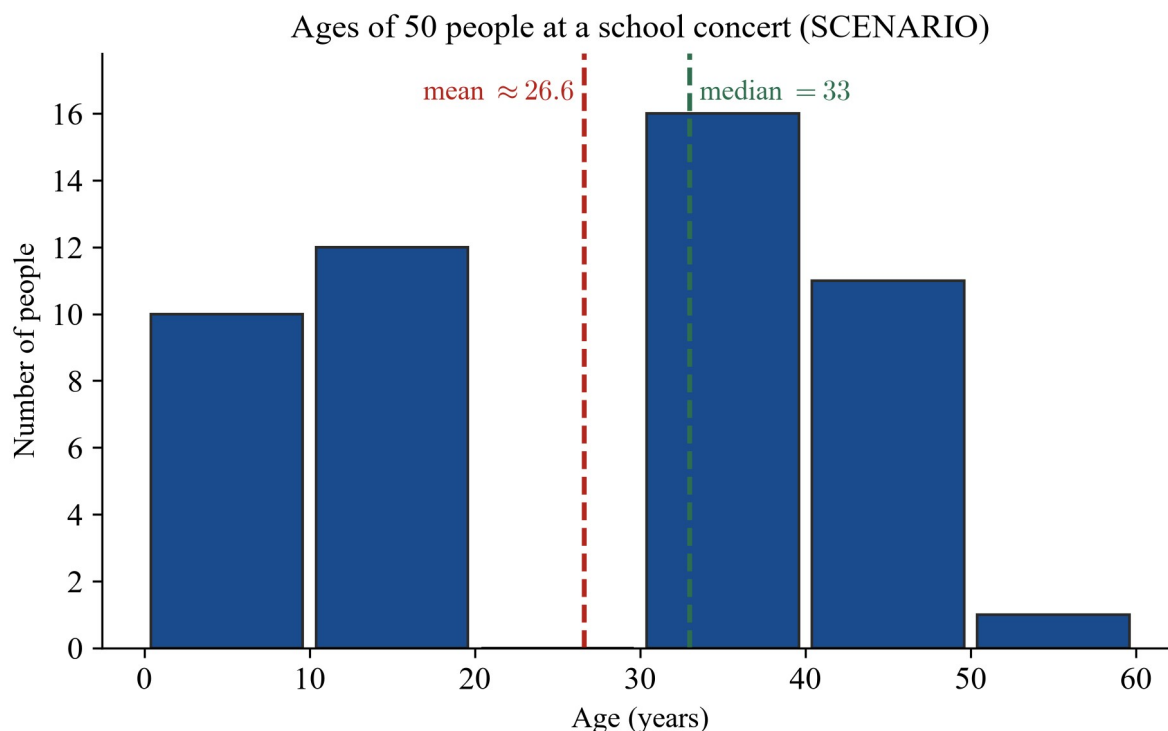
by 0.41 s on the means and 0.30 s on the medians. Spread: A’s times sit inside 2.1 s while B’s span 5.6 s — A is far more consistent. Shape: A is roughly symmetric; B carries a tail of slow times (38.6, 39.4, 40.7), so B’s set is positively skewed, and that tail drags B’s mean above B’s median. $\therefore$ Swimmer A is typically faster and much more consistent, so a selector who needs a reliable relay leg chooses A — B’s two slow races pull the mean slower and stretch the range, which is exactly the outlier effect working on real data.

Example 6 — unscaffolded

The first histogram below is real data — the daily rainfall, in millimetres, at Melbourne (Olympic Park) in July 2026. The second is a SCENARIO — the ages of the 50 people in the audience at a school concert. Name the shape of each distribution and give the evidence from the display. Source: Australian Bureau of Meteorology, “Daily Weather Observations for Melbourne (Olympic Park), Victoria, July 2026” (station 086338), prepared 14 August 2026. Retrieved 20 August 2026.



Histogram of Melbourne daily rainfall in July 2026 in two-millimetre classes: a very tall first bar of 24 dry or nearly dry days and a long low tail out to the 14.6 millimetre day



Histogram of the ages of 50 people at a school concert: one peak for children at 0 to 20 years, an empty class at 20 to 30, and a second peak for adults at 30 to 60 years

Rainfall: 16 of the 31 days recorded 0.0 mm, the median is 0.0 mm, and the mean is pulled up to about 1.4 mm by a few wet days, one of them 14.6 mm. The tall bar at the left and the long low tail to the right make this a strong **positive skew** — the tail sets the name. Concert ages: two peaks — children at 0–20 and adults at 30–60 — with an empty class between them, so the distribution is **bi-modal**. Its mean is 26.6 years and its median 33 years, and neither describes a typical person in the room: almost nobody at the concert was 27. $\therefore$ Shape is read from the display — a tail to the right is positive skew, two peaks is bi-modal — and in a bi-modal set the mean can be the least useful number of all.

Example 7 — unscaffolded

SCENARIO. The hourly pay, in dollars, of the eleven staff of a café, the owner included, is: 24, 24, 25, 25, 26, 26, 27, 28, 31, 33, 92. A local paper writes “average pay at the café is about \$33 an hour”. Analyse the choice of number.

The sum is $24 + 24 + 25 + 25 + 26 + 26 + 27 + 28 + 31 + 33 + 92 = 361$. The mean is $\frac{361}{11} \approx 32.82$ — the paper’s “\$33” is the mean, rounded. The median is the 6th of the 11 ordered values: 26. The range is $92 - 24 = 68$. Ten of the eleven staff earn between \$24 and \$33; the owner’s \$92 is an outlier, and it pulls the mean almost \$7 above the median. $\therefore$ The mean serves the “generous employer” story; the median, \$26, describes what a typical staff member actually earns. Quoting the mean here is a choice of representation, not a fact about the workers — the reader deserves both numbers.

Example 8 — unscaffolded

A radio presenter says: “Melbourne had a warm July in 2026 — more than half the days reached 15 °C or above.” Check the claim against the histogram of the 31 daily maxima. Source: Australian Bureau of Meteorology, “Daily Weather Observations for Melbourne (Olympic Park), Victoria, July 2026” (station 086338), prepared 14 August 2026. Retrieved 20 August 2026.

Count the days at 15.0 °C or above from the data behind the histogram: 15 of the 31 days. Half of 31 is 15.5, so “more than half” needs at least 16 days. 15 falls one day short; as a percentage, $\frac{15}{31} \approx 48.4\%$ — under half. $\therefore$ The claim fails by a single day: July 2026 was cool for the time of year, and 48.4% is not “more than half” — a near miss that shows why claims are settled by counts, not by impressions.

Practice questions

Scaffolded

Q1. The back-to-back stem-and-leaf plot of Example 1 shows the heights of 12 girls and 12 boys.

- How many heights are recorded on each side?
- Read off the tallest girl and the tallest boy.
- How many students in all are shorter than 160 cm?

Q2. A quiz out of 50 is marked for ten students, and the scores are displayed in a stem-and-leaf plot.

Stem	Leaves
2	1 4 7
3	0 2 2 5 8
4	1 6

Key: 3 / 2 means 32.

- Write the ten scores in order.
- Find the median score.
- Find the mode and the range.

Q3. The ages, in years, of eleven volunteers at a park working bee are: 18, 22, 25, 25, 31, 34, 34, 34, 39, 40, 42.

- Calculate the mean age to one decimal place.
- Find the median age.
- Find the range.
- Which measure, mean or median, better describes a typical volunteer? Give your reason.

Q4. Two teams' scores over seven games are shown back-to-back.

Team A	Stem	Team B
	0	9
8 8 5 2	1	4 6
7 4 1	2	0 2 5
	3	0

Key: 2 / 1 / 4 means 12 for Team A and 14 for Team B.

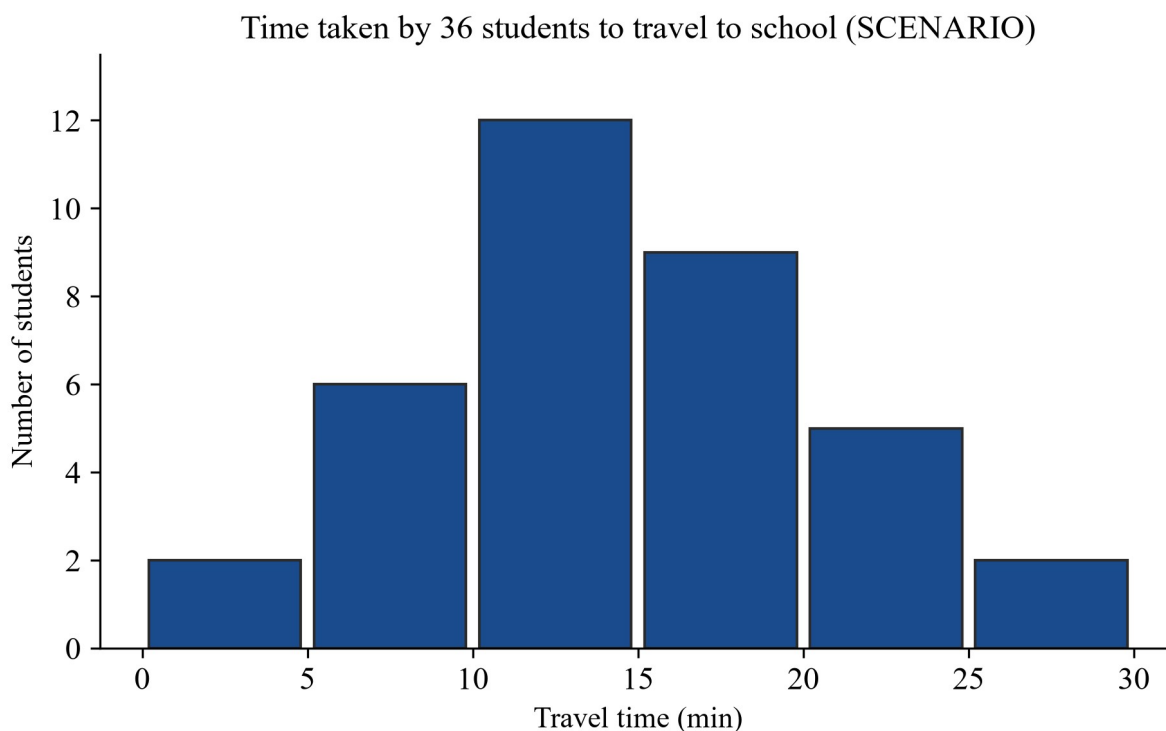
- List each team's seven scores.
- Find the median score of each team.
- Find the range of each team, and say which team is more consistent.

Q5. Find the mean, median and range of: 4, 7, 7, 9, 10, 12, 14, 17.

Q6. A set of seven values is: 7, 8, 8, 9, 10, 11, 45.

- Find the median and the range.
- Calculate the mean.
- Which of the three measures does the 45 distort most? Which does it distort least? Explain.

Q7. The histogram shows the time taken by 36 students to travel to school, in 5-minute classes (SCENARIO).



Histogram of the travel times of 36 students in five-minute classes, peaking at 10 to 15 minutes

- Check the total number of students by adding the frequencies.
- Write down the modal class.
- How many students take at least 15 minutes?

Q8. Use the same histogram.

- What percentage of the students take less than 20 minutes? Give your answer to one decimal place.
- A journalist writes “almost all students take less than 20 minutes to get to school”. Evaluate the sentence.

Q9. Name the shape (symmetric, positive skew, negative skew or bi-modal) of each described distribution.

- The two halves of the display are close to mirror images.
- A tail of high values runs off to the right, and the mean sits above the median.
- Two peaks appear, with a valley between them.
- A tail of low values runs off to the left, and the mean sits below the median.

Q10. The set 3, 5, 6, 6, 8, 9, 30 contains an outlier. The 30 is later corrected to 7.

- Find the mean, median and range before the correction.
- Find the mean, median and range after the correction.
- Which measure changed most, and which not at all?

Q11. Use the histogram of Melbourne’s July 2026 daily maxima (Example 3).

- Write down the modal class and the number of days in it.
- How many days had a maximum below 12 °C?
- How many days reached at least 16 °C?
- The mean (about 14.9 °C) sits above the median (14.6 °C). Which shape word does this suggest?

Q12. Two data sets are summarised without their values: set P has mean 12, median 12 and range 4; set Q has mean 12, median 12 and range 20.

- Which set is more spread out?
- In which set is the mean the more trustworthy picture of a typical value? Give your reason.
- Write one sentence comparing the two sets.

Un scaffolded

Q13. Mia’s goals over ten netball games: 8, 11, 12, 14, 15, 15, 17, 18, 20, 24. Zoe’s goals over ten games: 5, 9, 10, 12, 13, 14, 16, 19, 22, 28. Compare the two distributions using centre, spread and shape.

Q14. Build a back-to-back stem-and-leaf plot for the swimmers’ times of Example 5, using stems 35 to 40 and leaves as tenths of a second. Include a key.

Q15. Use the back-to-back plot of the July 2026 maxima (Example 2).

- Read off the median maximum for each city.
- Calculate the range for each city.
- How many Ballarat days reached at least 12.0 °C, and how many Melbourne days reached at least 15.0 °C?

Q16. Using the same plot, write a three-sentence comparison of the two cities’ July 2026 maxima: one sentence on centre, one on spread, one on shape.

Q17. The histogram shows the screen time, in minutes, of 20 students yesterday (SCENARIO), in 30-minute classes: 30–60 holds 4 students, 60–90 holds 8, 90–120 holds 6, and 120–150 holds 2.

- How many students are represented?
- Write down the modal class.
- How many students used at least 90 minutes?
- Name the shape of the distribution.

Q18. Use the two histograms of the July 2026 maxima drawn on one scale (Theory).

- Which city was typically warmer? Quote one number from each histogram to support your answer.

- b. Which city's month was more variable? Quote the ranges.
- c. Give one difference in shape between the two distributions.

Q19. Nine values are: 6, 9, 10, 11, 12, 13, 14, 15, 71.

- a. Find the mean (to one decimal place), the median and the range.
- b. Remove the 71 and find the mean, median and range of the eight remaining values.
- c. Describe the effect the outlier had on each of the three measures.

Q20. Five values have mean 12. Four of them are 9, 11, 13, 15. Find the fifth value.

Q21. The six values 3, 5, a , 8, 9, 12 have median 6.5.

- a. Find a .
- b. Does the range of the six values depend on a ? Explain.

Q22. Each line describes a distribution through its numbers only. Name the most likely shape.

- a. Mean 24, median 19 — one value is far above the rest.
- b. Mean 7.4, median 9 — one value is far below the rest.
- c. Mean and median both 12, and the display folds onto itself.
- d. The histogram has peaks at 8 and at 31, with almost nothing between.

Q23. Use the histogram of the ages of 50 people at a school concert (Example 6).

- a. Explain why the class 20–30 is empty.
- b. The mean age is 26.6 years. Explain why that number describes almost nobody in the room.
- c. Suggest a better way to summarise the audience's ages in one sentence.

Q24. Write down six whole-number values whose median is 8 and whose range is 10. Check your answer by calculating both summaries.

Q25. Write down two sets of five values that have the same range but different means. Show that both conditions hold.

Q26. Using the plot you built in Question 14, compare the two swimmers in two sentences: one about centre, one about spread.

Q27. Two cyclists record the distance, in kilometres, of 12 training rides each. Rider A: 24, 26, 27, 28, 29, 30, 30, 31, 32, 33, 34, 36. Rider B: 18, 22, 25, 27, 29, 30, 31, 33, 35, 38, 41, 45. Build a back-to-back stem-and-leaf plot of the two sets, with a key.

Q28. Every value of the set 8, 10, 12, 14, 16 is increased by 5.

- a. Find the mean, median and range of the new set.
- b. Which summaries changed, and which stayed the same? Generalise in one sentence.

Q29. Write down seven values with mean 10, median 9 and range 8. Check your answer by calculation.

Q30. A data set of seven distances, in kilometres, is 14, 18, 22, 25, 32, 36, 41. One value is recorded wrongly: 32 is written as 39.

- a. Find the mean, median and range of the true set.
- b. Find the mean, median and range of the set as recorded.
- c. Which measures would alert you that something was misrecorded, and which would hide the mistake?

Q31. Use the histogram of Melbourne's daily rainfall for July 2026 (Example 6).

- a. Write down the median daily rainfall and say, in words, what it tells you about the month.
- b. The mean (about 1.4 mm) sits well above the median. Explain why.
- c. Name the shape of the distribution.

Q32. Two classes sit the same test. Class A's ten scores: 55, 60, 62, 65, 68, 70, 72, 75, 78, 80. Class B's ten scores: 42, 48, 58, 62, 66, 70, 74, 85, 90, 95. Compare the two distributions using centre, spread and shape.

Q33. Three statements are made about the histogram of Melbourne's July 2026 maxima. Judge each as true or false, with a reason.

- a. "The mean maximum is above the median maximum, so the month is slightly positively skewed."
- b. "The modal class is 16–18 °C."
- c. "The range of the maxima is 8.5 °C."

Q34. A forecaster claims that "at least ten days of July 2026 stayed below 14 °C at Melbourne". Check the claim against the data behind the histogram.

Q35. Write down five values whose median is greater than their mean. Check by calculation, and name the shape such a set would show.

Q36. Use your plot from Question 27 to compare the two cyclists' training rides. Your comparison must include:

- the median and the mean distance for each rider;
 - the range for each rider;
 - the shape of each distribution, with the evidence;
 - whether any ride sits apart from the bulk of the data, and what it does to the measures;
 - a final sentence interpreting the comparison in context.
-

Answers

Q1 a 12 on each side; b tallest girl 171 cm, tallest boy 178 cm; c 6 students (4 girls, 2 boys). **Q2** a 21, 24, 27, 30, 32, 32, 35, 38, 41, 46; b 32; c mode 32, range 25. **Q3** a 31.3 years; b 34 years; c 24 years; d the median — the two youngest volunteers pull the mean below what most of the group is. **Q4** a A: 12, 15, 18, 18, 21, 24, 27; B: 9, 14, 16, 20, 22, 25, 30; b A 18, B 20; c A 15, B 21 — Team A is more consistent. **Q5** mean 10, median 9.5, range 13. **Q6** a median 9, range 38; b 14; c the range most (38 against 4 without it), the median least (it stays 9 either way). **Q7** a $2+6+12+9+5+2=36$; b 10–15 min; c 16 students. **Q8** a $\frac{29}{36} \approx 80.6\%$; b exaggerated — 80.6% is a large majority, but “almost all” is not true when about one in five takes 20 minutes or more. **Q9** a symmetric; b positive skew; c bi-modal; d negative skew. **Q10** a mean $\frac{67}{7} \approx 9.6$, median 6, range 27; b mean $\frac{44}{7} \approx 6.3$, median 6, range 6; c the range changed most (27 to 6), the median not at all. **Q11** a 14–16 °C, 12 days; b 2 days; c 8 days; d positive skew (mean above median, tail to the warm side). **Q12** a set Q; b set P — with a range of 4 the values hug the mean, while Q’s range of 20 hints at values far from 12; c e.g. “Both sets centre on 12, but Q’s values are far more spread out than P’s.” **Q13** Mia: mean 15.4, median 15, range 16. Zoe: mean 14.8, median 13.5, range 23. Mia is typically a little higher on both centre measures and clearly more consistent; both sets rise gently with no outlier — similar, slightly positively skewed shapes. **Q14** Plot with stems 35–40; A’s leaves: 36: 2 5 5 8; 37: 0 1 4 6 9; 38: 3. B’s leaves: 35: 1 9; 36: 4 8; 37: 2 5 8; 38: 6; 39: 4; 40: 7. Key: $2/36/4$ means 36.2 s for A and 36.4 s for B. **Q15** a Melbourne 14.6 °C, Ballarat 11.6 °C; b Melbourne 8.5 °C, Ballarat 5.1 °C; c Ballarat 12 days; Melbourne 15 days. **Q16** Centre: Melbourne’s maxima were typically about 3 °C warmer (medians 14.6 against 11.6). Spread: Melbourne varied more, range 8.5 °C against 5.1 °C. Shape: Melbourne’s month leans slightly positively skewed with its warm-day tail, while Ballarat’s sits in one roughly symmetric block. **Q17** a 20; b 60–90 min; c 8 students; d positive skew. **Q18** a Melbourne — median 14.6 °C against 11.6 °C (or mean 14.9 against 11.2); b Melbourne — range 8.5 °C against 5.1 °C; c Melbourne has a warm-day tail (slight positive skew); Ballarat’s distribution is one symmetric-looking block. **Q19** a mean 17.9, median 12, range 65; b mean 11.25, median 11.5, range 9; c the outlier lifted the mean by about 6.7 and stretched the range by 56, while the median moved by only 0.5. **Q20** 12 (the five values must total $5 \times 12 = 60$, and the four known values total 48). **Q21** a $a=5$; b no — as long as a lies between 3 and 12 the largest and smallest values stay 12 and 3, so the range stays 9. **Q22** a positive skew; b negative skew; c symmetric; d bi-modal. **Q23** a the audience is two groups — children and adults — and almost nobody is aged 20–30; b nearly everyone is a child (under 20) or an adult (over 30), so 26.6 sits in the empty valley between the peaks; c e.g. “the audience is mostly primary-aged children and adults in their 30s and 40s” — or report the two peaks instead of one average. **Q24** e.g. 4, 6, 8, 8, 10, 14 — median $\frac{8+8}{2}=8$, range $14-4=10$. **Q25** e.g. A: 10, 11, 12, 13, 14 (mean 12, range 4) and B: 20, 21, 22, 23, 24 (mean 22, range 4) — same range, different means. **Q26** Centre: A is typically faster (median 37.05 s against 37.35 s). Spread: A is far more consistent (range 2.1 s against 5.6 s). **Q27** Plot with stems 1–4; A’s leaves: 2: 4 6 7 8 9; 3: 0 0 1 2 3 4 6. B’s leaves: 1: 8; 2: 2 5 7 9; 3: 0 1 3 5 8; 4: 1 5. Key: $6/2/2$ means 26 km for A and 22 km for B. **Q28** a mean 17, median 17, range 8; b the mean and median each rise by 5; the range is unchanged — adding the same amount to every value shifts centre but not spread. **Q29** e.g. 6, 8, 9, 9, 11, 13, 14 — mean $\frac{70}{7}=10$, median 9, range $14-6=8$. **Q30** a mean ≈ 26.9 km, median 25 km, range 27 km; b mean ≈ 27.9 km, median 25 km, range 27 km; c the mean moves (by 1 km) and would alert you; the median and the range hide the mistake entirely. **Q31** a 0.0 mm — more than half the days had no rain at all; b a few wet days, one of them 14.6 mm, pull the mean up while the median stays at 0; c positive skew (a strong one). **Q32** A: mean 68.5, median 69, range 25. B: mean 69, median 68, range 53. The centres are almost identical, but B is far more spread (range 53 against 25) and positively skewed, with a tail of high scores (85, 90, 95) and low ones (42, 48); A is roughly symmetric. **Q33** a true — mean 14.9 above median 14.6 with a warm tail; b false — the modal class is 14–16 °C with 12 days, not 16–18 °C (6 days); c true — $19.5-11.0=8.5$ °C. **Q34** True — 11 of the 31 days stayed below 14.0 °C, which is at least ten. **Q35** e.g. 1, 8, 9, 9, 10 — median 9, mean $\frac{37}{5}=7.4$; the shape is negative skew (a low outlier drags the mean below the median). **Q36** A: mean 30 km, median 30 km, range 12 km. B: mean ≈ 31.2 km, median 30.5 km, range 27 km. A’s rides cluster close together and symmetrically around 30 km; B’s are positively skewed, with the long rides (38, 41, 45) pulling the mean above the median and stretching the range. The 45 km ride sits apart from B’s bulk and does the pulling. Interpretation: the riders log similar typical distances, but B’s training is built on occasional long rides while A rides much more evenly.

Fully-worked solutions

Q1. a Each side of a back-to-back plot holds one value per leaf: 12 girls' leaves and 12 boys' leaves. b Tallest girl: the top stem with a girl's leaf is 17, leaf 1 — 171 cm. Tallest boy: stem 17, leaf 8 — 178 cm. c Shorter than 160 cm means stem 15 only: girls 152, 155, 157, 158 (4) and boys 154, 158 (2). $\therefore$ a 12 each; b 171 cm and 178 cm; c $4 + 2 = 6$ students.

Q2. a Reading the leaves against the stems: 21, 24, 27, 30, 32, 32, 35, 38, 41, 46. b Ten values: the median is the mean of the 5th and 6th, both 32 — median 32. c The mode is the repeated value, 32; the range is $46 - 21 = 25$. $\therefore$ a as listed; b 32; c mode 32, range 25.

Q3. a Sum: $18 + 22 + 25 + 25 + 31 + 34 + 34 + 34 + 39 + 40 + 42 = 344$; mean $= \frac{344}{11} \approx 31.3$ years. b Eleven values: the 6th is the median — 34 years. c Range: $42 - 18 = 24$ years. d The mean (31.3) sits below the median (34) because the two youngest volunteers form a small low tail; the median describes the typical volunteer better. $\therefore$ a 31.3; b 34; c 24; d the median.

Q4. a Team A: 12, 15, 18, 18, 21, 24, 27. Team B: 9, 14, 16, 20, 22, 25, 30. b A: the 4th of seven is 18. B: the 4th of seven is 20. c A: $27 - 12 = 15$. B: $30 - 9 = 21$. The smaller range is the more consistent team. $\therefore$ a as listed; b 18 and 20; c 15 and 21 — Team A is more consistent.

Q5. Sum: $4 + 7 + 7 + 9 + 10 + 12 + 14 + 17 = 80$; mean $= \frac{80}{8} = 10$. Eight values: the middle two are 9 and 10, so the median is $\frac{9 + 10}{2} = 9.5$. Range: $17 - 4 = 13$. $\therefore$ Mean 10, median 9.5, range 13.

Q6. a Ordered data: 7, 8, 8, 9, 10, 11, 45 — the 4th value is the median, 9. Range: $45 - 7 = 38$. b Sum = 98; mean $= \frac{98}{7} = 14$. c Without the 45 the range would be $11 - 7 = 4$ and the mean about 8.8, so the range is distorted most; the median stays 9 either way, so it is distorted least. $\therefore$ a median 9, range 38; b 14; c range most, median least.

Q7. a $2 + 6 + 12 + 9 + 5 + 2 = 36$ students. b The tallest bar is the class 10–15 min. c At least 15 minutes: $9 + 5 + 2 = 16$ students. $\therefore$ a 36; b 10–15 min; c 16.

Q8. a Less than 20 minutes: $2 + 6 + 12 + 9 = 29$ students; $\frac{29}{36} \approx 80.6\%$. b 80.6% is a strong majority, but the remaining 7 students — about one in five — take 20 minutes or more, so “almost all” overstates it. $\therefore$ a 80.6%; b the sentence exaggerates: one in five is not “almost none”.

Q9. a Mirror halves with mean near median: symmetric. b A tail of high values to the right, mean above median: positive skew. c Two peaks with a valley: bi-modal. d A tail of low values to the left, mean below median: negative skew. $\therefore$ a symmetric; b positive skew; c bi-modal; d negative skew.

Q10. a Before: sum 67, mean $\frac{67}{7} \approx 9.6$; median (4th of 7) 6; range $30 - 3 = 27$. b After: sum 44, mean $\frac{44}{7} \approx 6.3$; median 6; range $9 - 3 = 6$. c The range fell from 27 to 6 — the biggest change; the mean fell too; the median did not move at all. $\therefore$ a 9.6, 6, 27; b 6.3, 6, 6; c range most, median not at all.

Q11. a The tallest bar is 14–16 °C with 12 days. b Below 12 °C: the 10–12 class holds 2 days. c At least 16 °C: $6 + 2 = 8$ days. d Mean above median with the tail at the warm end: slight positive skew. $\therefore$ a 14–16 °C, 12 days; b 2; c 8; d positive skew.

Q12. a The range measures spread: 20 against 4, so set Q is more spread out. b Set P — with all values within 4 of each other, the mean 12 is close to every value; Q's range of 20 means values can sit far from 12, so its mean is a weaker picture of a typical value. c Any sentence that keeps centre equal and spread different, e.g. “Both sets centre on 12, but Q's values are far more spread out than P's.” $\therefore$ a Q; b P; c as above.

Q13. Mia: sum 154, mean 15.4; median $\frac{15+15}{2}=15$; range $24-8=16$. Zoe: sum 148, mean 14.8; median $\frac{13+14}{2}=13.5$; range $28-5=23$. Centre: Mia is a little higher on both the mean and the median. Spread: Zoe's range of 23 against 16 makes Mia the more consistent scorer. Shape: both sets climb gently with a high-side tail and no value sitting apart — similar, slightly positively skewed shapes. $\therefore$ Mia typically scores slightly more and is clearly more consistent; the shapes are alike.

Q14. Stems are the whole seconds, 35 to 40; each leaf is one tenth of a second; on each side the smallest leaf sits next to the stem.

Swimmer A	Stem	Swimmer B
	35	1 9
8 5 5 2	36	4 8
9 6 4 1 0	37	2 5 8
3	38	6
	39	4
	40	7

Key: 2 / 36 / 4 means 36.2 s for Swimmer A and 36.4 s for Swimmer B. $\therefore$ The plot shows at a glance what the numbers confirm: A's leaves crowd into 36–37, B's spread from 35 to 40.

Q15. a Median = 16th of the 31 ordered leaves: Melbourne 14.6 °C; Ballarat 11.6 °C. b Melbourne: $19.5-11.0=8.5$ °C. Ballarat: $13.6-8.5=5.1$ °C. c Ballarat leaves at 12.0 and above: 12 days. Melbourne leaves at 15.0 and above: 15 days. $\therefore$ a 14.6 and 11.6 °C; b 8.5 and 5.1 °C; c 12 and 15 days.

Q16. Centre: Melbourne's month was typically about 3 °C warmer — median 14.6 °C against 11.6 °C, mean 14.9 against 11.2. Spread: Melbourne varied more — range 8.5 °C against Ballarat's 5.1 °C. Shape: Melbourne's tail of warm days (17–19.5 °C) gives a slight positive skew, while Ballarat's days sit in one block that is roughly symmetric. $\therefore$ Three sentences: warmer centre, wider spread, and a skew Melbourne has and Ballarat lacks.

Q17. a $4+8+6+2=20$ students. b The tallest bar: 60–90 min. c At least 90 minutes: $6+2=8$ students. d The tall bars sit at the left with a tail to the right: positive skew. $\therefore$ a 20; b 60–90 min; c 8; d positive skew.

Q18. a Melbourne: median 14.6 °C (mean 14.9) against Ballarat's 11.6 °C (mean 11.2) — Melbourne typically warmer by about 3 °C. b Melbourne's range is 8.5 °C against Ballarat's 5.1 °C — Melbourne more variable. c Melbourne shows a warm-day tail (slight positive skew); Ballarat's distribution is a single, roughly symmetric block. $\therefore$ a Melbourne; b Melbourne; c the warm tail at Melbourne.

Q19. a Sum = 161; mean = $\frac{161}{9} \approx 17.9$. Median (5th of 9) = 12. Range = $71-6=65$. b Without the 71: sum = 90; mean = $\frac{90}{8} = 11.25$. Median = $\frac{11+12}{2} = 11.5$. Range = $15-6=9$. c The outlier lifted the mean by $17.9-11.25 \approx 6.7$ and stretched the range by $65-9=56$, while the median moved by only 0.5. $\therefore$ a 17.9, 12, 65; b 11.25, 11.5, 9; c mean and range dragged, median almost untouched.

Q20. Five values with mean 12 must total $5 \times 12 = 60$. The four known values total $9+11+13+15=48$. The fifth value is $60-48=12$. $\therefore$ 12.

Q21. a With a between 5 and 8 the ordered list is 3, 5, a , 8, 9, 12; the median is $\frac{a+8}{2}=6.5$, so $a+8=13$ and $a=5$.

(Check: 3, 5, 5, 8, 9, 12 has median $\frac{5+8}{2}=6.5$.) b No. The range is largest minus smallest; as long as a stays between 3 and 12, the largest value is 12 and the smallest is 3, so the range stays $12-3=9$ whatever a is. $\therefore$ a $a=5$; b no — the range stays 9.

Q22. a Mean well above median with one high value: the tail points right — positive skew. b Mean below median with one low value: the tail points left — negative skew. c Mean equal to median with folding halves: symmetric. d Two peaks with a valley: bi-modal. $\therefore$ a positive skew; b negative skew; c symmetric; d bi-modal.

Q23. a The audience is two groups — school children and their parents and grandparents — and almost nobody present is aged 20–30, so that class is empty. b The mean mixes the two peaks: nearly everyone is under 20 or over 30, so a “typical” age of 26.6 matches almost no actual person. c Describe the two groups instead — e.g. “mostly children aged 5–18 and adults aged 30–52” — or report the two modal classes rather than one average. $\therefore$ A bi-modal set is summarised by its two peaks, not by one number in the valley.

Q24. One answer: 4, 6, 8, 8, 10, 14. Median: the middle two values are 8 and 8, so $\frac{8+8}{2}=8$. Range: $14-4=10$. $\therefore$

Both summaries check out; other answers are possible.

Q25. One answer: A: 10, 11, 12, 13, 14 and B: 20, 21, 22, 23, 24. Ranges: $14-10=4$ and $24-20=4$ — equal.

Means: $\frac{60}{5}=12$ and $\frac{110}{5}=22$ — different. $\therefore$ Shifting a whole set changes the mean and leaves the range, which is what the example exploits.

Q26. Centre: A’s median 37.05 s against B’s 37.35 s (means 37.13 against 37.54) — A is typically faster. Spread: A’s range 2.1 s against B’s 5.6 s — A is far more consistent. $\therefore$ Two sentences: faster centre, tighter spread, both in A’s favour.

Q27. Stems are the tens (1 to 4); leaves are the units; smallest leaf next to the stem on each side.

Rider A	Stem	Rider B
	1	8
9 8 7 6 4	2	2 5 7 9
6 4 3 2 1 0 0	3	0 1 3 5 8
	4	1 5

Key: 6 / 2 / 2 means 26 km for Rider A and 22 km for Rider B. $\therefore$ The plot already shows the difference: A’s leaves pile into stems 2–3, B’s run from stem 1 to stem 4.

Q28. a New set: 13, 15, 17, 19, 21. Mean $=\frac{85}{5}=17$; median 17; range $21-13=8$. b The old mean and median were 12 and 12, so both rose by exactly 5; the old range was $16-8=8$, unchanged. Generalising: adding the same amount to every value shifts the centre by that amount and leaves the spread alone. $\therefore$ a 17, 17, 8; b centre shifts, spread does not.

Q29. One answer: 6, 8, 9, 9, 11, 13, 14. Mean: sum = 70, so $\frac{70}{7}=10$. Median: the 4th of seven is 9. Range: $14-6=8$. $\therefore$ All three summaries check.

Q30. a True set ordered: 14, 18, 22, 25, 32, 36, 41. Sum 188; mean $\frac{188}{7}\approx 26.9$. Median (4th) 25. Range $41-14=27$.

b Recorded set ordered: 14, 18, 22, 25, 36, 39, 41. Sum 195; mean $\frac{195}{7}\approx 27.9$. Median 25. Range 27. c The mean rises by 1 km, so it would raise a suspicion; the median and the range are identical before and after, and would hide the mistake completely. $\therefore$ a 26.9, 25, 27; b 27.9, 25, 27; c the mean alerts, the median and range hide.

Q31. a The median is 0.0 mm: in the ordered month, the middle day had no rain — more than half of July’s days were dry. b Sixteen days sat at 0.0 mm, but a few wet days, topped by 14.6 mm, add to the total the mean divides up; the median, counting only order, never sees them. c A tall bar at the left and a long low tail to the right: positive skew, and a strong one. $\therefore$ a 0.0 mm, a mostly dry month; b the wet tail pulls the mean; c positive skew.

Q32. Class A: sum 685, mean 68.5; median $\frac{68+70}{2}=69$; range $80-55=25$. Class B: sum 690, mean 69; median $\frac{66+70}{2}=68$; range $95-42=53$. Centre: nearly identical (68.5 against 69). Spread: B's range is more than double A's — B is far more variable. Shape: A climbs evenly, roughly symmetric; B spreads to both tails (42, 48 below; 85, 90, 95 above) with the stronger high tail — positive skew. $\therefore$ Same centre, very different spread and shape: A is a tight, even class; B is a wide one with extreme scorers at both ends.

Q33. a True: mean 14.9°C sits above median 14.6°C and the tail runs to the warm days — slight positive skew. b False: the tallest bar is $14-16^{\circ}\text{C}$ with 12 days; the $16-18^{\circ}\text{C}$ class holds only 6. c True: range $=19.5-11.0=8.5^{\circ}\text{C}$. $\therefore$ a true; b false (modal class $14-16^{\circ}\text{C}$); c true.

Q34. Count the days below 14.0°C from the data behind the histogram: 11 days. The claim needs at least 10; 11 delivers it. $\therefore$ The claim is true — by one day more than stated.

Q35. One answer: 1, 8, 9, 9, 10. Median: the 3rd of five is 9. Mean: $\frac{1+8+9+9+10}{5}=\frac{37}{5}=7.4$. $9>7.4$ as required; the one low value drags the mean below the median, which is the signature of negative skew. $\therefore$ e.g. 1, 8, 9, 9, 10 — median 9, mean 7.4; negative skew.

Q36. Rider A: sum 360, mean $\frac{360}{12}=30$; median $\frac{30+30}{2}=30$; range $36-24=12$. Rider B: sum 374, mean $\frac{374}{12}\approx 31.2$; median $\frac{30+31}{2}=30.5$; range $45-18=27$. Shape: A's leaves crowd symmetrically about stem 3 — roughly symmetric; B's leaves run out a tail at stems 3–4 (38, 41, 45) — positive skew, with mean above median (31.2 against 30.5). Apart value: B's 45 km ride sits away from the bulk of B's rides; it helps pull B's mean above the median and stretches B's range. $\therefore$ The riders log similar typical distances (medians 30 and 30.5 km), but B's training rests on occasional long rides — wider range, positive skew — while A rides far more evenly; a comparison of distributions needs exactly these three strands: centre, spread, shape.

Choosing and justifying a data display

6B · Chapter 6 — Statistics · Level 9

Content description VC2M9ST04 (word for word): *choose appropriate forms of display or visualisation for a given type of data; justify selections and interpret displays for a given context*

Achievement standard anchor (AS 16): *Students compare and analyse the distributions of multiple numerical data sets, choose representations, describe features of these data sets using summary statistics and the shape of distributions, and consider the effect of outliers.*

Elaboration ids for linkage: VC2M9ST04.E01, VC2M9ST04.E02, VC2M9ST04.E03, VC2M9ST04.E04

Prerequisite pointer: 6A is assumed — the features of a distribution (centre and spread; shape: positive skew, negative skew, symmetric, bi-modal; outliers) and the displays used to compare them (histograms, dot plots, stem-and-leaf plots). They are used here as tools, not re-taught. 6A named the features a display can show; 6B decides which display to use, and says why.

Note on Level 9 statistics: the only summary numbers used at this level are the mean, the median and the range (D3.10). Quartiles and box plots belong to Level 10 and are not used in this section.

Note on sources: the two July 2026 weather data sets in this section are real. Source: Bureau of Meteorology, “Daily Weather Observations for Melbourne (Olympic Park), Victoria for July 2026”, station 086338, prepared 14 August 2026. Retrieved 20 August 2026. Every other quantitative context is a clearly-labelled SCENARIO with invented parameters (R2).

Theory

Where 6A leaves off

In 6A you learned to name the features of a distribution: its centre, its spread, its shape (positive skew, negative skew, symmetric, bi-modal), and any outliers. You also learned the displays that show those features — histograms, dot plots and stem-and-leaf plots. 6B asks the question that comes *before* all of that: you have a set of data and a question to answer, so **which display do you use, and why?**

The content description’s verbs are *choose* and *justify*. A justification in this subtopic always names three things:

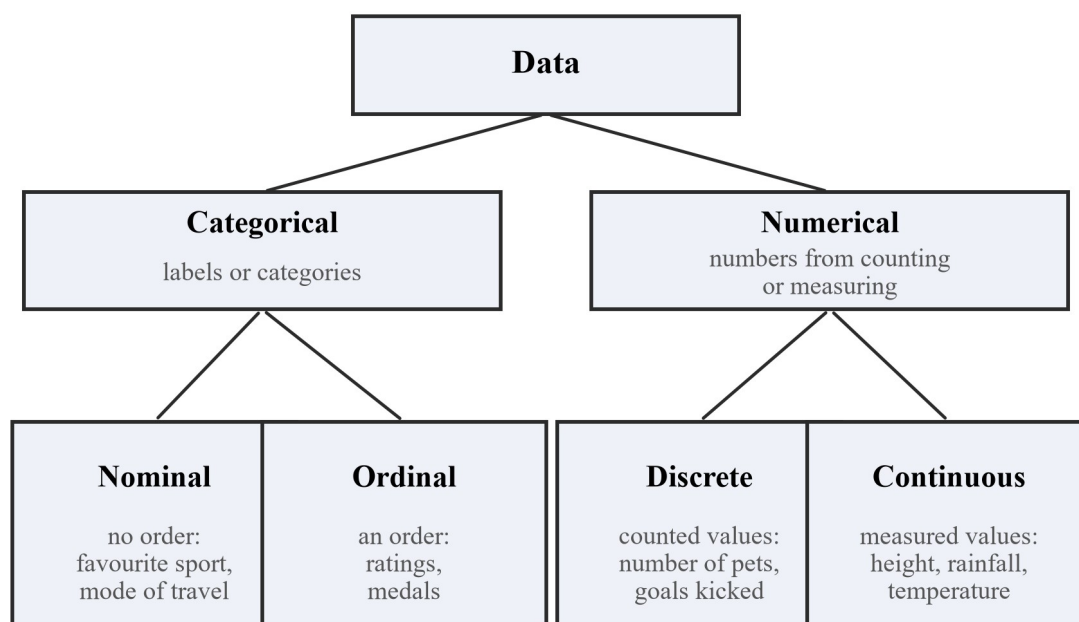
1. the **type of the data**;
2. what the **question** asks the display to show;
3. why the **chosen display** shows exactly that.

The four data types

Every choice in this subtopic starts with the data type. There are two families, each split in two.

- **Categorical data** has values that are labels — words, not numbers.
 - **Nominal** — the labels have no natural order: eye colour, suburb, favourite sport.
 - **Ordinal** — the labels have a natural order, but the gaps between them are not numbers: medals (gold, silver, bronze), ratings (poor, fair, good).
- **Numerical data** has values that are numbers.
 - **Discrete** — the values come from counting: whole numbers with gaps between them, such as the number of pets in a household.
 - **Continuous** — the values come from measuring: any value in a range is possible, such as height, time or temperature.

The four data types — ask what the values are, then choose



A tree diagram that splits data into categorical and numerical, then into nominal and ordinal, and into discrete and continuous, each with a worked example.

Two quick tests sort almost anything. *Could the answer be halfway between two values?* If yes, the data is continuous; if the answer must jump from one whole value to the next, it is discrete. *Is the value a label rather than a number?* If yes, it is categorical — and the order test (do the labels rank?) separates nominal from ordinal.

Matching displays to data types

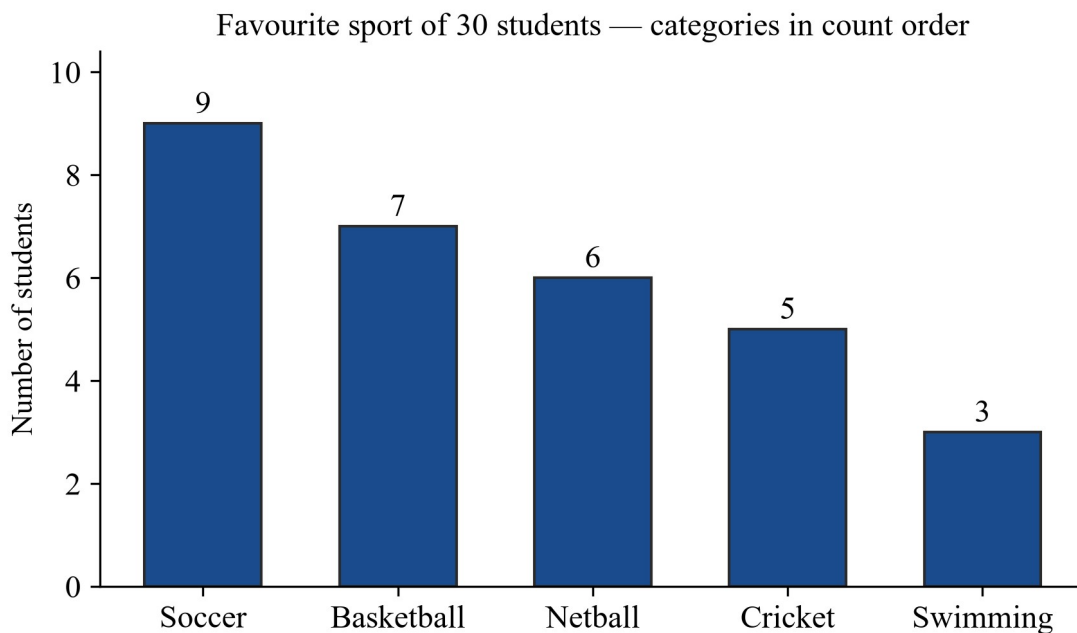
The type of the data decides the family of displays that can work. The question then decides which one of that family to use.

Data type	Displays that work	What the display shows well
Categorical (nominal)	bar chart; pie chart when there are few categories and shares of the whole matter	how the counts compare across categories
Categorical (ordinal)	bar chart with the categories kept in order	how the counts compare, <i>and</i> the order
Numerical (discrete), small data set	dot plot; stem-and-leaf plot	every raw value; clusters, gaps, centre and spread
Numerical (discrete or continuous), larger data set	histogram	the shape, centre and spread, and any outliers
Numerical, recorded in order over time	line graph	the trend — rises, falls, and when they happen
Totals made of parts, across categories	stacked bar chart	the totals, and the parts inside each total
Parts recorded over time	stacked area chart	the trend of the total, and how the mix changes

Two warnings go with the table. A line graph joins points in order, so it is only honest when the order means something — almost always time. Joining categories with line segments invents an order that is not there (Question 18). And a pie chart shows shares of a whole; the moment the counts themselves must be compared, a bar chart is easier to read, because the eye compares lengths far better than it compares angles.

Nominal data: bar charts

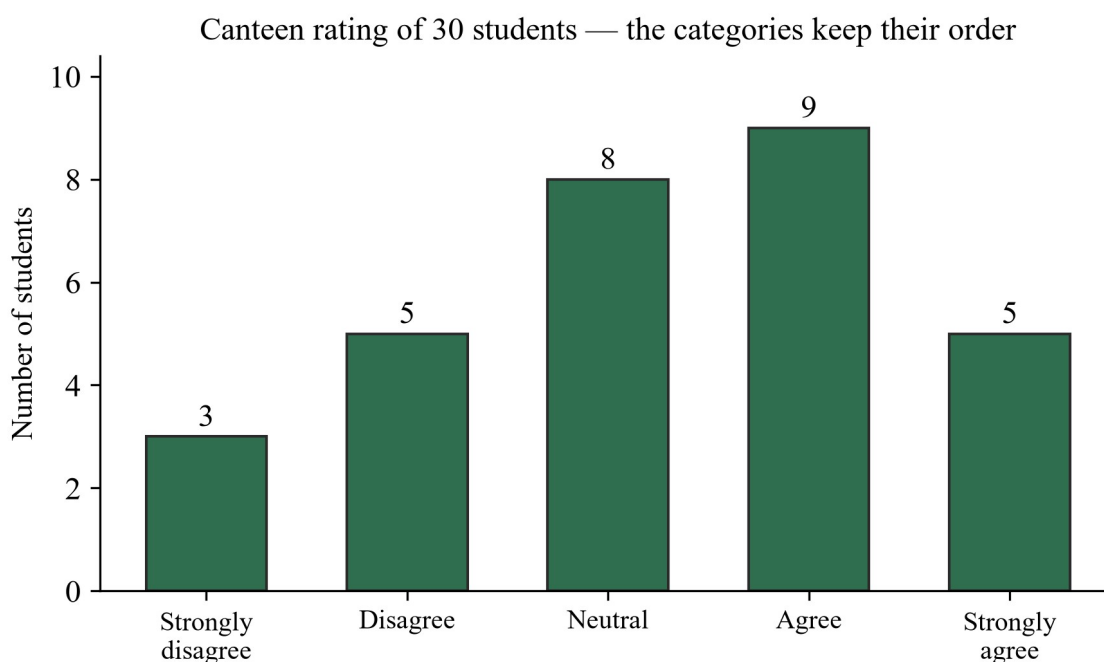
SCENARIO: thirty students name their favourite sport. The values are labels with no order, so the data is nominal. A bar chart shows how the counts compare; the bars can stand in any order, although sorting them by size makes the comparison easiest to read.



A bar chart of thirty students' favourite sport: soccer 9, basketball 7, netball 6, cricket 5, swimming 3, with the bars sorted from largest to smallest.

Ordinal data: keep the order

SCENARIO: thirty students rate a statement from *strongly disagree* to *strongly agree*. The values are still labels, but now the order is the point of the data. The display must keep the categories in their natural order — alphabetical order or order by size would destroy the information.



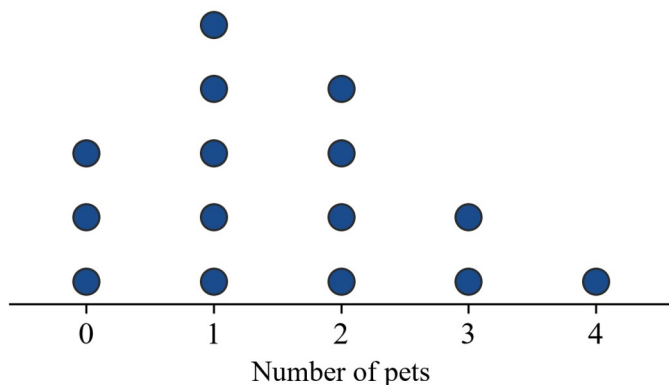
A bar chart of thirty students' ratings of a statement, with the five categories kept in order from strongly disagree to strongly agree.

Small numerical data sets: dot plots and stem-and-leaf plots

When the values are numbers and there are not too many of them, a dot plot or a stem-and-leaf plot keeps every raw value on show. Both show clusters and gaps at a glance, and both let you read off the median and the range directly. A stem-and-leaf plot has one extra strength: it keeps the exact values, sorted.

Number of pets for 15 students — two displays of the same data

Dot plot — every value shows



Stem-and-leaf plot

```

0 | 0 0 0
1 | 1 1 1 1 1
2 | 2 2 2 2
3 | 3 3
4 | 4
  
```

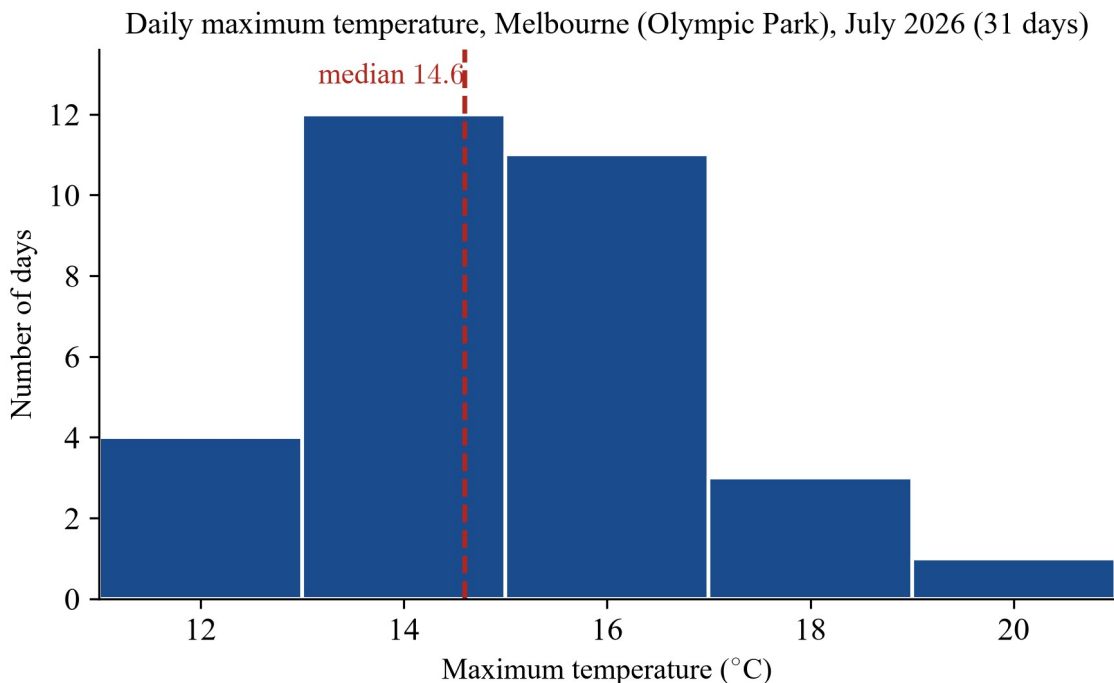
Key: 2 | 2 means 2 pets

The same fifteen household-pet counts drawn two ways: a dot plot above, and a stem-and-leaf plot with a key below it.

Histograms: continuous data in groups

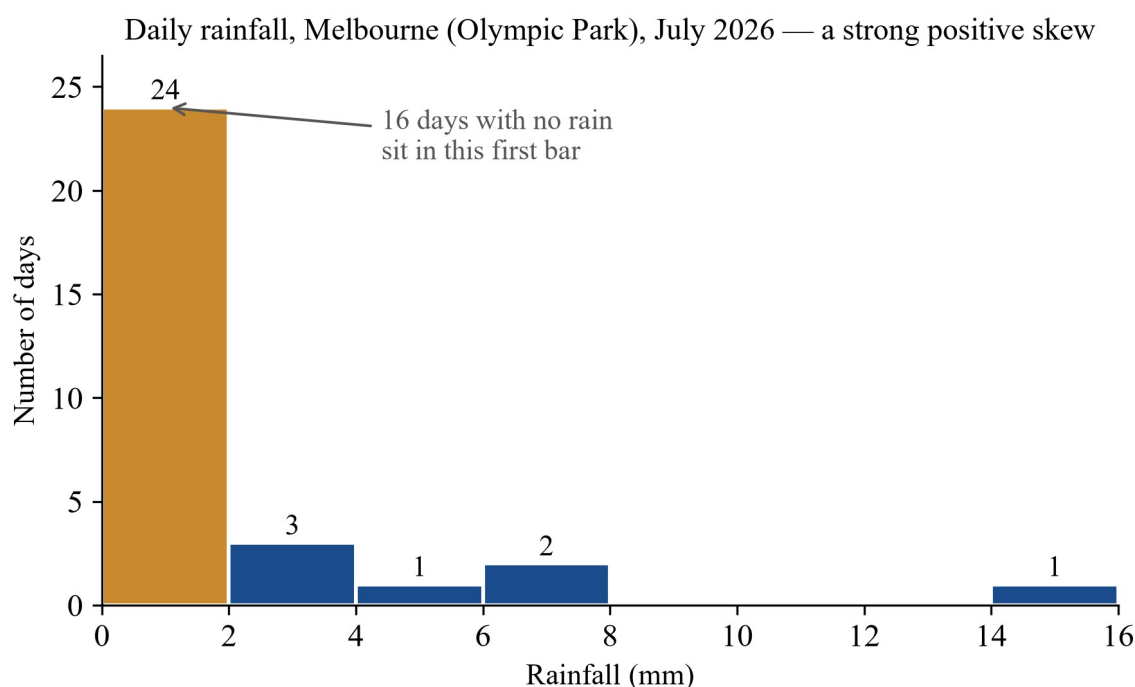
When there are many values, or the values are continuous, listing every value stops being useful. A histogram groups the values into equal-width intervals and shows how many fall in each group — so the shape of the distribution appears at a glance, and the centre, spread and outliers can be read from it. The mean, median and range are still calculated from the raw values, not read off the groups.

The figure below shows the real daily maximum temperatures recorded at Melbourne (Olympic Park) in July 2026. The median, 14.6 °C, is marked. Most days sit between 13 and 17 °C, with a small tail towards the warmer days.



A histogram of the 31 July 2026 daily maximum temperatures at Melbourne (Olympic Park), grouped into 2-degree intervals from 11 to 21, with the median of 14.6 degrees marked.

The same month's rainfall tells a very different story. Sixteen of the thirty-one days had no rain at all, and one wet day reached 14.6 mm — the distribution is strongly positively skewed. The histogram makes that skew obvious in a way no single number could.

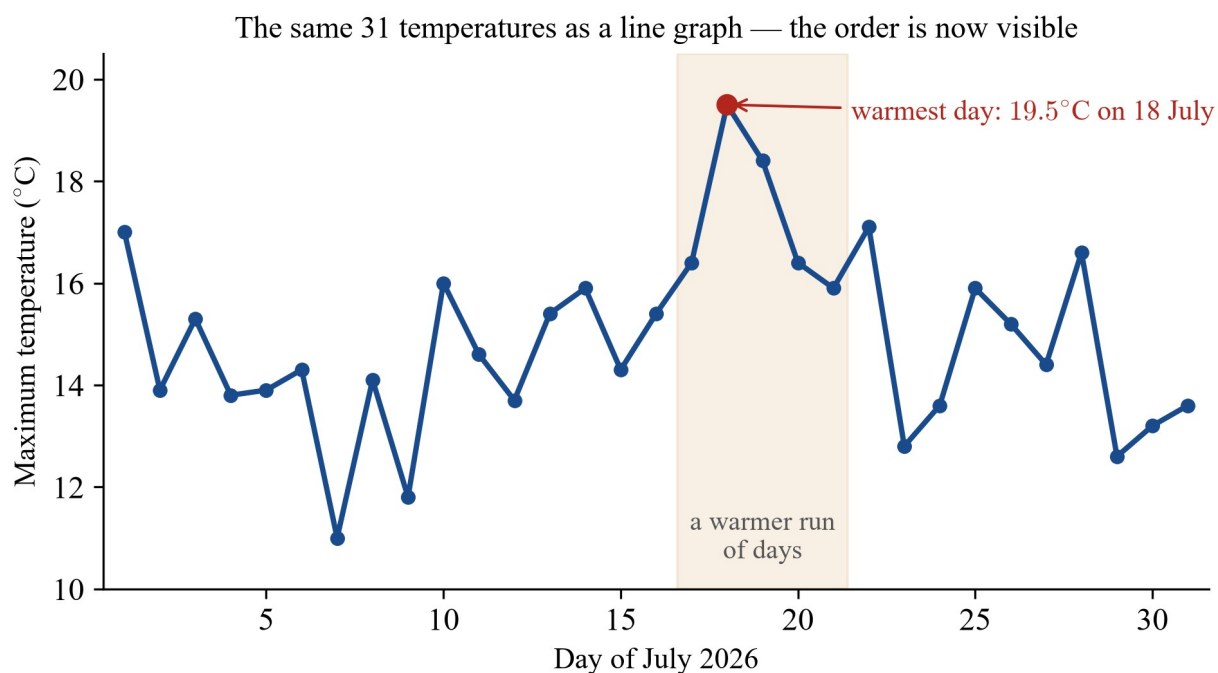


A histogram of the 31 July 2026 daily rainfalls at Melbourne (Olympic Park), grouped into 2-millimetre intervals from 0 to 16, showing a tall first bar of 24 days and a note that 16 days had no rain.

Same data, different display

The July maximum temperatures can also be drawn in the order they were recorded, as a line graph. Both displays are honest; they answer different questions.

- The **line graph** shows what happened *when*: the day-to-day rises and falls, and the warm spell from the 17th to the 21st.
- The **histogram** shows what the month *was like overall*: the shape of the distribution, its centre and its spread.

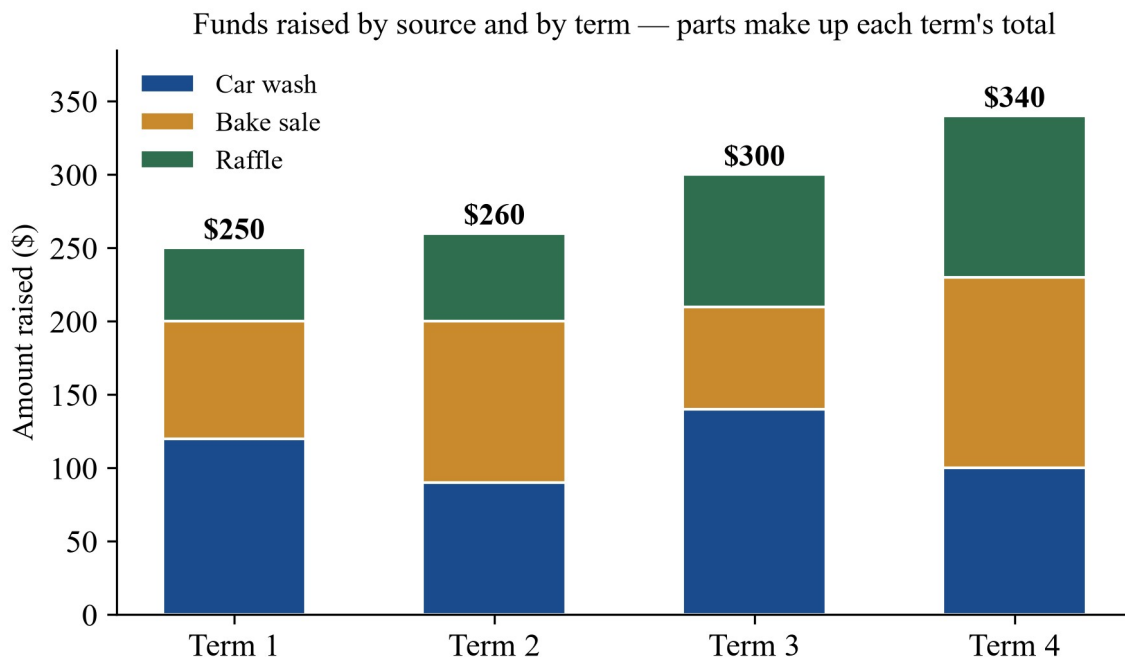


A line graph of the 31 July 2026 daily maximum temperatures at Melbourne (Olympic Park) in the order they were recorded, with the warmest day, 19.5 degrees on 18 July, marked.

Choosing between them is not a matter of taste: it follows from the question. “Which days were warm?” asks for the line graph; “what is a typical July day?” asks for the histogram.

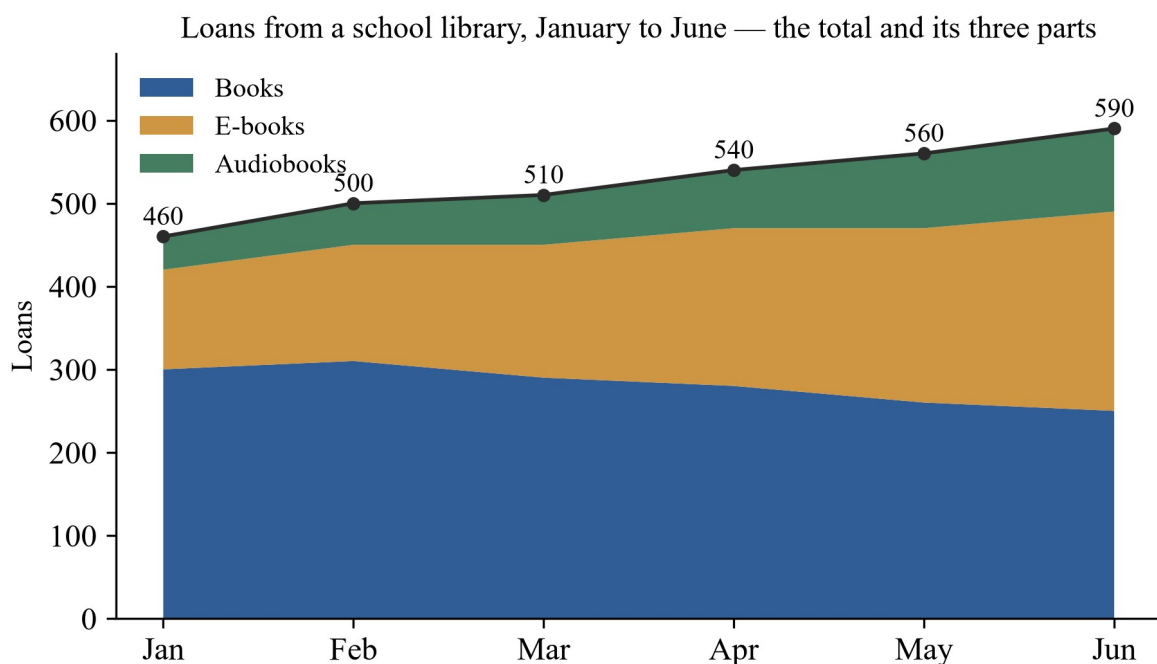
Categories with parts inside them

Sometimes each category is a total that can be split into smaller parts — four terms of fundraising, split by source. A **stacked bar chart** draws each total as one bar split into segments, so it supports two comparisons at once: totals across categories (the bar heights), and parts within each category (the segments).



A stacked bar chart of a school’s fundraising across four terms, each bar split into car wash, bake sale and raffle, with the term totals 250, 260, 300 and 340 dollars marked above the bars.

When the categories are points in time and there are several of them, a **stacked area chart** does the same job: the height of the whole shape tracks the total over time, and the bands inside it track how the mix changes. A single line graph of the totals alone would hide the mix; separate line graphs of the parts would hide the total.



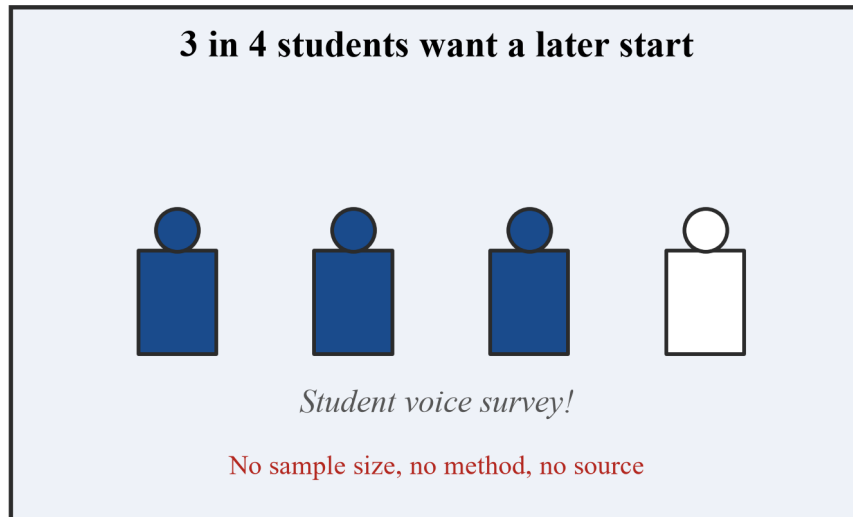
A stacked area chart of a library’s monthly loans from January to June, split into books, e-books and audiobooks, with the monthly total drawn as a line above the bands.

Non-standard representations: infographics

An **infographic** is a display built for a purpose and an audience: a poster, a social-media card, a campaign image. It usually mixes a big claim, one or two numbers and a picture. Infographics can communicate a statistical result quickly and well — and they can also hide everything that would let a reader check the claim. Three questions evaluate one:

1. What statistical question does it answer, and is the data there to answer it?
2. Who is the intended audience, and what does the display want them to do?
3. Could a reader check the claim — is the sample size, the method and the source given?

A poster-style display (SCENARIO) — made to persuade



An infographic poster reading “3 in 4 students want a later start”, with four person icons of which three are shaded, and a red line underneath noting that no sample size, no method and no source are given.

6D takes this further, where displays are analysed for the point of view they promote. In 6B the test is fit for purpose: does the display suit the data type and the question?

The justification checklist

This checklist is the working tool of the subtopic. Every worked example and every question below runs some part of it.

1. Name the data type: nominal, ordinal, discrete or continuous.
2. Name what the question asks the display to show: counts across categories, the shape of a distribution, a trend over time, or parts inside totals.
3. Choose the display that shows that feature, using the table above.
4. If two displays could work, name what each one adds — and pick by the question, not by taste.
5. Write the justification in one sentence: data type + question + display + what the display shows.

Worked examples

Example 1 — scaffolded

SCENARIO. Classify each variable as nominal, ordinal, discrete numerical or continuous numerical:

a a student’s blood group (A, B, AB, O) **b** the size of a school uniform (XS, S, M, L, XL) **c** the number of apps on a student’s phone **d** the mass of a student’s school bag

Working

Comment

a The values are labels: A, B, AB, O

The values are words, not numbers — the data is categorical.

Working	Comment
A is not “more than” B — the labels have no natural order	Categorical nominal .
b The labels XS, S, M, L, XL are words, but they rank from smallest to largest	Categorical, with a natural order: ordinal .
c The number of apps is counted: 0, 1, 2, 3, ...	Numbers from counting, with gaps between the values: numerical discrete .
d The mass is measured; any value in a range is possible, such as 4.35 kg	Numbers from measuring: numerical continuous .
∴ a nominal, b ordinal, c discrete, d continuous	Label or number first; then the order test for categories, and the count-or-measure test for numbers.

Example 2 — scaffolded

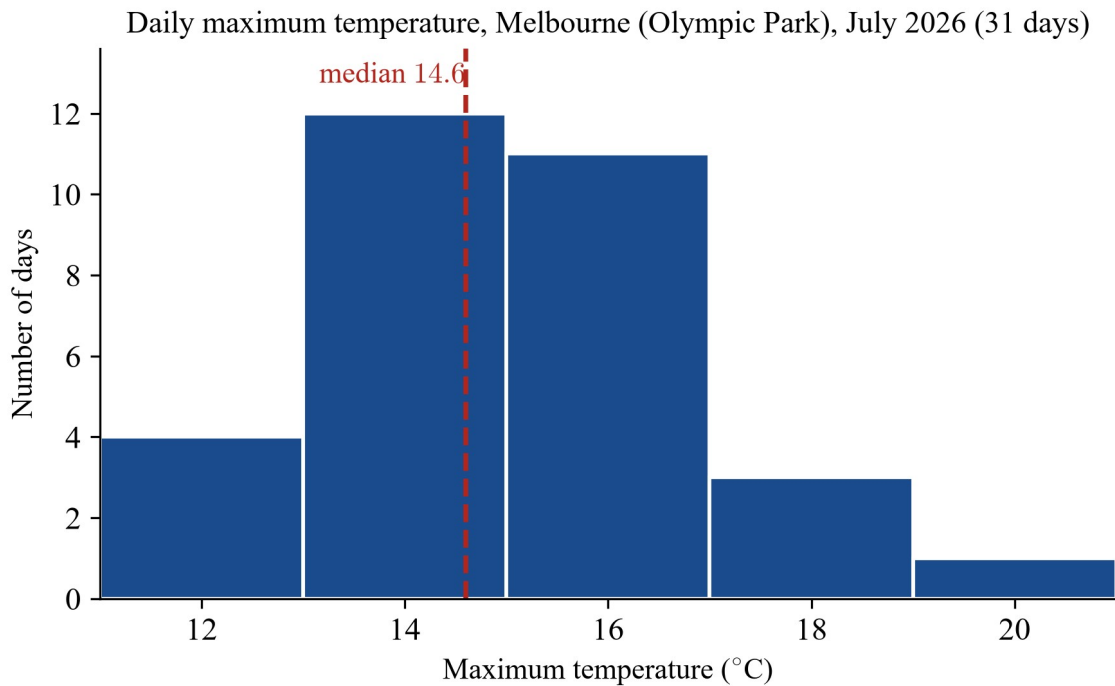
The daily maximum temperatures, in °C, recorded at Melbourne (Olympic Park) in July 2026 were:

17.0, 13.9, 15.3, 13.8, 13.9, 14.3, 11.0, 14.1, 11.8, 16.0, 14.6, 13.7, 15.4, 15.9, 14.3, 15.4, 16.4, 19.5, 18.4, 16.4, 15.9, 17.1, 12.8, 13.6, 15.9, 15.2, 14.4, 16.6, 12.6, 13.2, 13.6

Source: Bureau of Meteorology, “Daily Weather Observations for Melbourne (Olympic Park), Victoria for July 2026”, station 086338, prepared 14 August 2026. Retrieved 20 August 2026.

a Classify the data. b Choose a display and justify the choice. c Find the mean, median and range, and say what they tell you about the month.

Working	Comment
a The values are measured temperatures; any value in a range is possible	Numerical continuous .
b There are 31 values — too many for a tidy dot plot, and they need grouping	A histogram groups continuous data into equal-width intervals.
Interval counts: 4, 12, 11, 3, 1 for 11–13, 13–15, 15–17, 17–19, 19–21 °C	Most days sit in the two middle intervals, 13–17 °C.
c Mean: $\frac{462.0}{31} \approx 14.9$ °C	The sum of the 31 values is 462.0.
Median: the 16th of the 31 ordered values = 14.6 °C	With 31 values, the median is the middle one.
Range: $19.5 - 11.0 = 8.5$ °C	Warmest day minus coolest day.
∴ A histogram suits this data: the values are continuous, there are 31 of them, and the question asks for the shape as well as the centre and spread. The mean (≈ 14.9 °C) sits just above the median (14.6 °C), so the shape is roughly symmetric with a slight lean towards the warmer days	Checklist: data type (continuous) + question (shape, centre, spread) + display (histogram) + what it shows.



A histogram of the 31 July 2026 daily maximum temperatures at Melbourne (Olympic Park), grouped into 2-degree intervals from 11 to 21, with the median of 14.6 degrees marked.

Example 3 — unscaffolded

The daily rainfalls, in mm, recorded at Melbourne (Olympic Park) in July 2026 were:

6.6, 4.0, 7.0, 0.2, 1.6, 0, 0, 0, 0.2, 0, 0, 14.6, 1.4, 0.4, 0.2, 0.2, 0, 0, 0, 0, 0, 3.0, 0, 0, 0.6, 0, 0, 2.2, 2.2, 0

Source: Bureau of Meteorology, as in Example 2.

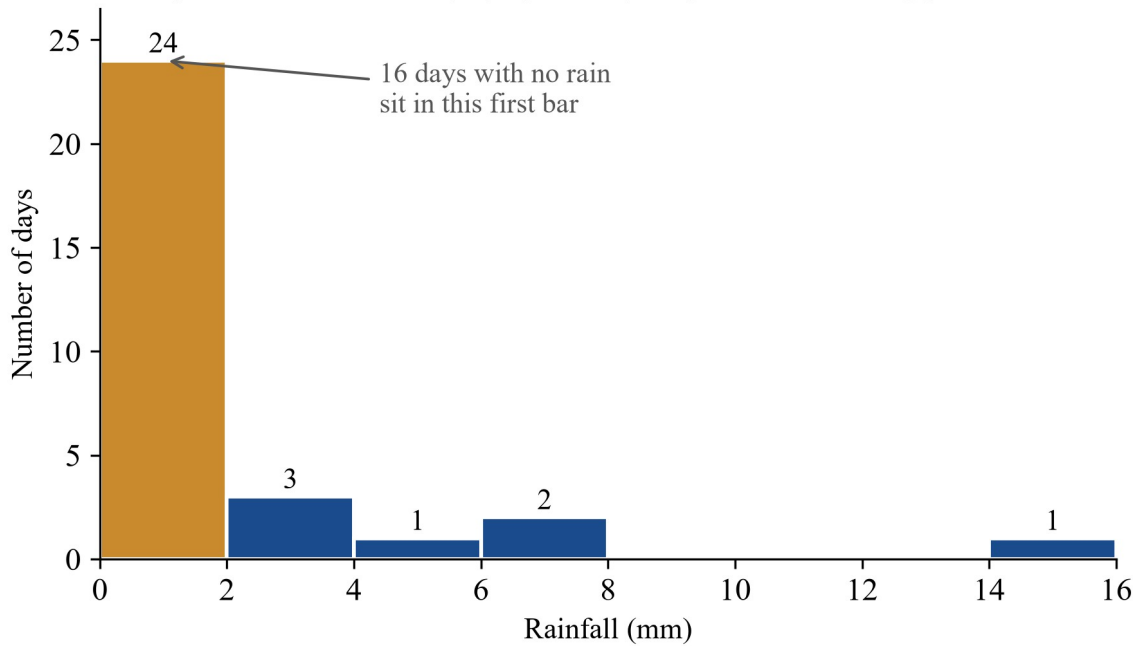
Choose a display for these data and justify the choice. Then find the mean, median and range, and describe a typical July day.

The values are measured amounts — any value in a range is possible — so the data is numerical continuous, and there are 31 values, so the values need grouping. A histogram in 2-mm intervals gives counts 24, 3, 1, 2, 0, 0, 0, 1 for 0–2, 2–4, 4–6, 6–8, 8–10, 10–12, 12–14, 14–16 mm. The sum of the 31 values is 44.4 mm, so the mean is

$\frac{44.4}{31} \approx 1.43$ mm. The median is the 16th of the 31 ordered values, which is 0 mm — sixteen of the thirty-one days

had no rain at all. The range is $14.6 - 0 = 14.6$ mm, and it comes almost entirely from one wet day. $\therefore$ The histogram is the right display: the data is continuous and grouped, and the shape is the story — sixteen dry days piled into the first interval, with a long tail out to 14.6 mm, a strong positive skew. The median (0 mm) describes a typical July day; the mean (≈ 1.43 mm) is pulled up by a few wet days and paints the month as wetter than it feels. A line graph would show *when* it rained, but the question asks about *amounts*, so the histogram answers it.

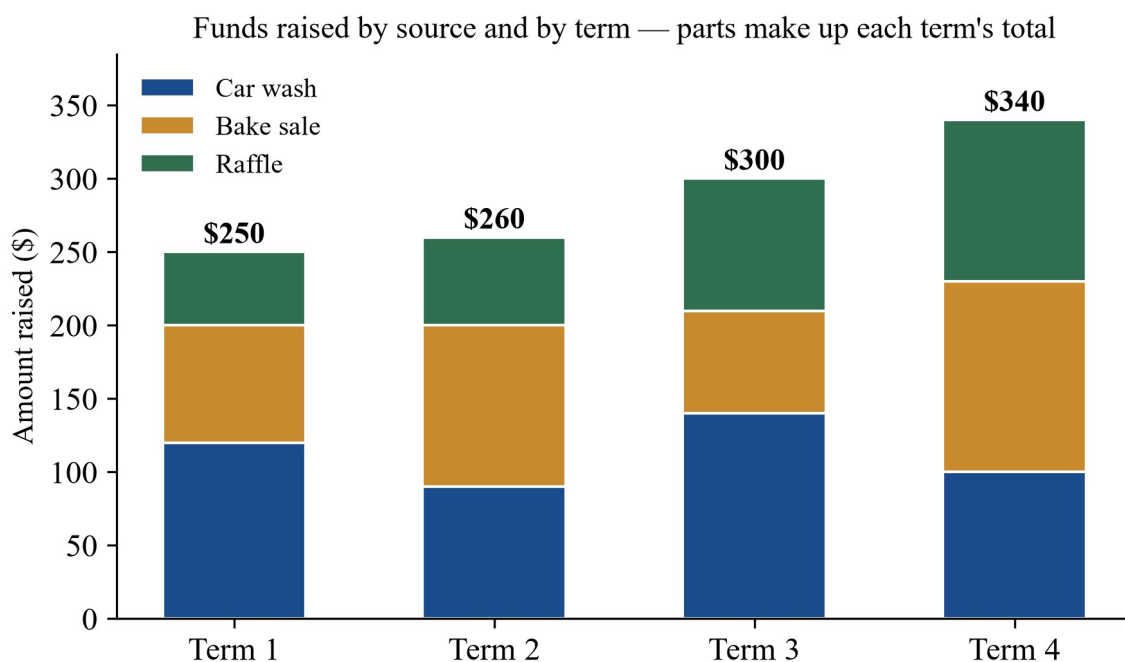
Daily rainfall, Melbourne (Olympic Park), July 2026 — a strong positive skew



A histogram of the 31 July 2026 daily rainfalls at Melbourne (Olympic Park), grouped into 2-millimetre intervals from 0 to 16, showing a tall first bar of 24 days and a note that 16 days had no rain.

Example 4 — unscaffolded

SCENARIO. A school council records the money it raises in each of four terms, split by source: car wash, bake sale and raffle. The figure shows the results.



A stacked bar chart of a school's fundraising across four terms, each bar split into car wash, bake sale and raffle, with the term totals 250, 260, 300 and 340 dollars marked above the bars.

a Find the total raised in each term, and the grand total. **b** Justify the choice of a stacked bar chart for the question “how did the total change across the year, and what made up each term?” **c** The council secretary wants to redraw the same information with weekly totals across the whole year. Name a better display for that version, and justify it.

a Reading the bar heights: Term 1 raised \$ 250, Term 2 raised \$ 260, Term 3 raised \$ 300 and Term 4 raised \$ 340. The grand total is $250 + 260 + 300 + 340 = \1150 . The biggest single term-source in the figure is the Term 3 car wash (\$ 140).

b The question has two parts, and the stacked bar chart answers both: the bar heights compare the totals across the four terms, and the segments show what made up each term's total. A plain bar chart would show only the totals; three separate bar charts (one per source) would show only the parts. The categories are terms — four of them, with no time flow to join — so bars suit them.

c Weekly totals across a whole year give about 52 points in order over time. Time is continuous, and the questions are the same two as before: the trend of the total, and the changing mix of sources. A stacked area chart does both jobs in one display: the top edge of the shape tracks the weekly total, and the bands inside it track each source's share. Fifty-two thin bars would be hard to read, and a single line graph of the totals would hide the mix. $\therefore$ Four categories with parts inside them: stacked bar chart. Many points in order over time with parts inside them: stacked area chart. The data type and the question choose the display in both cases.

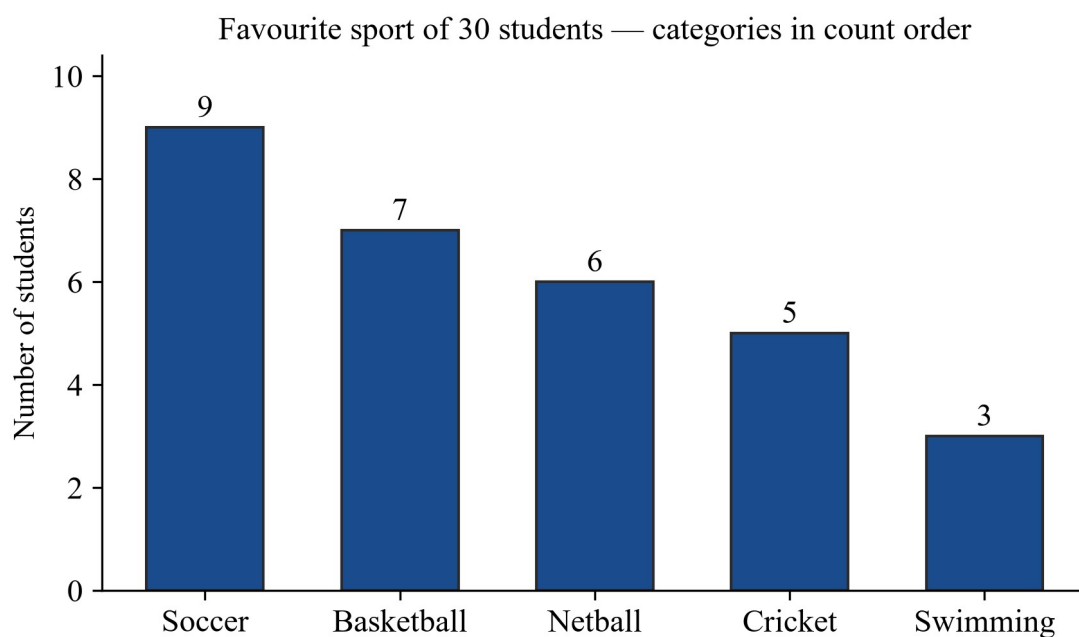
Practice questions

Scaffolded

Q1. Classify each variable as nominal, ordinal, discrete numerical or continuous numerical. Give a one-line reason for each.

- The colours of the four school houses (red, blue, green, gold).
- The medal won in an event at the athletics carnival (gold, silver, bronze).
- The number of goals kicked by a team across a season.
- The height of each student in the class.
- A student's rating of the new library (poor, fair, good).
- The number of wet days in each month of this year.

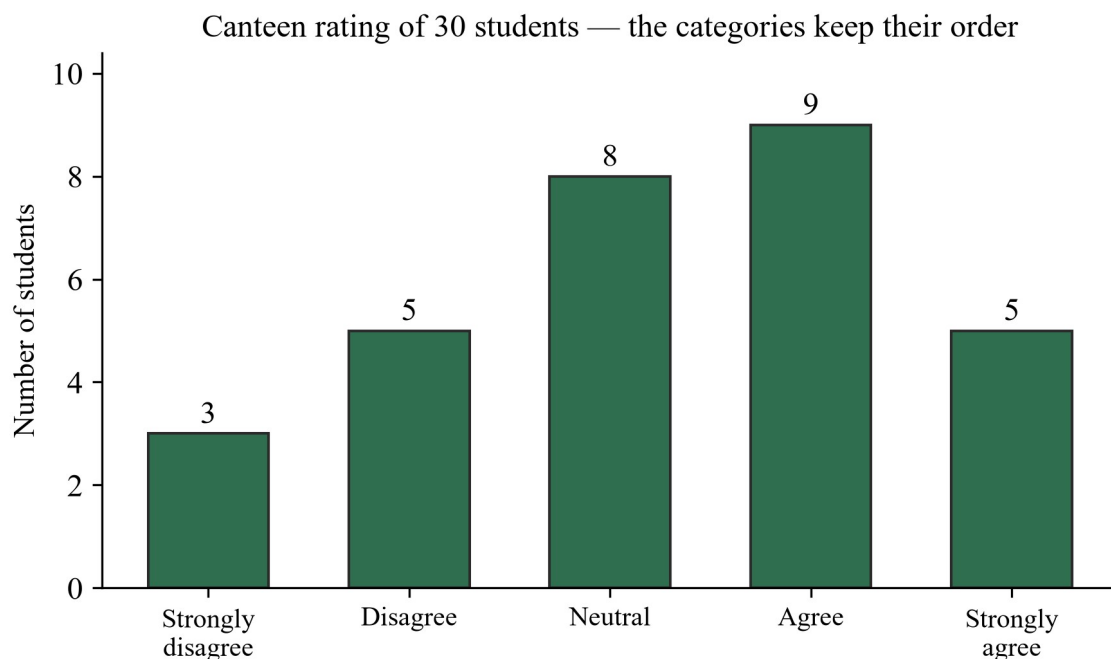
Q2. SCENARIO. Thirty students were asked to name their favourite sport. The figure shows the results.



A bar chart of thirty students' favourite sport: soccer 9, basketball 7, netball 6, cricket 5, swimming 3, with the bars sorted from largest to smallest.

- What is the data type of "favourite sport"? Give your reason.
- The figure uses a bar chart. Explain why a line graph would be wrong here.
- What fraction of the students chose soccer? Give your answer as a percentage.

Q3. SCENARIO. Thirty students rated the statement “the canteen offers healthy choices” from *strongly disagree* to *strongly agree*. The figure shows the results.

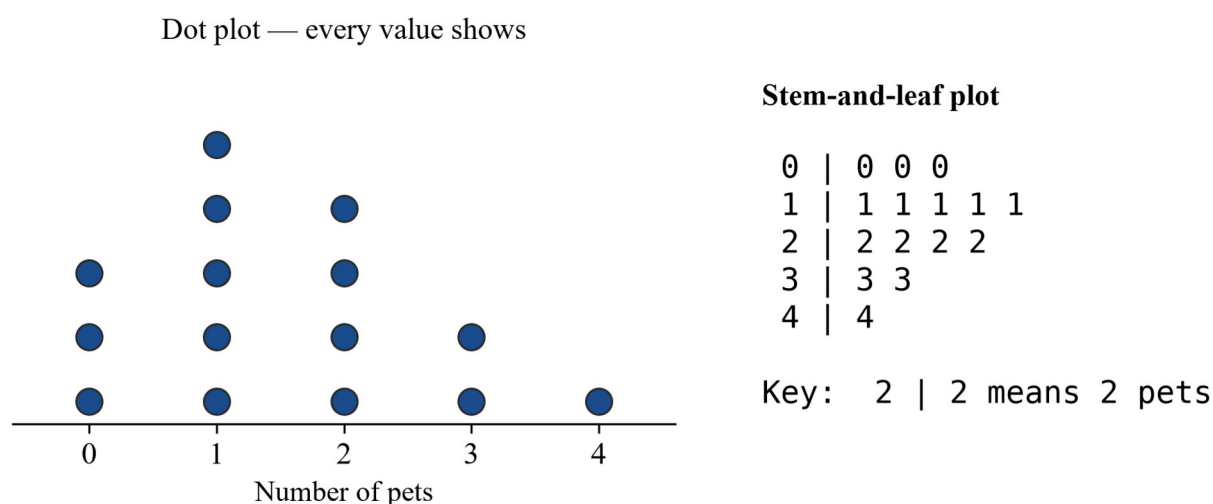


A bar chart of thirty students’ ratings of a statement, with the five categories kept in order from *strongly disagree* to *strongly agree*.

- What is the data type of the ratings? Why must the categories be kept in this order?
- Explain why an ordered bar chart is a better choice here than a pie chart.
- What fraction of the students agreed or strongly agreed? Give your answer as a percentage correct to one decimal place.

Q4. SCENARIO. Fifteen students were asked how many pets live in their household. The raw data is: 2, 0, 1, 3, 0, 1, 1, 2, 4, 0, 2, 1, 3, 2, 1. The figure shows the same data as a dot plot and as a stem-and-leaf plot.

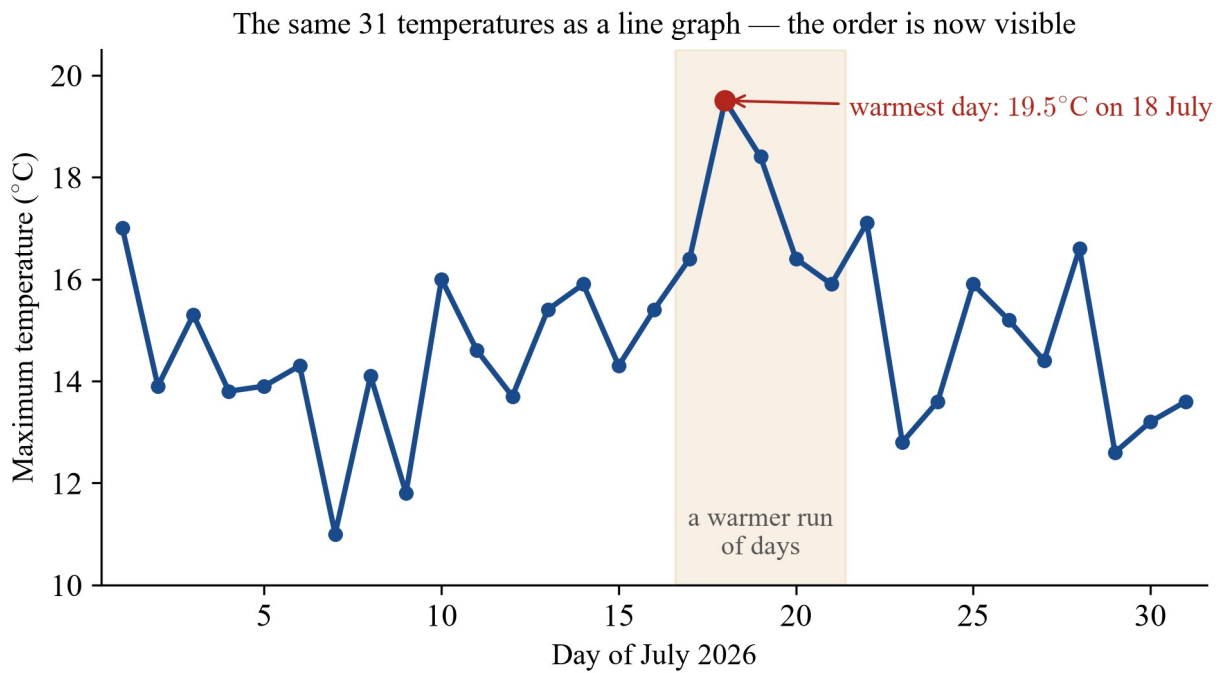
Number of pets for 15 students — two displays of the same data



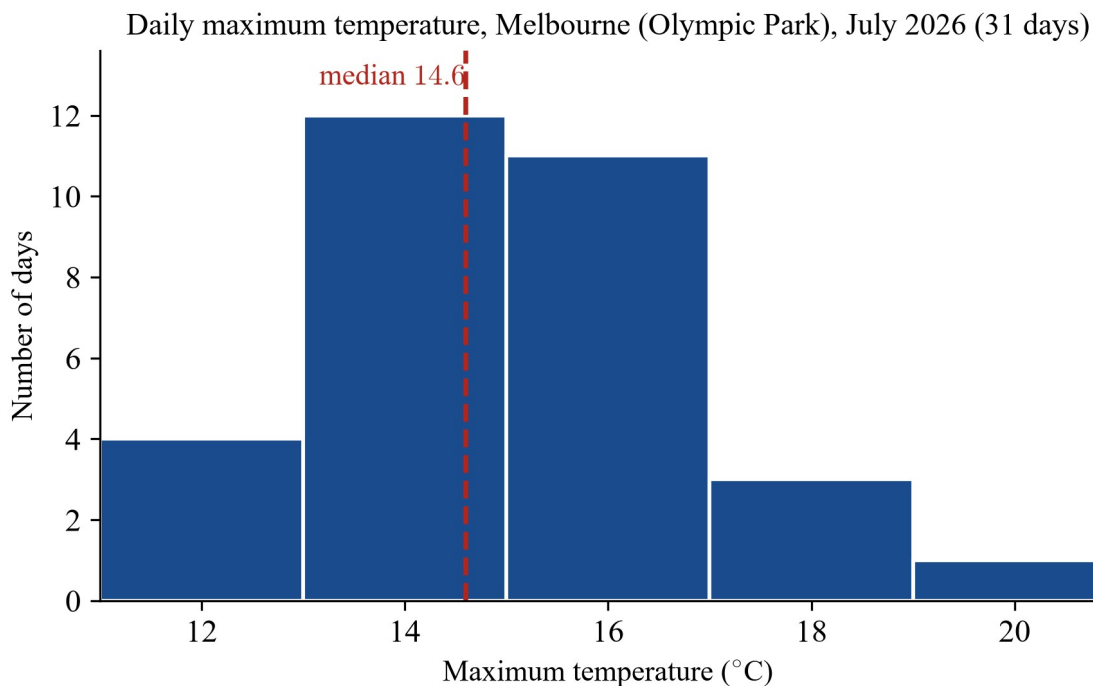
The same fifteen household-pet counts drawn two ways: a dot plot above, and a stem-and-leaf plot with a key below it.

- What is the data type of “number of pets”? Give your reason.
- Read the median and the range from either display.
- Give two reasons why a dot plot or a stem-and-leaf plot is a good choice for these 15 values.

Q5. The July 2026 maximum temperatures from Example 2 are drawn two ways: as a line graph and as a histogram.



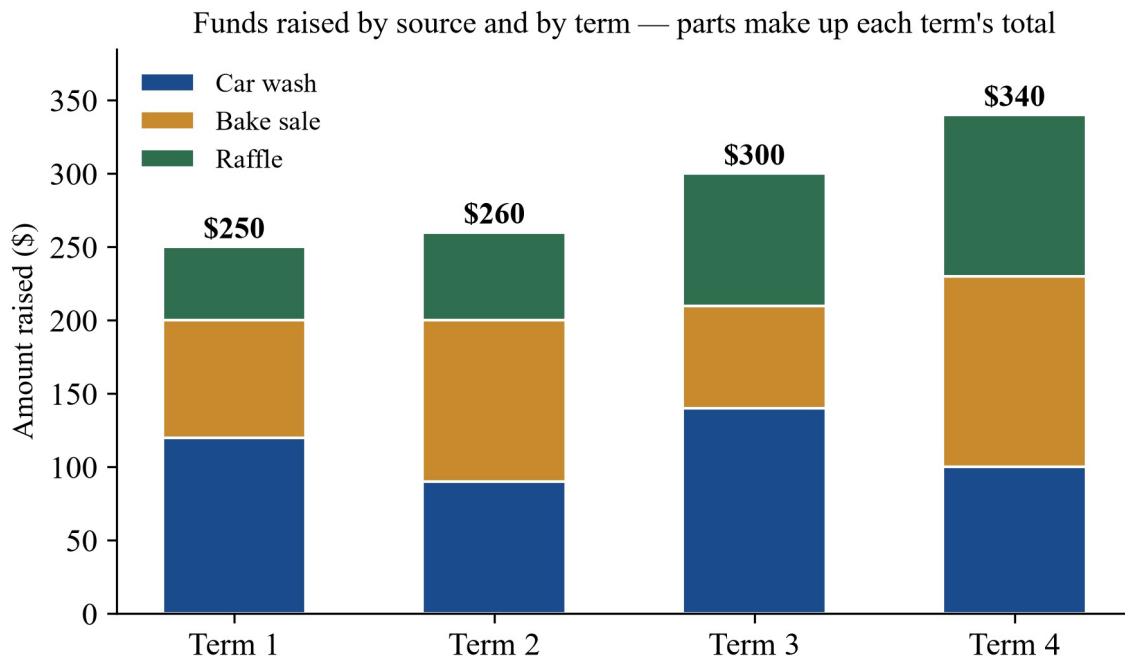
A line graph of the 31 July 2026 daily maximum temperatures at Melbourne (Olympic Park) in the order they were recorded, with the warmest day, 19.5 degrees on 18 July, marked.



A histogram of the 31 July 2026 daily maximum temperatures at Melbourne (Olympic Park), grouped into 2-degree intervals from 11 to 21, with the median of 14.6 degrees marked.

- What is the data type of daily maximum temperature?
- Write down two things the line graph shows that the histogram hides.
- Write down two things the histogram shows that the line graph hides.

Q6. The figure shows the fundraising data from Example 4 as a stacked bar chart.



A stacked bar chart of a school's fundraising across four terms, each bar split into car wash, bake sale and raffle, with the term totals 250, 260, 300 and 340 dollars marked above the bars.

- Read off the total raised in each term.
- Which term raised the most? Which source contributed the most to that term's total?
- Name the two comparisons this display supports.

Unscaffolded

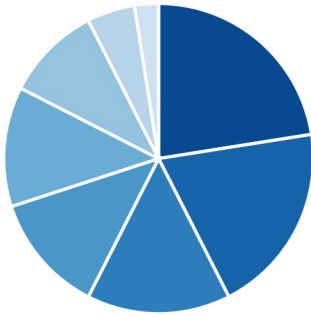
Q7. For each collection of data, choose a display and give a one-line justification that names the data type.

- The eye colours of 25 students.
- The breed of dog owned by each student in a Year-9 class.
- The height of one seedling, measured once a week for eight weeks.
- The marks out of 50 for a class of 28 students.
- The daily minimum temperatures at a weather station across a whole year.
- The medals (gold, silver, bronze) won by each house at the swimming carnival.

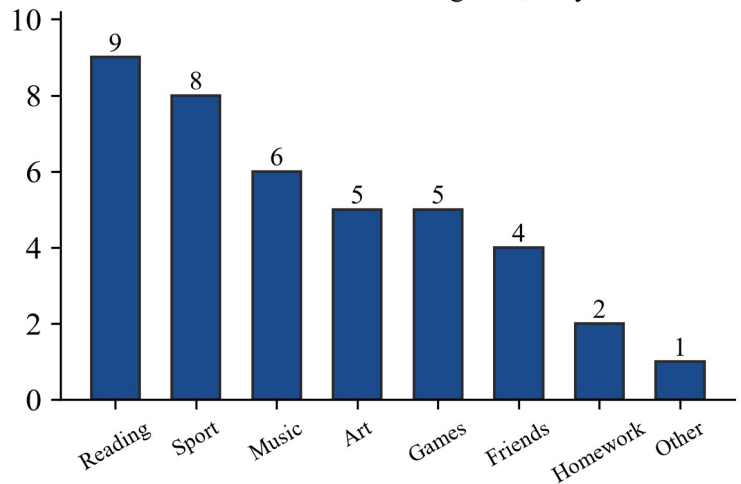
Q8. SCENARIO. Forty students were asked what they usually do at lunchtime. The figure shows the same counts as a pie chart and as a bar chart.

Lunchtime activities of 40 students — same data, two displays

Pie chart — 8 slices are hard to rank



Bar chart — the same 8 categories, easy to rank



The same forty lunchtime-activity counts drawn two ways: a pie chart on the left, and a bar chart sorted from largest to smallest on the right.

- Which of the two displays makes it easier to compare the sizes of the categories? Give your reason.
- Name the most common activity, and its percentage of the 40 students.
- Give one situation in which a pie chart would still be the better choice.

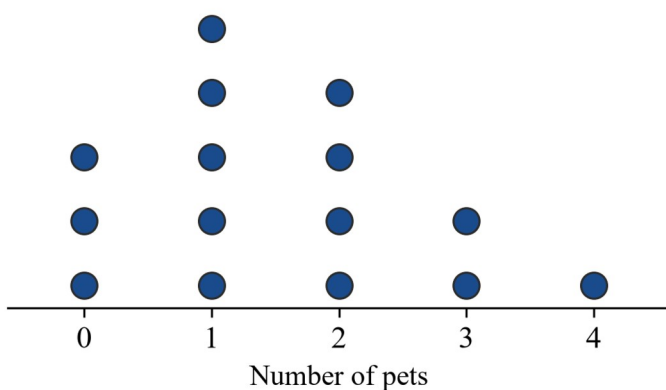
Q9. SCENARIO. After sport, twelve students measured their pulse rates in beats per minute: 62, 68, 70, 72, 74, 75, 76, 78, 80, 82, 86, 90.

- Construct a stem-and-leaf plot for these values, and include a key.
- Justify the choice of a stem-and-leaf plot for this data set.
- Find the mean, the median and the range. Does the data show a slight positive skew? Give your reason.

Q10. The pet data from Question 4 is shown in the figure as a dot plot and as a stem-and-leaf plot.

Number of pets for 15 students — two displays of the same data

Dot plot — every value shows



Stem-and-leaf plot

```

0 | 0 0 0
1 | 1 1 1 1 1
2 | 2 2 2 2
3 | 3 3
4 | 4
    
```

Key: 2 | 2 means 2 pets

The same fifteen household-pet counts drawn two ways: a dot plot above, and a stem-and-leaf plot with a key below it.

- Calculate the mean, the median and the range of the 15 values.

- Both displays show every raw value. Give one advantage of the dot plot and one advantage of the stem-and-leaf plot.
- If the same question had been asked of 200 students instead of 15, would you still choose a dot plot or a stem-and-leaf plot? Choose a display and justify your choice.

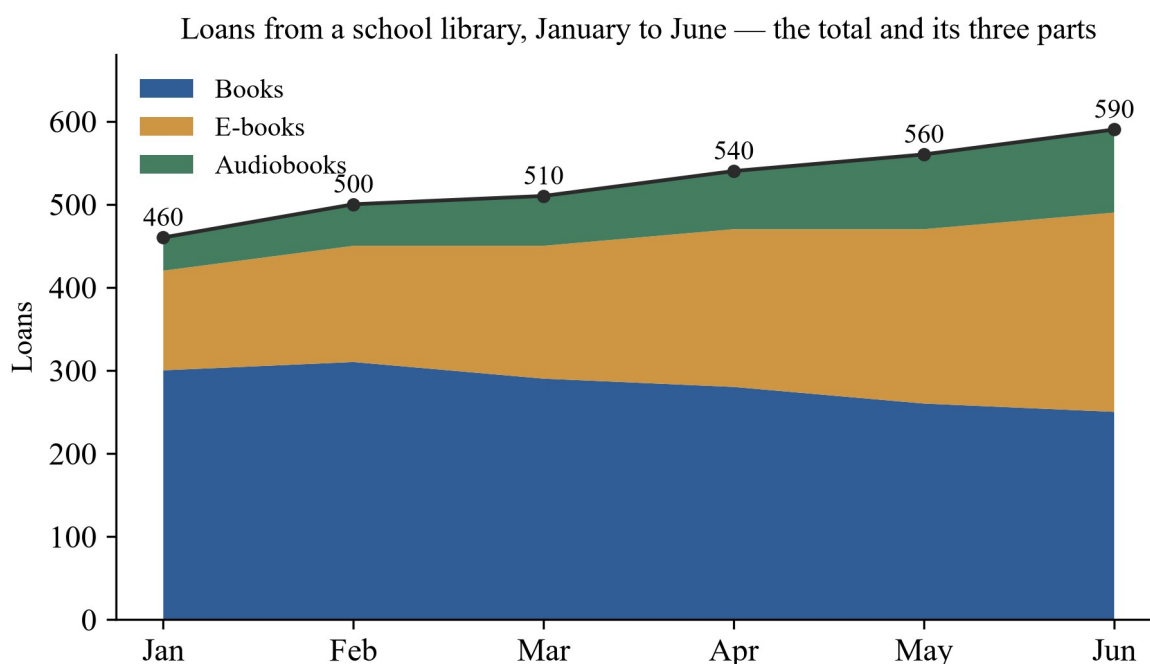
Q11. SCENARIO. A small shop records its total takings, in dollars, over eight weeks: 420, 460, 440, 510, 530, 490, 560, 600.

- Are the weekly takings categorical or numerical?
- Justify the choice of a line graph for these values.
- Find the overall change from week 1 to week 8, and the mean weekly takings. What does the mean hide?

Q12. Use the fundraising stacked bar chart from Question 6.

- Which source raised the most across the whole year? Give the amount.
- Find the car wash's share of the grand total, as a fraction in lowest terms and as a percentage correct to one decimal place.
- The raffle's four term totals are 50, 60, 90, 110. Write down one story about the raffle that the stacked bar chart tells and the four term totals alone do not.

Q13. SCENARIO. A library records its loans each month from January to June, split by format: books, e-books and audiobooks. The figure shows the data as a stacked area chart, with the total loans drawn as a line.

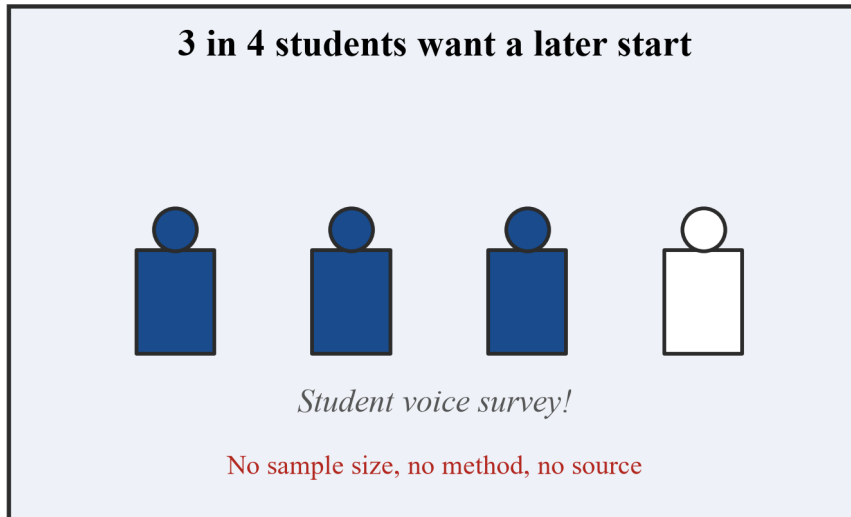


A stacked area chart of a library's monthly loans from January to June, split into books, e-books and audiobooks, with the monthly total drawn as a line above the bands.

- Read off the total loans in January and in June. What is the change across the six months?
- Which format grew the most over the six months, and by how much? Which format shrank, and by how much?
- Justify the choice of an area chart here rather than three separate line graphs.
- By June, had e-books overtaken books? Give the two numbers.

Q14. SCENARIO. A student council poster is shown in the figure.

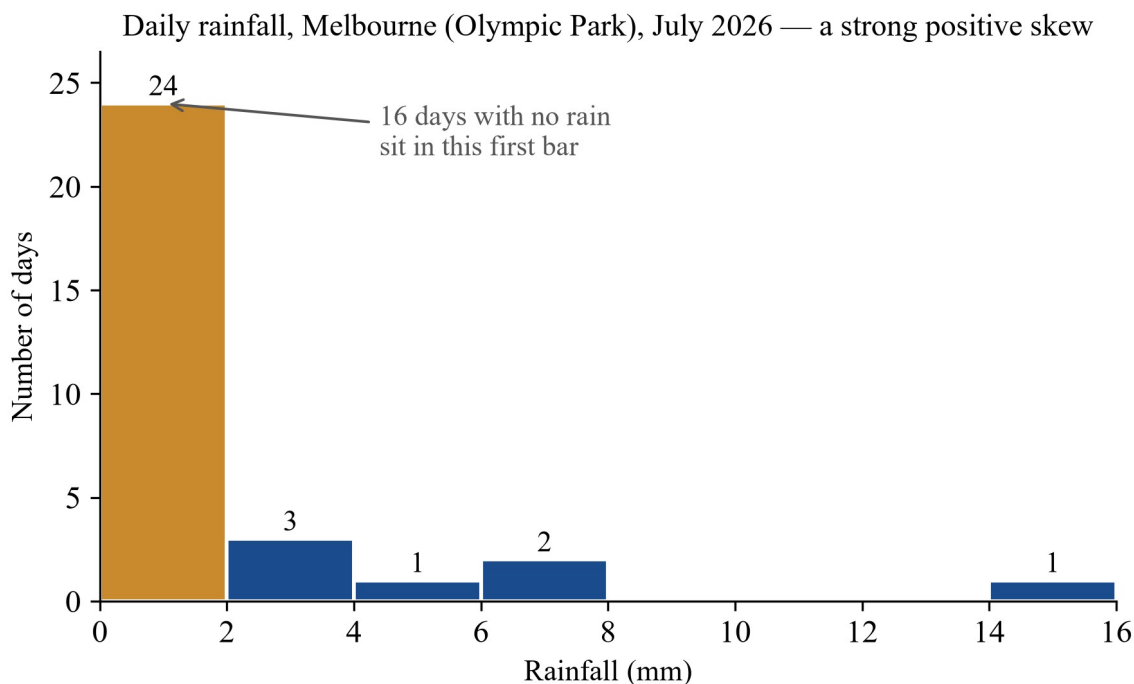
A poster-style display (SCENARIO) — made to persuade



An infographic poster reading “3 in 4 students want a later start”, with four person icons of which three are shaded, and a red line underneath noting that no sample size, no method and no source are given.

- Write “3 in 4” as a percentage.
- What is the purpose of the poster, and who is its intended audience?
- If the survey had asked 25 students, show why “3 in 4” cannot be exact. Name two things the poster should add before the claim can be trusted.

Q15. The July 2026 rainfall from Example 3 is shown in the histogram.

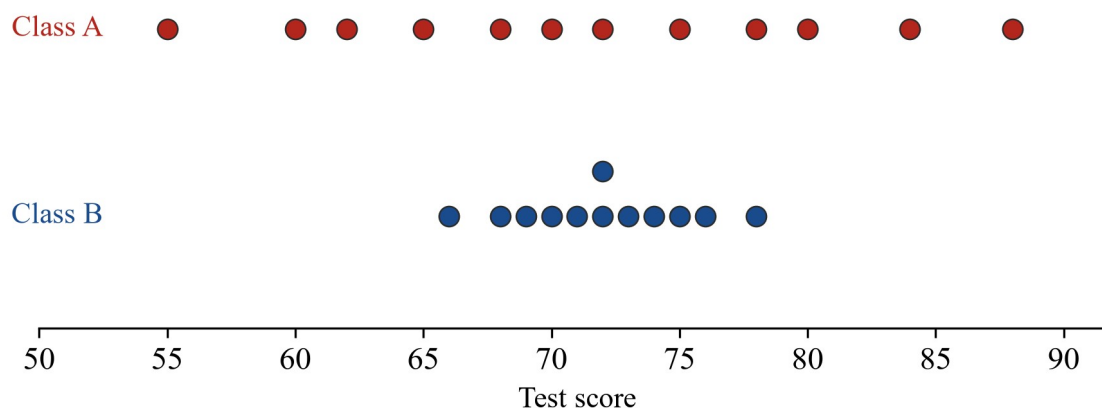


A histogram of the 31 July 2026 daily rainfalls at Melbourne (Olympic Park), grouped into 2-millimetre intervals from 0 to 16, showing a tall first bar of 24 days and a note that 16 days had no rain.

- How many days had less than 2 mm of rain?
- The mean is about 1.43 mm and the median is 0 mm. Which measure better describes a typical July day? Give your reason.
- Name the shape of this distribution, and explain why quoting the mean alone would mislead a reader.

Q16. SCENARIO. Two classes of 12 students each sit the same quiz. The figure shows both sets of scores as dot plots on a shared axis.

Twelve scores in each class — same centre, different spread



Two dot plots on a shared horizontal axis from 50 to 90, one above the other, showing class A's scores spread wide and class B's scores clustered tightly.

- Find the median and the range for each class.
- Which class's results are more consistent? Give your reason.
- Justify the choice of two dot plots on a shared axis for comparing the two classes.

Q17. SCENARIO. One hundred and twenty students were asked how they travel to school. The counts are: bus 34, car 30, walking 26, train 14, cycling 10, other 6.

- What is the data type of "method of travel to school"?
- Choose a display and justify your choice. Explain why a histogram would not be appropriate.
- What percentage of the students travel by bus? Give one reason why a pie chart is a weaker choice than yours for these data.

Q18. SCENARIO. A school newsletter reports a survey of favourite sports in Year 9. The counts are: soccer 9, basketball 7, netball 6, cricket 5, swimming 3. The newsletter draws them as a line graph, with the points joined by line segments.

- What does joining the points suggest about the sports that is not true?
- Name the correct display for these data.
- Justify your choice using the checklist: name the data type, the question the display must answer, and why your display answers it.

Answers

Q1 a Nominal — labels with no natural order; b ordinal — labels with a natural order; c discrete — counted whole numbers; d continuous — measured, any value in a range is possible; e ordinal — labels with a natural order; f discrete — counted whole numbers. **Q2** a Nominal — the sports are labels with no order; b the sports have no order, and joining points invents one; c $\frac{9}{30} = 30\%$. **Q3** a Ordinal — the order of the ratings is part of the data; b a pie chart shows shares of a whole but loses the order of the ratings; c $\frac{14}{30} \approx 46.7\%$. **Q4** a Discrete — pets are counted in whole numbers; b median 1 pet, range 4; c there are only 15 values, and every raw value can stay on show; clusters and gaps are visible at a glance. **Q5** a Continuous; b e.g. the warm spell from the 17th to the 21st, the warmest day (19.5°C on 18 July), the day-to-day rises and falls; c e.g. the shape of the distribution (most days between 13 and 17°C), and the centre and spread of the month. **Q6** a \$ 250, \$ 260, \$ 300, \$ 340; b Term 4; the bake sale (\$ 130); c totals across terms, and the parts inside each term. **Q7** a Bar chart — eye colour is nominal; b bar chart — breed is nominal; c line graph — height is continuous and measured in order over time; d stem-and-leaf plot or dot plot — 28 discrete values is a small set; e histogram — 365 continuous values need grouping; f ordered bar chart — the medals are ordinal, so the order must be kept. **Q8** a The bar chart — the eye compares bar lengths more easily than pie angles; b reading, $\frac{9}{40} = 22.5\%$; c e.g. when the message is a share of a whole and there are only a few categories. **Q9** a Stems 6–9: 6|28; 7|024568; 8|026; 9|0; key 7|5 = 75 beats per minute; b 12 values is a small set, and the stem-and-leaf plot keeps every value sorted and on show; c mean $\frac{913}{12} \approx 76.1$, median 75.5, range 28 — yes, a slight positive skew, since the mean sits just above the median. **Q10** a Mean $\frac{23}{15} \approx 1.53$ pets, median 1 pet, range 4; b the dot plot shows the build-up at each value at a glance; the stem-and-leaf plot keeps the exact values, sorted; c neither — 200 values need grouping, so a histogram is the better choice. **Q11** a Numerical; b the values are recorded in order over time and the question is about the trend; c change \$ 180 (from \$ 420 to \$ 600); mean \$ 501.25 — the mean mixes the low early weeks with the high late weeks and hides the rise. **Q12** a Car wash, \$ 450; b $\frac{450}{1150} = \frac{9}{23} \approx 39.1\%$; c the raffle grew every term — a rising trend inside one part of the display. **Q13** a 460 and 590 loans; up 130; b e-books, up 120 (from 120 to 240); books, down 50 (from 300 to 250); c the area chart shows the total's trend and the changing mix in one display; three separate line graphs would show the parts but not the total; d no — 240 e-books against 250 books. **Q14** a 75%; b to persuade the school community to support a later start; c 75% of 25 is 18.75 — not a whole number of students, so the fraction cannot be exact; the poster should add the sample size and the method (who was asked, and how). **Q15** a 24 days; b the median — half the days had 0 mm, so 0 mm is the typical day; c positive skew — the mean is pulled up by a few wet days, while a typical day is dry. **Q16** a Class A: median 71, range 33; class B: median 72, range 12; b class B — the smaller range means the scores sit closer together; c the shared axis lets the centres and spreads of the two classes be compared at a glance. **Q17** a Nominal; b a bar chart (sorted by size) — the values are labels and the question asks how the counts compare; a histogram groups numerical values into intervals, and there are no numbers to group; c $\frac{34}{120} \approx 28.3\%$; six pie slices are hard to compare by eye, and the counts themselves are what matter, not shares of a whole. **Q18** a An order, and change from one sport to the next — but the sports are labels with no order; b a bar chart; c the data is nominal, the question is “how do the counts compare?”, and a bar chart compares the counts without inventing an order.

Fully-worked solutions

Q1. a The house colours are labels with no order — nominal. b Gold, silver and bronze are labels, but they rank — ordinal. c Goals are counted: 0, 1, 2, ... — discrete. d Heights are measured, and any value in a range is possible — continuous. e Poor, fair and good are labels that rank — ordinal. f Wet days are counted in whole numbers — discrete. $\therefore$ The sorting rule: label or number first, then the order test for labels, the count-or-measure test for numbers.

Q2. a The sports are labels with no natural order — nominal. b A line graph joins points in order; the sports have no order, so any joined line would be invented. c $\frac{9}{30} = \frac{3}{10} = 30\%$. $\therefore$ Nominal data with counts to compare: bar chart; no line segments, because there is no order to join.

Q3. a The ratings are ordinal: they are labels, and their order (strongly disagree through to strongly agree) is part of the data. b A pie chart shows shares of a whole; it cannot keep the categories in order, so the pattern — how opinions climb from disagreement to agreement — would be lost. c Agree or strongly agree: $9 + 5 = 14$ of 30: $\frac{14}{30} \approx 46.7\%$. $\therefore$ Ordinal data needs a display that keeps the order — an ordered bar chart.

Q4. a The number of pets is counted in whole numbers — discrete numerical. b From the dot plot (or the stem-and-leaf plot): the middle of the 15 values is 1, so the median is 1 pet; the range is $4 - 0 = 4$. c There are only 15 values, so every raw value can stay on show; and both displays show clusters and gaps at a glance — here the build-up at 1 pet. $\therefore$ Small discrete data sets suit displays that keep every value: dot plots or stem-and-leaf plots.

Q5. a Daily maximum temperature is measured — continuous numerical. b The line graph keeps the order of the days, so it shows the day-to-day rises and falls, the warm spell from the 17th to the 21st (all at least 15.9°C), and that the warmest day was 19.5°C on 18 July. c The histogram groups the values, so it shows the shape of the month — most days between 13 and 17°C — and with it the centre (median 14.6°C) and the spread (range 8.5°C). $\therefore$ Same data, different question: “when?” asks for the line graph, “what was the month like?” asks for the histogram.

Q6. a Bar heights: Term 1 \$250, Term 2 \$260, Term 3 \$300, Term 4 \$340. b Term 4 raised the most (\$340); inside Term 4, the bake sale contributed the most (\$130, against \$110 raffle and \$100 car wash). c The bar heights compare the totals across terms; the segments compare the parts within each term. $\therefore$ One stacked bar chart, two comparisons: totals across categories and parts within categories.

Q7. a Bar chart — eye colour is nominal: labels with no order, and the question asks how the counts compare. b Bar chart — breed is nominal for the same reason. c Line graph — the height of one seedling is continuous and measured in order over time, and the question is about growth. d Stem-and-leaf plot or dot plot — 28 discrete marks is a small set, so every value can stay on show. e Histogram — 365 continuous values need grouping to show the shape of the year’s temperatures. f Ordered bar chart — the medals are ordinal, so the display must keep gold, silver, bronze in order. $\therefore$ In every case the data type narrows the choice, and the question picks inside it.

Q8. a The bar chart. Comparing counts means comparing lengths, and the eye reads lengths far better than it reads pie angles — with eight categories, several of the slices are close in size and hard to rank. b Reading is the most common: $\frac{9}{40} = 22.5\%$. c When the message is a share of a whole and there are few categories — e.g. “half of our energy comes from one source” with two or three slices. $\therefore$ Counts to compare: bar chart. A share of a whole with few parts: pie chart.

Q9. a Stems in tens, leaves in units:

```
6 | 2 8
7 | 0 2 4 5 6 8
8 | 0 2 6
9 | 0
```

Key: $7|5 = 75$ beats per minute. b There are only 12 values, and the stem-and-leaf plot keeps every one of them, sorted, while still showing the shape — a build-up in the 70s. c Mean $\frac{913}{12} \approx 76.1$; median $\frac{75+76}{2} = 75.5$; range $90 - 62 = 28$. The mean sits just above the median, with one high value (90) pulling at it — a slight positive skew. $\therefore$ Small numerical data set with shape to show: stem-and-leaf plot.

Q10. a Sum $2 + 0 + \dots + 1 = 23$, so the mean is $\frac{23}{15} \approx 1.53$ pets. The 8th of the 15 ordered values is 1, so the median is 1 pet. The range is $4 - 0 = 4$. b The dot plot shows how the values pile up at each number — the cluster at 1 pet — at a glance. The stem-and-leaf plot keeps the exact values, sorted, so nothing is lost. c With 200 values, both displays

would pile up past reading: 200 dots or 200 leaves. Grouping is needed, so the better choice is a histogram — discrete or continuous, large data sets are grouped to show the shape. $\therefore$ Display choice depends on size as well as type: 15 values keep every raw value on show; 200 values need a histogram.

Q11. a The takings are dollar amounts — numerical. b The values are recorded in order over time, and the question is about how sales move week to week — a trend is a job for a line graph. c The change is $600 - 420 = \$180$. The mean is $\frac{4010}{8} = \$501.25$. The mean blends the low early weeks with the high late weeks, so it hides the rise — two shops with the same mean can have flat sales or rising sales. $\therefore$ Order matters here, so the line graph earns its place; the mean alone would not tell the story.

Q12. a Across the year: car wash $120 + 90 + 140 + 100 = \$450$; bake sale $80 + 110 + 70 + 130 = \$390$; raffle $50 + 60 + 90 + 110 = \$310$. The car wash raised the most. b The grand total is $\$1150$, so the car wash's share is $\frac{450}{1150} = \frac{9}{23} \approx 39.1\%$. c The raffle grew every term — 50, 60, 90, 110. The term totals alone say nothing about any source; the segments tell the raffle's rising story. $\therefore$ The stacked bar chart does more than total: it shows the trend hiding inside a part.

Q13. a The total line sits at 460 in January and 590 in June: up 130 loans. b E-books grew the most: $240 - 120 = 120$ loans. Books shrank: $300 - 250 = 50$ loans. Audiobooks grew $100 - 40 = 60$. c The question asks about the total *and* about the formats inside it, across time. The area chart shows both: the top edge tracks the total, and the bands track each format. Three separate line graphs would show the formats but not the total; a single line of the total would hide the formats. d No: in June, 240 e-books against 250 books — close, but not yet past. $\therefore$ Parts over time, with the total in the question: stacked area chart.

Q14. a “3 in 4” is $\frac{3}{4} = 75\%$. b The purpose is to persuade: the council wants the school to back a later start. The intended audience is the school community — students, staff and families who read the poster. c If 25 students were asked, “3 in 4” would mean $\frac{3}{4} \times 25 = 18.75$ students — not a whole number, so the fraction cannot be exact for that sample size. The poster should add the sample size (how many students were asked) and the method (who was asked, and how they were chosen); a source would let a reader check the claim. $\therefore$ The number may be right; without a sample size and a method, no reader can check it — and the check above shows it cannot be exact for 25 students.

Q15. a The first interval, 0–2 mm, holds 24 days. b The median. Sixteen of the thirty-one days had 0 mm, so half the month had no rain at all: 0 mm is the typical day. The mean of about 1.43 mm is dragged up by a few wet days. c The distribution is positively skewed — a tall pile of dry days with a long tail towards the wet ones. Quoting the mean alone would make July sound wetter than it feels, because one day of 14.6 mm does more work on the mean than sixteen dry days can undo. $\therefore$ Skewed data needs its shape shown — the histogram — and the median to describe the typical value.

Q16. a Class A: the middle of the 12 ordered scores is $\frac{70+72}{2} = 71$, and the range is $88 - 55 = 33$. Class B: the middle is $\frac{72+72}{2} = 72$, and the range is $78 - 66 = 12$. b Class B. The medians are almost the same (71 against 72), but class B's range is far smaller, so its scores sit much closer together. c The shared axis lines the two groups up, so the centres and the spreads can be compared at a glance: class A's dots stretch wide, class B's cluster. (A back-to-back stem-and-leaf plot, from 6A, would do the same job.) $\therefore$ Comparing two groups: same axis, side by side — the comparison is the point of the display.

Q17. a Method of travel is a label with no order — nominal. b A bar chart, with the bars sorted from largest to smallest: the values are labels, and the question asks how the counts compare. A histogram would not be appropriate because a histogram groups numerical values into intervals — there are no numbers here to group. c $\frac{34}{120} \approx 28.3\%$ travel by bus. A pie chart is a weaker choice: the six slices are close in size and hard to rank by eye, and the question is about the counts themselves, not about shares of a whole. $\therefore$ Nominal data, counts to compare: sorted bar chart.

Q18. a Joining the points suggests the sports come in an order, and that something changes from soccer to basketball to netball. The sports are labels with no order — soccer is not “before” basketball — so the line segments say something false. b A bar chart. c Data type: nominal — labels with no order. Question: how do the counts compare across the sports? Display: a bar chart compares the counts with one bar per sport and invents no order, because the bars stand apart. $\therefore$ A line graph implies order and change; nominal data has neither, so the display must not draw either.

Reading survey reports: how the data was obtained, and estimating population means and medians

6C · Chapter 6 — Statistics · Level 9

Content description VC2M9ST01 (word for word): *analyse reports of surveys in digital media and elsewhere for information on how data was obtained around everyday questions and issues involving at least one numerical and at least one categorical variable, to estimate population means and medians*

Achievement standard anchor (AS 17): *They explain how sampling techniques and representation can be used to support or question conclusions or to promote a point of view.*

Elaboration ids for linkage: VC2M9ST01.E01, VC2M9ST01.E02, VC2M9ST01.E03, VC2M9ST01.E04, VC2M9ST01.E05

Prerequisite pointer: Level 8 Statistics is assumed — population vs sample and collection techniques (VC2M8ST01) and investigations with samples and fair inference (VC2M8ST04) — together with 6A (centre, spread, shape and outliers: mean, median and range only, no quartiles and no box plots, D3.10) and 6B (choosing a display). The mean-and-median tools are used here, not re-taught; the new work is reading *someone else's* report and using a sample to estimate a population value.

Notes on the elaborations: the section's skill work is delivered on two real, named, dated, retrievable reports — the ABS *Average Weekly Earnings* release (E01: linking claims to statistics and representative data) and the BOM *Victoria in 2025* annual rainfall report (E05's rainfall context, on Australian stations; the international comparison of E05 is not taken up, as the sourced BOM dataset is Victorian) — and on clearly-labelled SCENARIO opinion-poll reports in the E04 style (single-use plastics), discussing methods of data collection and the reasonableness of conclusions (E01, E04). The example contexts of E02 (Asia-region trade growth) and E03 (ages of residents across named countries) are not taken up; per D5, elaborations illustrate a content description, they are not the content description itself.

Notes on sources: real figures carry a source line (R1). Invented parameters appear only inside labelled SCENARIO items (R2). No Aboriginal or Torres Strait Islander content appears in this section (CONTENT_POLICY Rule 1).

Theory

A report is a claim with data behind it

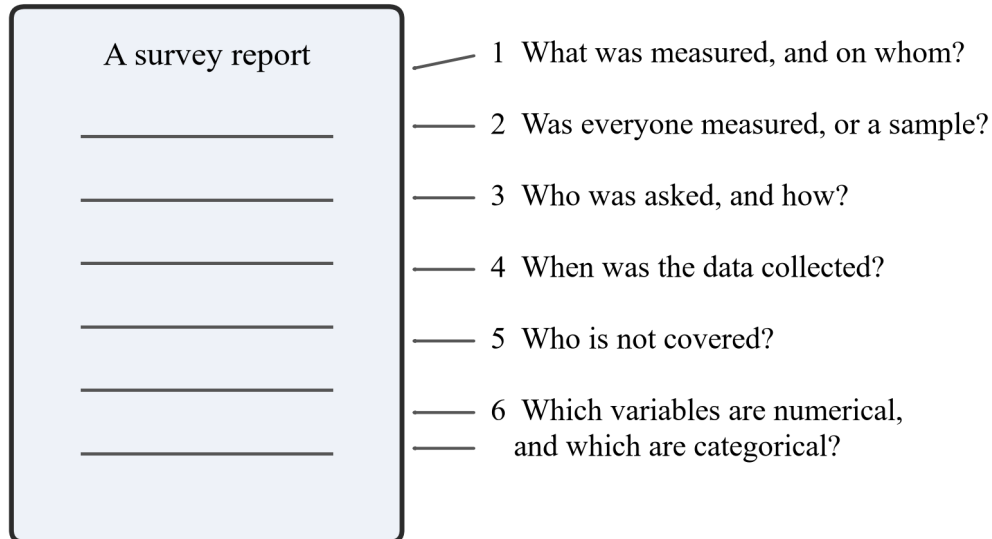
News sites, advertisers and governments publish survey reports every week: *Australians earn \$2,160 a week on average, most shoppers back a plastics ban, 2025 was a dry year in Victoria*. The content description's verb is *analyse*: take the report apart, ask how the data was obtained, and check whether the numbers support the claim. Some data is obtained by **asking** — a survey — and some by **measuring** — a weather station reading a rain gauge each day. Both arrive as reports, and both get read the same way.

Review from Level 8: the **population** is the whole group a report wants to speak about; a **sample** is the part the data actually comes from; a **census** measures every member of the population. Every report that uses a sample is publishing an **estimate** of a population value, and an estimate is close to the truth but is not the truth.

The reading checklist

Six questions unlock almost every report. They are the working tool of this subtopic, and every worked example and question below runs some of them.

Read a survey report like a statistician: six questions

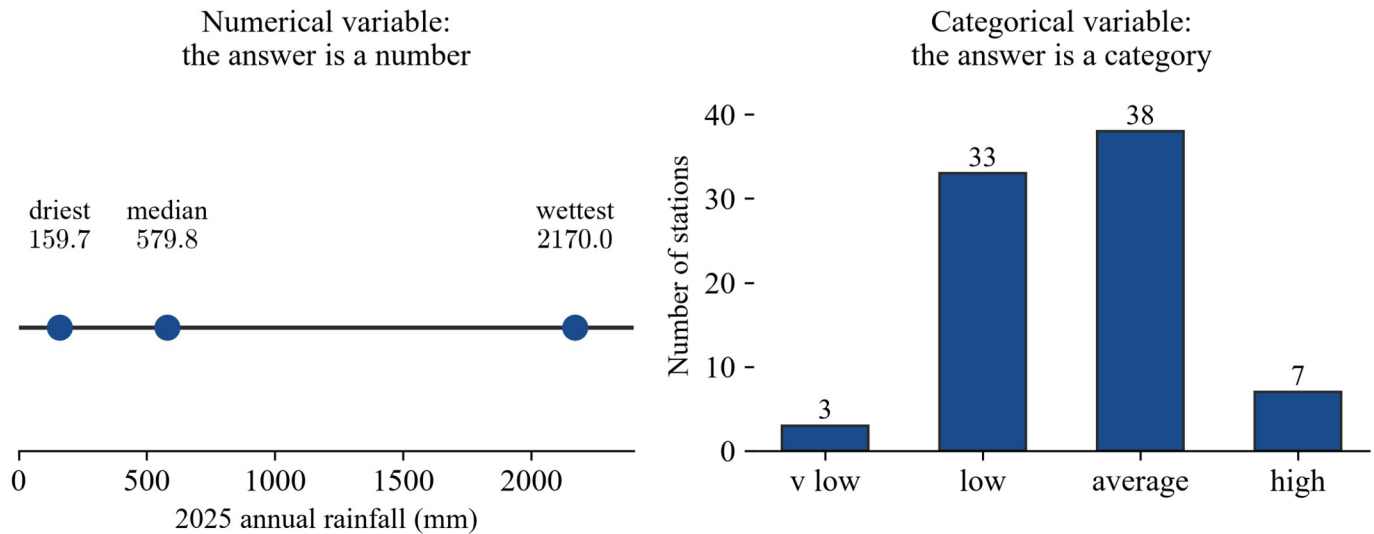


Schematic of a survey report with six reader questions: what was measured and on whom; everyone or a sample; who was asked and how; when; who is not covered; which variables are numerical and which categorical.

1. What was measured, and on whom?
2. Was everyone measured, or a sample?
3. Who was asked, and how?
4. When was the data collected?
5. Who is not covered?
6. Which variables are numerical, and which are categorical?

A **variable** is anything the report records for each person, place or thing in the data. A **numerical** variable's values are numbers you could sensibly average — rainfall in millimetres, earnings in dollars. A **categorical** variable's values are labels or categories — sex, state, or a ranking such as *low / average / high*. The content description expects a report to carry at least one of each, and a careful reader names both before trusting anything else.

Two kinds of variable in one report — BOM, Victoria in 2025, 81 stations



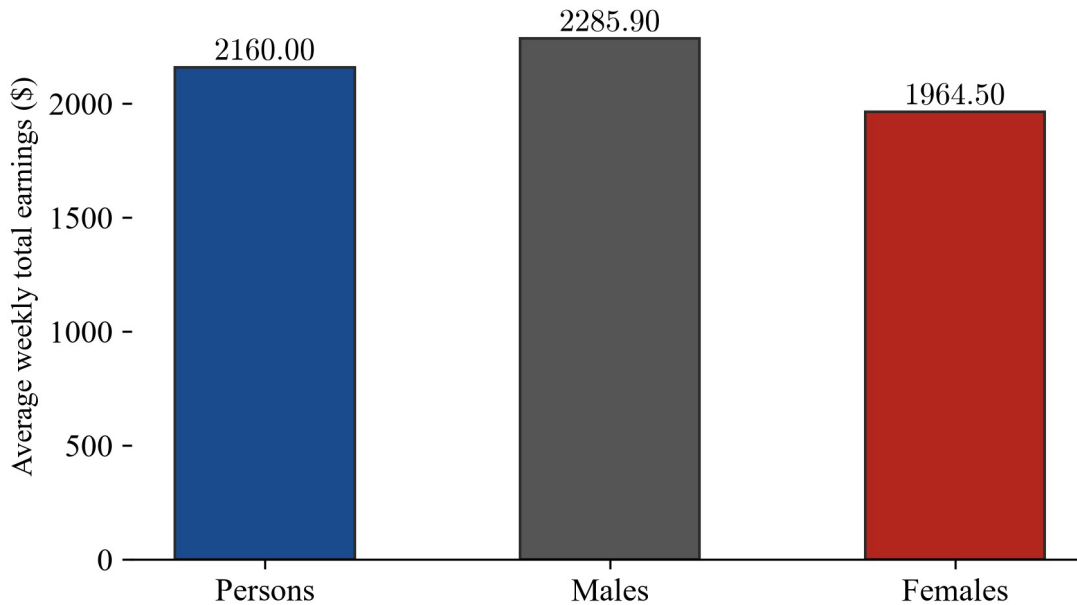
Two panels from the BOM Victoria in 2025 report: a number line of 2025 annual rainfall showing the driest station (159.7 mm), the median (579.8 mm) and the wettest (2170.0 mm), beside a bar graph of the categorical decile ranks (3 very low, 33 low, 38 average, 7 high).

Real report 1 — ABS, Average Weekly Earnings

The Australian Bureau of Statistics (ABS) publishes *Average Weekly Earnings* twice a year. Source: Australian Bureau of Statistics, *Average Weekly Earnings, Australia, May 2026*, 2026. Retrieved 20 August 2026. The companion methodology page records how the data was obtained: it is a survey of **employing businesses**, not of workers — businesses report, through an online form, the total pay and the number of employees paid in the reference week (the last pay period ending on or before the third Friday of May). The businesses are a probability sample drawn from the ABS business register, stratified by sector (public or private), industry, state or territory and employment size, and the survey aims for a response rate of 90 %. The scope is full-time adult employees, so the survey *leaves out* whole groups: members of the defence forces, workers in agriculture, forestry and fishing, the self-employed, and unpaid directors, among others.

Source: Australian Bureau of Statistics, *Average Weekly Earnings, Australia, May 2026 — methodology*, 2026. Retrieved 20 August 2026.

Average weekly total earnings, full-time adults,
Australia, May 2026 (seasonally adjusted)



Bar graph of average weekly total earnings for full-time adults, Australia, May 2026, seasonally adjusted: 2160.00 dollars for persons, 2285.90 dollars for males and 1964.50 dollars for females.

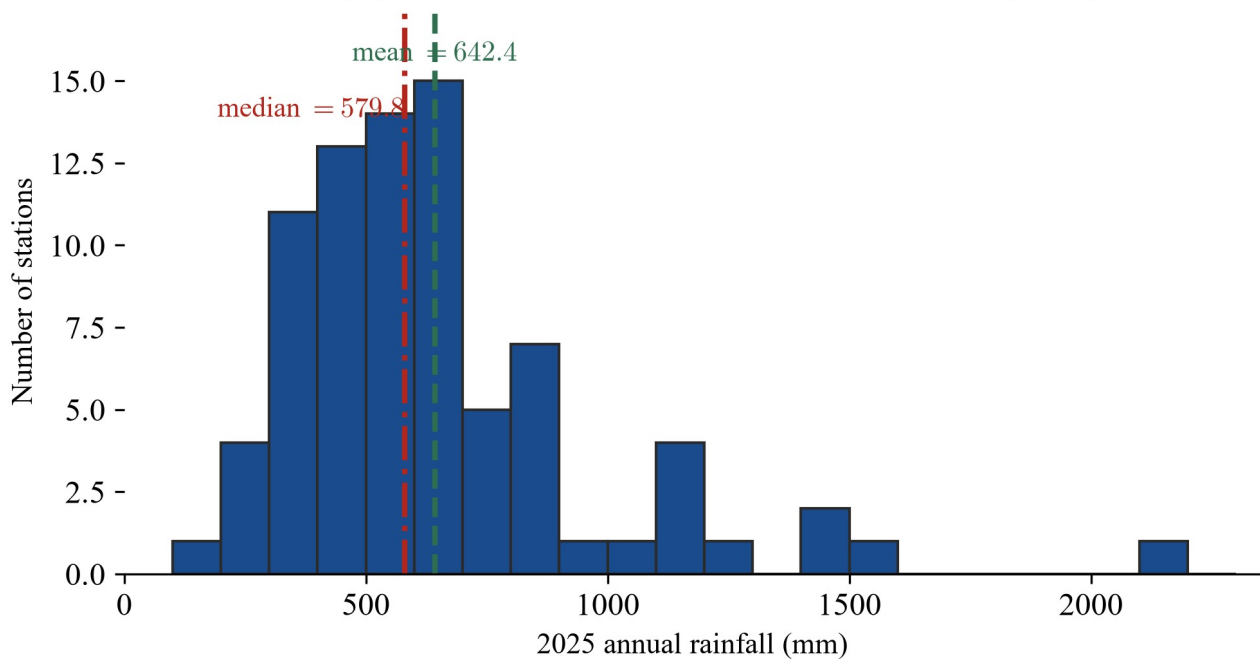
Run the checklist on the release. The **numerical variable** is weekly earnings in dollars; the **categorical variables** are sex (male / female / persons) and state or territory. And the headline number is a **mean**: the ABS divides the total weekly earnings it collected by the number of employees it counted. For full-time adults in May 2026 the seasonally adjusted mean total earnings were \$2,160.00 for persons, \$2,285.90 for males and \$1,964.50 for females. Notice what the release does *not* publish: a median. A mean divided off totals says nothing about the worker in the middle — and earnings data is famously skewed by a few very large salaries, so the median would give a fairer picture of the typical worker. That absence is itself a finding when you analyse the report.

Real report 2 — BOM, Victoria in 2025

The Bureau of Meteorology (BOM) publishes an annual climate summary for each state. Source: Bureau of Meteorology, *Victoria in 2025 (Annual Climate Summary)*, 2026. Retrieved 20 August 2026. Its table lists, for 81 Victorian weather stations with complete records, the annual rainfall for 2025 beside each station's long-term average, plus a **categorical** decile rank: *very low*, *low*, *average* or *high*. This data was obtained by measurement, not by asking: a rain gauge is read at each station every day, and the year's daily readings are added. For the year as a whole, the report states that Victoria's area-averaged rainfall was 544.6 mm, 18% below the 1961–1990 average — 2025 was a dry year.

One station's record shows what those averages sit on. Source: Bureau of Meteorology, *Monthly Climate Statistics for Melbourne Airport (086282)*, 2026. Retrieved 20 August 2026. Over 1970–2026, Melbourne Airport's mean annual rainfall is 536.2 mm; its wettest year was 820.8 mm (1978) and its driest 310.2 mm (1997) — year-to-year variation is large, which is exactly why one year is read against a long-term average.

The population: 2025 annual rainfall at all 81 stations (BOM)



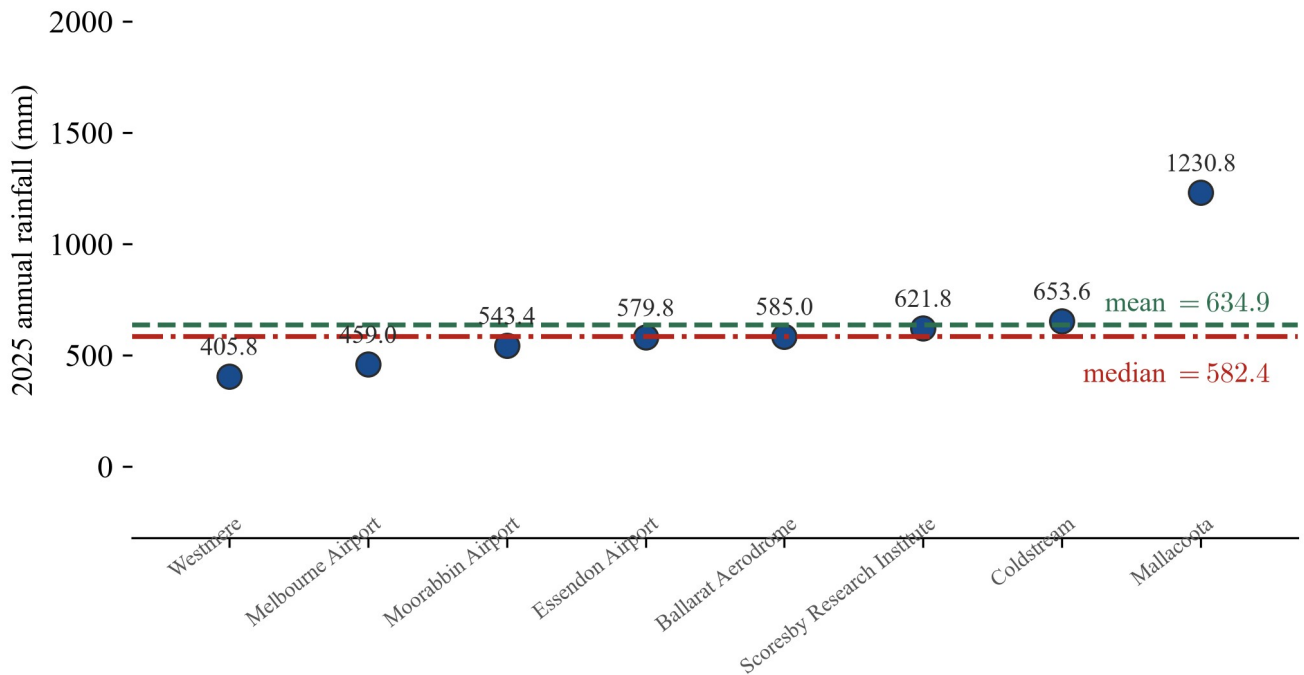
Histogram of 2025 annual rainfall at all 81 Victorian stations, with the mean (642.4 mm) to the right of the median (579.8 mm) and a long tail of wet alpine stations.

Treat the 81 stations as a **population**. Its 2025 rainfall has mean 642.4 mm and median 579.8 mm. The histogram shows the shape: most stations sit between 300 and 900 mm, while a few wet alpine stations — Mount Baw Baw at 2170.0 mm, Mount Hotham at 1506.8 mm — drag the mean above the median. The BOM's own report makes the point plainly: *the median is sometimes more representative than the mean of long-term average rain.*

Estimating a population mean and median from a sample

You will not usually be handed the whole population. The skill the content description asks for is estimating the population mean and median from a sample: compute the sample mean and the sample median, and offer them as estimates.

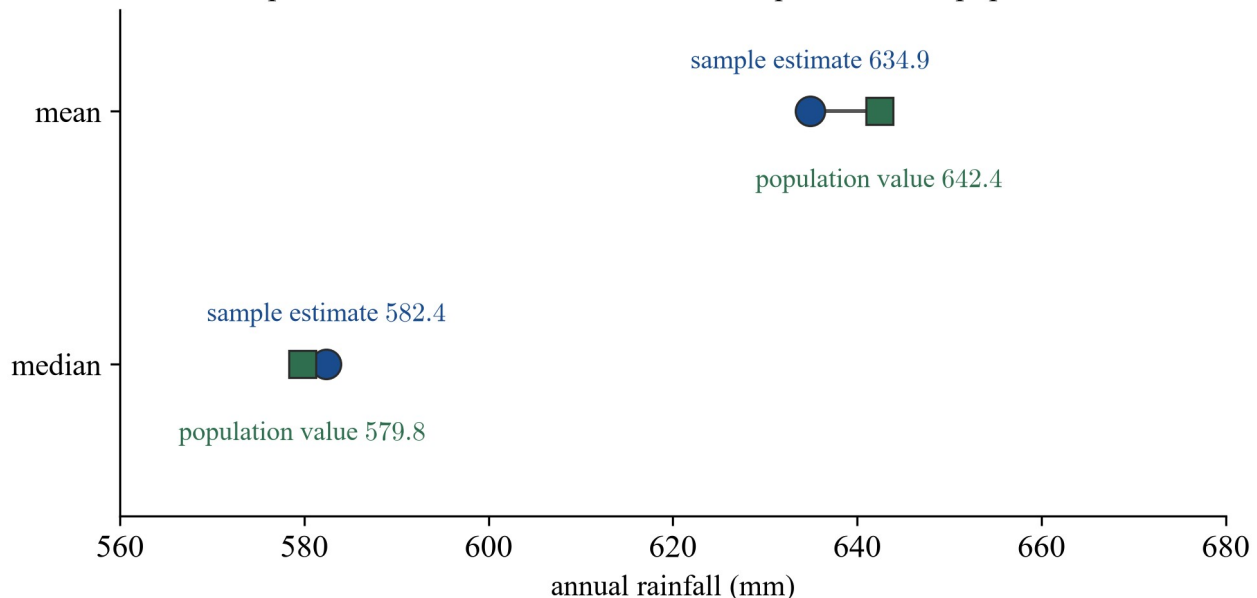
2025 rainfall for a simple random sample of 8 stations



Dot plot of 2025 rainfall for a random sample of 8 stations, with the sample mean (634.9 mm) and sample median (582.4 mm) marked.

Take a simple random sample of 8 of the 81 stations — Mallacoota 1230.8, Coldstream 653.6, Scoresby Research Institute 621.8, Ballarat Aerodrome 585.0, Essendon Airport 579.8, Moorabbin Airport 543.4, Melbourne Airport 459.0, Westmere 405.8 (values in mm). The sample mean is 634.9 mm and the sample median is 582.4 mm, so the estimates are: population mean $\approx$ 634.9 mm, population median $\approx$ 582.4 mm.

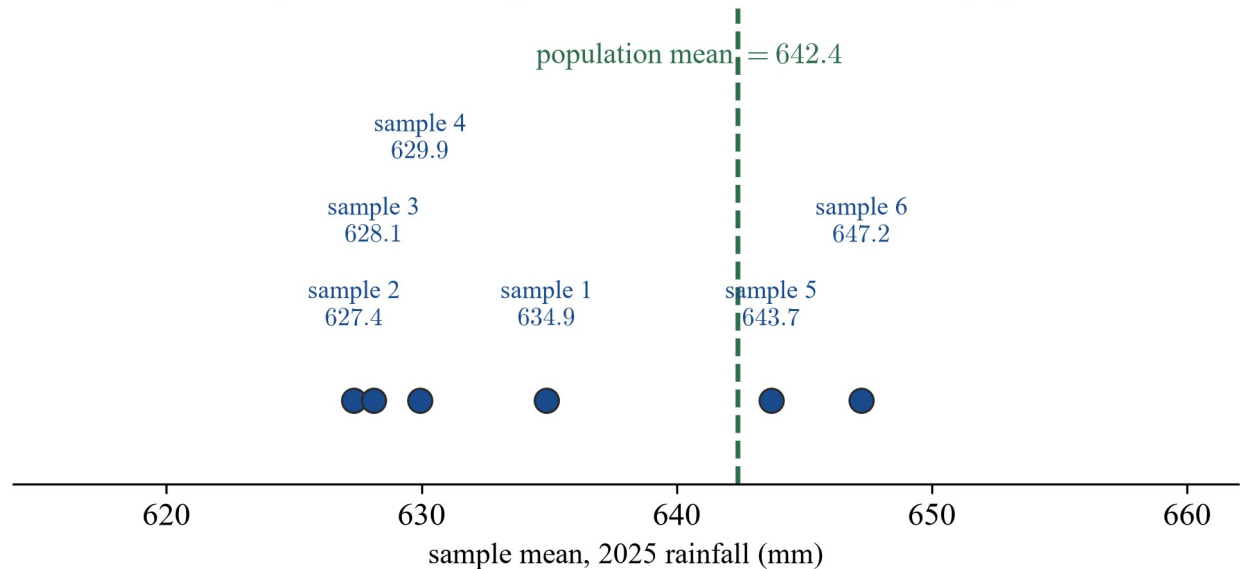
The sample estimates are close to — but not equal to — the population values



Dumbbell chart comparing the sample estimates (mean 634.9 mm, median 582.4 mm) with the population values (mean 642.4 mm, median 579.8 mm).

Against the true population values — mean 642.4 mm, median 579.8 mm — the sample mean is out by about 7.5 mm and the sample median by 2.6 mm. Both estimates are close; neither is exact. That is the whole story of estimation: a sample statistic stands in for a population value, and it carries an error you can only see when the population is known.

Six random samples of 8 stations gave six different estimates of the population mean



Number line showing six different sample means (627.4 to 647.2 mm) from six random samples of 8 stations, all near the population mean of 642.4 mm.

Different samples give different estimates. Six simple random samples of 8 stations give six sample means between 627.4 mm and 647.2 mm — a spread of about 20 mm around the population mean of 642.4 mm. None of the six is wrong; each is one fair draw. So a careful report says *about*, and a careful reader asks how large the sample was, because larger samples vary less (6D takes that further).

Mean or median: which story the data tells

When the data is roughly symmetric, the mean and median agree, and either summarises the centre. When a few values sit far from the rest, the mean is pulled towards them and the median almost does not move — the Level 9 rule from 6A. For the 81 station rainfalls the mean exceeds the median by 62.6 mm, entirely because of the wet tail. Quoting one or the other is a choice, and a report that quotes only the mean of a skewed quantity is promoting the bigger story.

Judging a claim

A report’s conclusion is only as strong as its method. Before believing *Australians earn...* or *shoppers back...*, ask the checklist: who was measured, by what method, with whom left out, and does the published statistic even match the published data? A claim with no method attached cannot be checked — and a sample that misses whole groups describes the groups it reached, not the population it names. That is AS 17’s skill: using how the data was obtained to support a conclusion, or to question one.

Worked examples

Example 1 — scaffolded

Read the ABS Average Weekly Earnings release like a statistician: how was the data obtained, what are the variables, and what is the published “average”? Figures as published for full-time adults, May 2026, seasonally adjusted.

Working	Comment
Report: ABS, <i>Average Weekly Earnings, Australia, May 2026</i> , released 13 August 2026	A real, named, dated report — the claim can be checked against it.
How obtained: a survey of employing businesses, run online, twice a year (May and November)	The ABS asks businesses, not workers, for total pay and employee counts in the reference week.
Sample: a probability sample of businesses from the	A sample, not a census of businesses — and every

Working	Comment
ABS business register; target response rate 90 %	business has a known chance of selection.
Scope: full-time adult employees only; defence, agriculture, the self-employed and unpaid directors excluded	Checklist step 5: the “average weekly earnings” figure never describes those groups.
Numerical variable: weekly earnings (in dollars); categorical variables: sex; state or territory	One numerical and (at least) one categorical variable, as the content description expects.
Published figure: \$ 2,160.00 for persons = total earnings $\div$ number of employees	The headline number is a mean, built off totals.
The release publishes no median	With skewed earnings data, the typical worker sits below the mean — the report’s silence on the median is itself a finding.
$\therefore$ The report estimates a population mean from a well-documented business survey; a reader quoting it must say “full-time adult employees”, not “Australians”, and must know that no median is on offer.	Method, scope and statistic, all named.

Example 2 — scaffolded

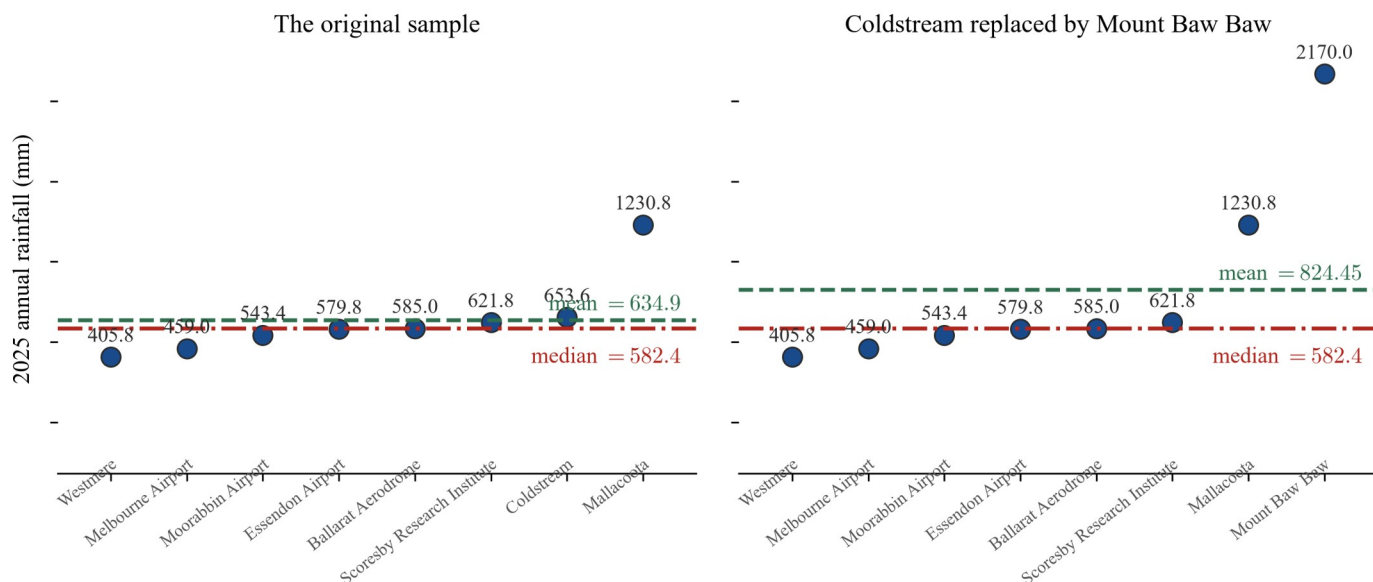
The BOM Victoria in 2025 table lists 81 stations with complete records. A simple random sample of 8 stations recorded the 2025 rainfalls shown in the figure (values in mm). Use the sample to estimate the population mean and median, then compare your estimates with the true population values: mean 642.4 mm, median 579.8 mm.

Working	Comment
Population: the 81 stations in the BOM table; the values wanted are the mean and median of their 2025 rainfalls	The sample will stand in for these 81 stations.
Sample values, in order: 405.8, 459.0, 543.4, 579.8, 585.0, 621.8, 653.6, 1230.	Ordering first makes the median easy.
Sum: $405.8 + 459.0 + 543.4 + 579.8 + 585.0 + 621.8 + 653.6 +$	One step per line.
Sample mean: $5079.2 \div 8 = 634.9$ mm	The estimate of the population mean.
Sample median: middle two values 579.8 and 585.0; $(579.8 + 585.0) \div 2 = 582.4$ mm	Eight values — even — so the median is the mean of the 4th and 5th.
Errors against the population: mean $642.4 - 634.9 = 7.5$ mm out; median $582.4 - 579.8 = 2.6$ mm out	Close, not equal — estimation always carries an error.
$\therefore$ From 8 of the 81 stations, the sample mean 634.9 mm estimates the population mean and the sample median 582.4 mm estimates the population median; both land within a few millimetres of the truth, and neither hits it.	A good estimate is not an exact one.

Example 3 — unscaffolded

Same sample, one change: suppose the draw had included the alpine station Mount Baw Baw (2170.0 mm) instead of Coldstream (653.6 mm). Calculate the mean and the median again, and decide which measure better represents a typical station. (The BOM’s own summary notes that the median is sometimes more representative than the mean of long-term average rain.)

New values, in order: 405.8, 459.0, 543.4, 579.8, 585.0, 621.8, 1230.8, 2170.0. New sum: $5079.2 - 653.6 + 2170.0 = 6595.6$. New mean: $6595.6 \div 8 = 824.45$ mm. New median: the middle two values are still 579.8 and 585.0, so the median is still 582.4 mm. One station changed and the mean jumped $824.45 - 634.9 = 189.55$ mm, while the median did not move at all. Seven of the eight stations sit below 824.45 mm, so the mean now sits well above most of the sample; the median still sits among the typical stations.



One value changed. The mean jumped from 634.9 to 824.45 mm; the median stayed at 582.4 mm.

Two dot plots of the same 8-station sample before and after replacing Coldstream with Mount Baw Baw (2170.0 mm): the mean moves from 634.9 to 824.45 mm while the median stays at 582.4 mm.

∴ For a skewed quantity like rainfall, the median is the more representative summary: the mean is pulled towards the wet tail, exactly as the BOM’s own note says.

Example 4 — unscaffolded

SCENARIO. A local news site reports: “Survey finds the average weekly earnings of Year-9 part-time workers is \$180.” The fine print says 25 Year-9 students with part-time jobs at one school recorded one week’s pay, that the total recorded was \$4,150, and that each student also named their job type (retail, hospitality, or other). Analyse the report.

Check the statistic against the data: $4,150 \div 25 = 166$, not 180 — the published “average” does not match the published total. Check the numerical variable and the categorical variable: weekly earnings (in dollars) is numerical; job type is categorical — the report at least lets both be checked. Check the method: the students were the ones with jobs, at one school, in one week. Students with no part-time job earn \$0 and were never asked, so “the average weekly earnings of Year-9 students” would have to count them; the report’s \$166 (or \$180) describes only the part-time workers of one school in one week. One week also matters: a student paid fortnightly may have recorded \$0 in the recorded week. ∴ The claim fails twice — the headline number contradicts the report’s own total, and a sample of part-time workers at one school cannot speak for Year-9 students generally. The method does not support the conclusion.

Practice questions

Scaffolded

Q1. The ABS *Average Weekly Earnings* release, May 2026. Answer from the report and its methodology page.

- Who did the ABS ask for the data — workers, or businesses?
- In which two months each year is the survey run?
- Name one group of workers the survey leaves out.
- Is the published \$2,160.00 a mean or a median? How can you tell from the methodology?

Q2. Name one numerical variable and one categorical variable in each report.

- The ABS *Average Weekly Earnings* release.

- b. The BOM *Victoria in 2025* station table. (Hint: a numerical variable's values are numbers you could average; a categorical variable's values are labels.)

Q3. SCENARIO. A school survey asks a random sample of 5 students how many minutes their trip to school takes: 12, 15, 20, 25, 33.

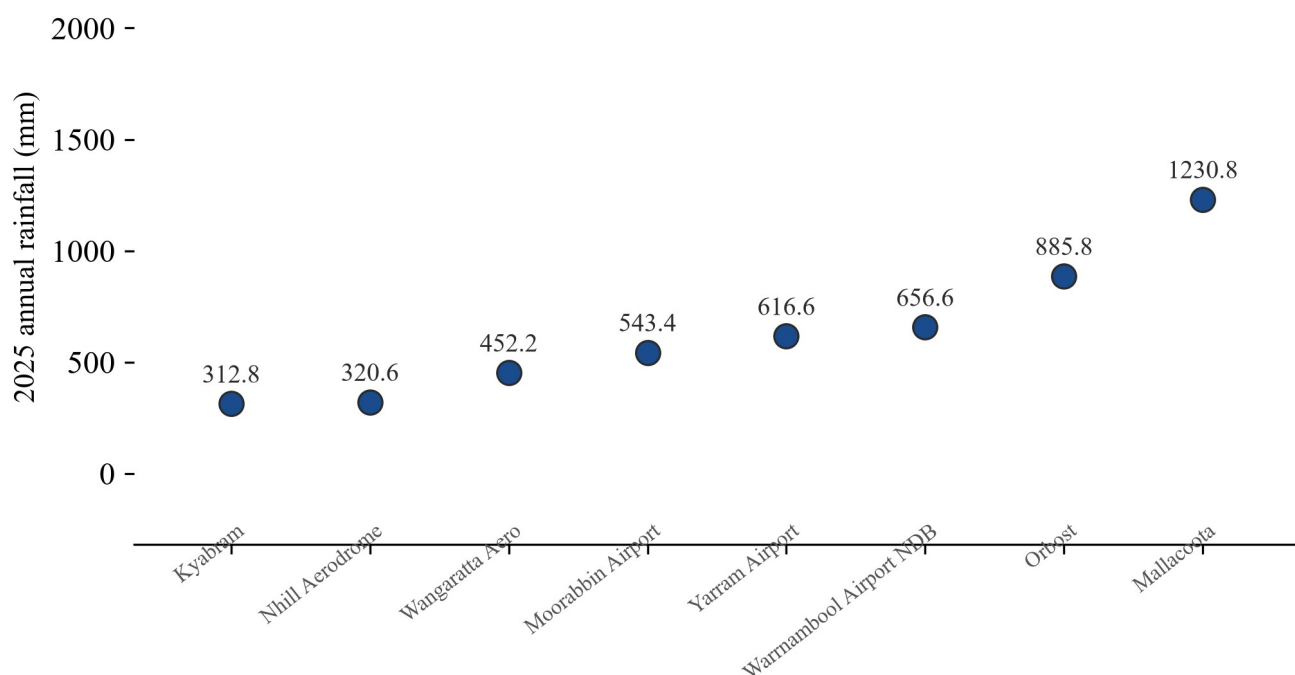
- Find the sum of the five times.
- Hence find the sample mean.
- The sample mean estimates which population value? Say the population in words.

Q4. SCENARIO. Eight students record their maths homework last week, in minutes: 14, 16, 16, 18, 21, 22, 25, 60. The values are already in order.

- Write down the two middle values.
- Hence find the sample median.
- The sample median estimates which population value?

Q5. The figure shows the 2025 rainfalls (mm) for a second simple random sample of 8 stations.

2025 rainfall for another simple random sample of 8 stations



Dot plot of 2025 rainfall for a second random sample of 8 stations: 312.8, 320.6, 452.2, 543.4, 616.6, 656.6, 885.8 and 1230.8 mm.

- List the eight values in order.
- Calculate the sample mean.
- Calculate the sample median.
- The population mean is 642.4 mm and the population median is 579.8 mm. Which of your two estimates landed closer?

Q6. Back to the Q4 data: 14, 16, 16, 18, 21, 22, 25, 60.

- Calculate the sample mean.
- Which single value is pulling the mean away from the other seven?
- A student writes “the typical student did 24 minutes of homework”. Using your mean from part a and your median from Q4, explain why the median is the fairer summary here.

Un scaffolded

Q7. ABS *Average Weekly Earnings*, May 2026, seasonally adjusted, full-time adults: males \$ 2,285.90, females \$ 1,964.50.

- Calculate the gap between the male and female means.
- Express the female mean as a percentage of the male mean, to one decimal place.
- Earnings data is skewed by a few very large salaries. Explain why a median weekly earnings figure would say more about the typical worker than the mean does.

Q8. ABS *Average Weekly Earnings*, May 2026.

- Name two groups the survey leaves out.
- A headline reads “Australians earn \$ 2,160 a week on average”. Give two reasons the headline overreaches what the release measured.
- The May 2026 ordinary-time mean for Victoria was \$ 2,041.10 (persons), against \$ 2,083.70 for Australia. Calculate the difference, and say which is larger.

Q9. SCENARIO. A school library takes a random sample of 10 students and records books borrowed last month: 0, 1, 1, 2, 2, 2, 3, 3, 4, 6. Calculate the sample mean and the sample median, and state what each estimates.

Q10. SCENARIO. The librarian, who borrowed 18 books, is added to the Q9 list of 10 students, making 11 values. Recalculate the mean and the median. Which estimate moved more, and what does that tell you about quoting the mean of a skewed set?

Q11. BOM *Victoria in 2025*: five stations’ 2025 rainfalls were Mildura Airport 159.7 mm, Bendigo Airport 396.2 mm, Melbourne Airport 459.0 mm, Cape Otway Lighthouse 823.8 mm and Falls Creek 1466.4 mm. Treat them as a sample of the 81-station population (mean 642.4 mm, median 579.8 mm). Calculate the sample mean and median, and comment on how each estimate compares with the population value.

Q12. The same five stations’ long-term average rainfalls are Mildura Airport 287.4 mm, Bendigo Airport 510.7 mm, Melbourne Airport 537.3 mm, Cape Otway Lighthouse 893.7 mm and Falls Creek 1393.3 mm. The population of 81 long-term averages has mean 703.5 mm and median 649.8 mm. Calculate the sample mean and median estimates and comment on which estimate landed closer this time. Does one sample always give a good estimate of every value?

Q13. SCENARIO. An online ad reads: “9 out of 10 people surveyed preferred the new recipe!” No method, no date, and nothing about the 10 beyond what the slogan says. List three checklist questions the ad leaves unanswered, and say why each one matters.

Q14. SCENARIO. A report on single-use plastics reads: “Survey: 62% of shoppers support a ban.” The method: 200 shoppers were approached outside one shopping centre on a Saturday morning; 124 said they support the ban.

- Verify the 62 %.
- Name two groups of shoppers the method misses.
- The report’s conclusion names shoppers in general. Is that conclusion reasonable from this method? Explain.

Q15. The BOM *Victoria in 2025* population of 81 station rainfalls has mean 642.4 mm and median 579.8 mm. A journalist must quote one number as “the typical annual rainfall at a Victorian station”. Which should be quoted, and why? Refer to the shape of the distribution.

Q16. Melbourne Airport’s mean annual rainfall over 1970–2026 is 536.2 mm, and the mean number of days with at least 1 mm of rain is 86.1 per year.

- Estimate the mean weekly rainfall by dividing the annual mean by 52, to one decimal place.
- A presenter says “Melbourne gets about 10 mm of rain in a typical week”. Use the rain-day figure to explain why that sentence misleads, even though the arithmetic in part a is sound.

Q17. BOM *Victoria in 2025*: of the 81 stations with complete records, 3 were ranked *very low*, 33 *low*, 38 *average* and 7 *high* for 2025 rainfall.

- a. How many stations ranked below average?
- b. Express this as a fraction of the 81 stations in simplest form, and as a percentage to one decimal place.
- c. The summary states that Victoria's area-averaged 2025 rainfall was 544.6 mm, 18 % below the 1961–1990 average. Do the station ranks support that statement? Explain.

Q18. SCENARIO. A news site reports: “Survey: the average Year-9 student spends 4.5 hours per week on homework.” The method: 50 Year-9 students at one Melbourne school kept a homework diary for one week; the diaries totalled 225 hours; each student also recorded whether they had a quiet place to work (yes or no).

- a. Describe how the data was obtained, and give one weakness of the method.
 - b. Name the numerical variable and the categorical variable.
 - c. Check the reported 4.5 against the published total.
 - d. Does the sample support the conclusion “the average Year-9 student spends 4.5 hours per week on homework”? Explain.
-

Answers

Q1 a Businesses (employing businesses), not workers; b May and November; c any one of: the defence forces, agriculture/forestry/fishing workers, the self-employed, unpaid directors; d a mean — it is total earnings divided by the number of employees, and no median is published. **Q2** a Numerical: weekly earnings (in dollars); categorical: sex (or state/territory). b Numerical: annual rainfall (mm); categorical: decile rank (very low/low/average/high) — or the station itself. **Q3** a 105 minutes; b 21 minutes; c the mean trip time of all students at the school. **Q4** a 18 and 21; b 19.5 minutes; c the median homework time of all students in the group the sample came from. **Q5** a 312.8, 320.6, 452.2, 543.4, 616.6, 656.6, 885.8, 1230.8; b 627.35 mm; c 580.0 mm; d the median (580.0 vs 579.8, out by 0.2 mm; the mean is out by 15.05 mm). **Q6** a 24 minutes; b the 60; c the median (19.5) sits among the seven ordinary values; the mean is dragged up by the one 60-minute student. **Q7** a \$321.40; b 85.9%; c a few large salaries pull the mean up, so the median better describes the worker in the middle. **Q8** a Any two excluded groups; b the figure covers full-time adult employees only (not all Australians), and it is a mean of a sample of businesses, not a census; c \$42.60, the Australian figure is larger. **Q9** Mean 2.4 books, estimating the mean books borrowed by all students; median 2 books, estimating the median. **Q10** Mean ≈ 3.8 books, median 2 books; the mean moved (from 2.4 to ≈ 3.8) and the median did not — the mean is not robust to one extreme value. **Q11** Mean 661.02 mm (out by 18.62 mm), median 459.0 mm (out by 120.8 mm); the mean landed closer this time. **Q12** Mean 724.48 mm (out by 20.98 mm), median 537.3 mm (out by 112.5 mm); the mean landed closer; no — one small sample can estimate one value well and another badly. **Q13** Any three of: who were the 10 (and who chose them)?; how were they asked (could saying “yes” be rewarded)?; when and where was it done?; is 10 enough to stand in for everyone?; who did not respond? Each matters because the method decides who the 9-out-of-10 describes. **Q14** a $\frac{124}{200} = 62\%$; b e.g. weekday shoppers, online shoppers, people who avoid that centre on Saturdays; c no — the sample describes Saturday shoppers at one centre, not shoppers in general. **Q15** The median, 579.8 mm — the distribution is skewed by the wet alpine tail, which pulls the mean above what most stations receive. **Q16** a $536.2 \div 52 \approx 10.3$ mm; b rain falls on only about 86 days a year — most weeks get 0 mm and wet weeks get far more than 10 mm, so no week is “typical”. **Q17** a 36; b $\frac{4}{9} \approx 44.4\%$; c yes — nearly half the stations ranked low or very low, consistent with a year 18% below average. **Q18** a Diaries kept by 50 Year-9 students at one school for one week; weakness: one school (or one week) is not all Year-9 students (or all weeks); b numerical: hours of homework; categorical: quiet place to work (yes/no); c $225 \div 50 = 4.5$ — the arithmetic checks out; d no — one school’s one week cannot speak for Year-9 students everywhere.

Fully-worked solutions

Q1. a The methodology page says the data is collected from employing businesses; workers are never asked. b The survey is conducted twice a year, in May and November. c For example, the self-employed (also acceptable: defence forces, agriculture/forestry/fishing workers, unpaid directors). d A mean: the ABS divides total weekly earnings by the number of employees — and the release publishes no median at all. $\therefore$ The release is a business survey publishing means, twice a year, over a defined scope.

Q2. a Weekly earnings in dollars is numerical (it makes sense to average it); sex — male/female/persons — and state/territory are categorical. b Annual rainfall in millimetres is numerical; the decile rank (very low, low, average, high) is categorical. $\therefore$ Both reports carry at least one variable of each kind, as the content description expects.

Q3. a $12 + 15 + 20 + 25 + 33 = 105$ minutes. b Mean $= 105 \div 5 = 21$ minutes. c The sample mean estimates the mean trip time of the population — all students at the school. $\therefore$ 21 minutes is an estimate of the population mean, not the population mean itself.

Q4. a With 8 ordered values the middle two are the 4th and 5th: 18 and 21. b Median $= (18 + 21) \div 2 = 19.5$ minutes. c The median homework time of the whole group the sample was drawn from. $\therefore$ Even counts take the mean of the two middle values.

Q5. a In order: 312.8, 320.6, 452.2, 543.4, 616.6, 656.6, 885.8, 1230.8 (mm). b Sum = 5018.8; mean $= 5018.8 \div 8 = 627.35$ mm. c Middle two: 543.4 and 616.6; median $= (543.4 + 616.6) \div 2 = 580.0$ mm. d Median

error: $580.0 - 579.8 = 0.2$ mm; mean error: $642.4 - 627.35 = 15.05$ mm — the median landed closer. $\therefore$ One random sample can nail one population value and miss the other; estimates are judged on average, not one draw.

Q6. a Sum = 192; mean = $192 \div 8 = 24$ minutes. b The single 60-minute value; the other seven sit between 14 and 25. c Seven of the eight students did 25 minutes or less, so the median 19.5 describes the typical student; the mean 24 is dragged up by the one 60, and “typical = 24” overstates six of the eight. $\therefore$ With a skewed set, the median is the fairer summary of “typical”.

Q7. a $2,285.90 - 1,964.50 = 321.40$: the gap is \$321.40 per week. b $\frac{1,964.50}{2,285.90} \times 100 \approx 85.9399\ldots \approx 85.9\%$. c A small number of very large salaries pulls the mean above what most workers receive; the median — the worker in the middle — is not pulled, so it describes the typical worker better. $\therefore$ The mean answers “what is the total divided up?”; the median would answer “what does the typical worker get?”.

Q8. a Any two of: members of the permanent defence forces; employees in agriculture, forestry and fishing; the self-employed; unpaid directors; employees of private households. b First, the scope is full-time adult employees — part-time workers, teenagers, the self-employed and people without jobs are not described by the figure. Second, the figure is a mean built from a sample of businesses, so it is an estimate with a method attached, not a fact about every Australian. c $2,083.70 - 2,041.10 = 42.60$: the Australian ordinary-time mean is \$42.60 larger than Victoria’s. $\therefore$ Scope, method and geography all limit what the headline may claim.

Q9. Sum = $0 + 1 + 1 + 2 + 2 + 2 + 3 + 3 + 4 + 6 = 24$; mean = $24 \div 10 = 2.4$ books. With 10 ordered values the middle two are the 5th and 6th, both 2; median = 2 books. $\therefore$ The sample mean 2.4 estimates the population mean books borrowed and the sample median 2 estimates the population median.

Q10. New sum = $24 + 18 = 42$ over 11 values; mean = $\frac{42}{11} \approx 3.818\ldots \approx 3.8$ books. With 11 ordered values the median is the 6th value: 2 books. The mean moved from 2.4 to ≈ 3.8 while the median did not move at all. $\therefore$ One extreme value pulls the mean and leaves the median — quoting the mean of a skewed set overstates the typical case.

Q11. Sum = $159.7 + 396.2 + 459.0 + 823.8 + 1466.4 = 3305.1$; sample mean = $3305.1 \div 5 = 661.02$ mm. Median of the five: the 3rd value, 459.0 mm. Against the population: the mean is out by $661.02 - 642.4 = 18.62$ mm; the median is out by $579.8 - 459.0 = 120.8$ mm. $\therefore$ This sample’s mean is a good estimate and its median a poor one — the sample happened to miss the middle of the distribution.

Q12. Sum = $287.4 + 510.7 + 537.3 + 893.7 + 1393.3 = 3622.4$; sample mean = $3622.4 \div 5 = 724.48$ mm. Median: the 3rd of the five ordered values, 537.3 mm. Against the population (mean 703.5, median 649.8): the mean is out by 20.98 mm, the median by 112.5 mm — again the mean lands closer. $\therefore$ The same five stations estimate the mean well and the median badly; one sample does not guarantee a good estimate of every population value.

Q13. Who were the 10 people, and who chose them? If the maker picked friends or fans, the sample leans before a question is asked. How were they asked? A tasting with a free sample rewards “yes”, and “preferred” may simply mean “chose the free one”. Is 10 people enough, and do they stand in for everyone the ad addresses? A sample of 10 varies so widely that 9-out-of-10 proves little. $\therefore$ A claim with no usable method cannot be checked, and an unchecked claim is advertising, not evidence.

Q14. a $\frac{124}{200} = \frac{62}{100} = 62\%$ — the arithmetic checks. b For example: shoppers who shop on weekdays or evenings, shoppers who buy online, and people who never visit that centre — none of them entered the sample. c No. The method describes people outside one centre on a Saturday morning; “shoppers” in general includes everyone the method missed, and Saturday shoppers may differ (more leisure shoppers, fewer time-poor workers). $\therefore$ The percentage is honest; the conclusion names a population the method never measured.

Q15. The histogram is skewed: most stations sit between 300 and 900 mm while a few alpine stations exceed 1,400 mm and one reaches 2170.0 mm. Those wet stations pull the mean (642.4 mm) above the median (579.8 mm); 642.4 mm is more rain than 50 of the 81 stations actually received. The median sits at the middle station and is untouched by the tail — exactly the point the BOM makes in its own note. $\therefore$ Quote the median, 579.8 mm, as the typical station’s year; the mean overstates what most stations receive.

Q16. a $536.2 \div 52 = 10.311... \approx 10.3$ mm per week on average. b Rain fell on a mean of only 86.1 days per year — roughly one day in four — so in a typical week it either does not rain at all or it rains well over 10 mm. The 10.3 mm is a mean spread over weeks that never look like it. $\therefore$ A mean per week is a true average but a false picture of a typical week; the presenter's sentence misleads even though the arithmetic is right.

Q17. a $3 + 33 = 36$ stations ranked very low or low. b $\frac{36}{81} = \frac{4}{9} \approx 0.4444... \approx 44.4\%$. c Yes: 44.4 % of stations ranked below average and only 7 ranked high, matching a state summary of 544.6 mm against a much higher 1961–1990 average — two views of the same dry year. $\therefore$ The categorical ranks and the area-averaged figure agree, which is what a coherent report should do.

Q18. a The data was obtained by diary: 50 Year-9 students at one school recorded one week of homework. Weakness: one school is not all schools (and one week is not all weeks — exam weeks and holiday weeks differ). b Numerical: hours of homework per week. Categorical: quiet place to work (yes/no). c $225 \div 50 = 4.5$ hours — the reported average matches the published total, so the statistic is at least honest. d No. The sample covers one school in one week; “the average Year-9 student” names a population the diary never measured, so the conclusion overreaches even though the arithmetic is right. $\therefore$ Honest arithmetic inside a weak method still produces an unsupported conclusion — method first, numbers second.

Sampling methods, sample-to-sample variation, and representation that promotes a point of view

6D · Chapter 6 — Statistics · Level 9

Content description VC2M9ST02 (word for word): *analyse how different sampling methods, and different samples using the same method, can affect the results of surveys and how choice of representation can be used to support a particular point of view*

Achievement standard anchor (AS 17): *They explain how sampling techniques and representation can be used to support or question conclusions or to promote a point of view.*

Elaboration ids for linkage: VC2M9ST02.E01, VC2M9ST02.E02, VC2M9ST02.E03

Prerequisite pointer: Level 8 Statistics is assumed — population vs sample and collection techniques (VC2M8ST01), distributions from primary and secondary sources with random and non-random sampling (VC2M8ST02), variation between random samples of the same size and the effect of sample size (VC2M8ST03), and investigations with samples and fair inference (VC2M8ST04). The method names and the repeated-sampling experiment are rehearsed below as tools, not re-taught. Level 8 ran the experiment; 6D does the consequences — which sampling method, and whose point of view.

Note on VC2M9ST02.E03 (exclusion recorded): E03's material on cultural bias relating to Aboriginal and Torres Strait Islander Peoples is excluded from this book under CONTENT_POLICY Rule 1 (D10; recorded as ACCEPTED here, since TOC.md must not be edited from inside a section). The content description is delivered in full via VC2M9ST02.E01 (critique of media visualisations) and VC2M9ST02.E02 (the landline / unknown-number survey mechanism) plus original scenarios, per D5 — elaborations are illustrations of a CD, not the CD itself.

Note on sources: network sources were unreachable at authoring, so every quantitative context in this section is a clearly-labelled SCENARIO with invented parameters (R2). No real data ships and therefore no source line is owed (R1).

Theory

Where Level 8 leaves off

At Level 8 you met the basic objects of statistics. The **population** is the whole group you want to say something about; a **sample** is the part of the population you actually collect data from; a **census** asks every member of the population. You also met the sampling methods by name, and you ran the repeated-sampling experiment: draw many samples of the same size from the same population, and the results come out different from sample to sample.

Level 9 builds the consequences on top of that experiment. The content description's verb is *analyse*: take a survey or a display apart, see what each choice does to the result, and judge the effect. This subtopic analyses three things:

1. how the **method** used to choose a sample leans the result;
2. how **sample-to-sample variation** limits what one sample can show;
3. how a choice of **representation** — a graph, a number, a headline — can promote a point of view.

The sampling methods, in short

These are the Level 8 definitions, rehearsed as tools.

Random methods — every member of the population has a known chance of being chosen:

- **Simple random sampling** — every sample of the chosen size is equally likely; names are drawn from a hat or chosen by random numbers.
- **Systematic sampling** — pick a random starting point, then take every k th name on the list.
- **Stratified sampling** — split the population into groups, and take a random sample from each group in proportion to the group's size.

- **Cluster sampling** — choose whole groups at random, then include everyone in the chosen groups.

Non-random methods:

- **Quota sampling** — fill a set number of places from each group, but let the person collecting the data choose who fills them.
- **Convenience sampling** — ask whoever is easiest to reach.
- **Judgement sampling** — the person collecting the data picks who they think is typical.

Only random methods give every member of the population a known chance of being chosen, so only random methods support a claim about the whole population. The non-random methods are still used in real life because they are cheap and quick — but the lean they bring must be named, and that is the work of this subtopic.

Bias and representativeness

In this subtopic, **bias** means a lean that is built into the way data is collected or displayed, so that results are pulled in one direction every time the method is run. Bias is not the same as a mistake: a survey can be carried out carefully and still be biased, because of who the method could reach. Bias is also not the same as sample-to-sample variation, which pulls results in *different* directions on different runs.

A sample is **representative** of the population when its make-up matches the population's make-up closely enough for the question being asked. Representative is the goal; bias is what moves a sample away from it.

Three common ways a collection method builds in a lean:

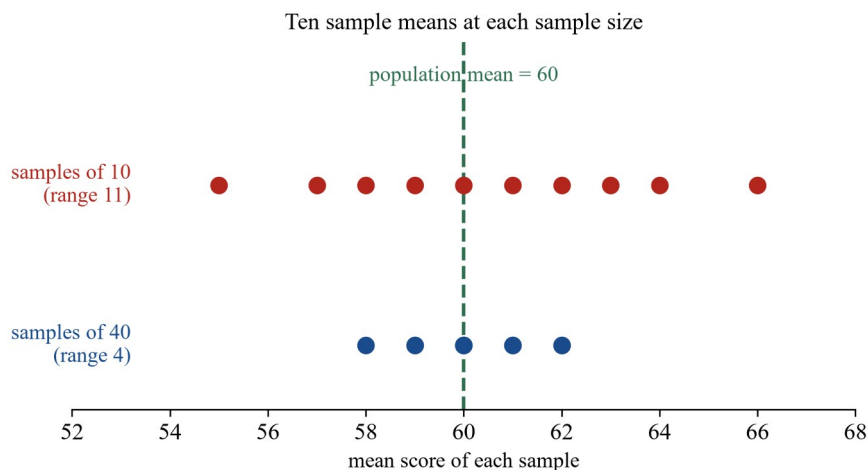
- **Convenience samples over-represent people who are easy to reach.** Asking shoppers at one store at 10 am on a Tuesday misses everyone who works there at that time.
- **Voluntary responses over-represent people with strong feelings.** People who care about a question — or are angry about it — are far more likely to reply than people who do not.
- **Contact methods reach only the people who have, and answer, that contact.** A phone poll only reaches people whose number it can call, and who answer it.

Sampling variation: different samples, different results

Even a perfectly fair method gives different results on different runs — that is the Level 8 experiment. Three consequences matter for the rest of this subtopic.

1. A result from one sample is an **estimate** of the population value. It sits near the truth, but it is not the truth.
2. Smaller samples vary more. The range of possible results is wider for a small sample than for a large one, so a large sample gives a tighter estimate.
3. When two fair surveys give different results, the difference is not automatically a mistake or a discovery. First check whether a difference of that size could come from sampling variation alone.

Differences between percentages are measured in **percentage points** (pp): the difference between 52 % and 47 % is 5 percentage points.



The figure shows ten sample means at each of two sample sizes, all drawn from the same population with mean 60. Both sets sit around 60 — a fair method does not lean — but the samples of 10 spread over a range of 11 points while the samples of 40 spread over a range of only 4. Bigger samples vary less.

Coverage and response: who the method actually reaches

A method can be random in principle and still miss people in practice. Two Level 8 ideas turn up constantly in real surveys.

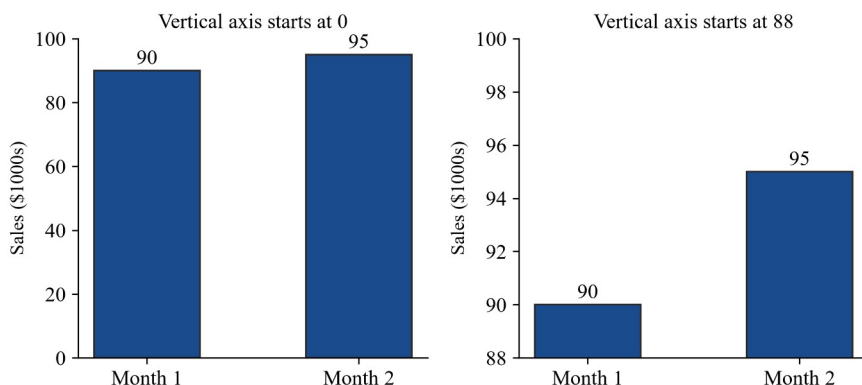
- **Coverage** — who the method can reach at all. SCENARIO: a town survey phones landline numbers in working hours. Households with no landline never enter the sample, and households where everyone works daytime rarely answer. The sample that answers is older and more home-based than the town.
- **Response** — who, among the people reached, actually takes part. People who do not take part produce **non-response**, and people who rush to take part in an open call produce **self-selection**. If the people who answer differ from the people who do not, the result leans — even when the contact list was fair.

Fewer people answer calls from numbers they do not know than used to, and fewer households keep a landline than used to. Both changes make phone surveys harder to trust than they were, unless the survey also chases up the people who did not answer.

How a display can promote a point of view

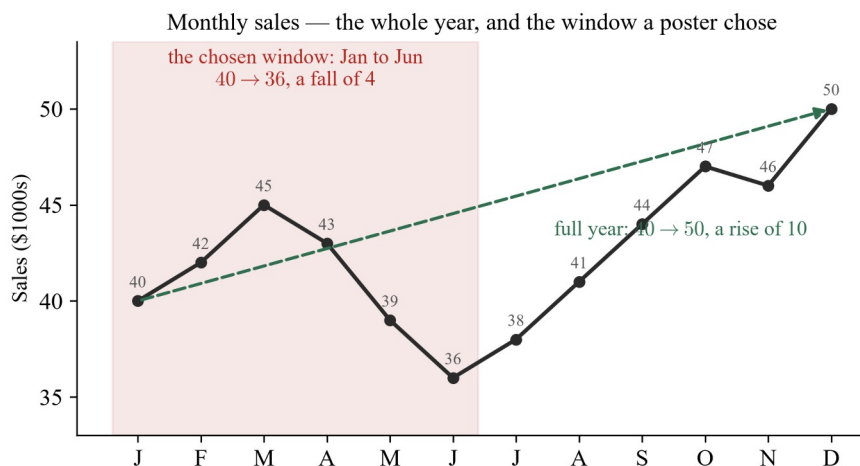
The same numbers can be drawn in ways that lead a reader to different conclusions. Five techniques do most of the work.

1. **A cut axis.** The vertical axis of a bar graph starts above zero, so small differences look large. A bar graph should start at zero unless the break is clearly shown and explained.



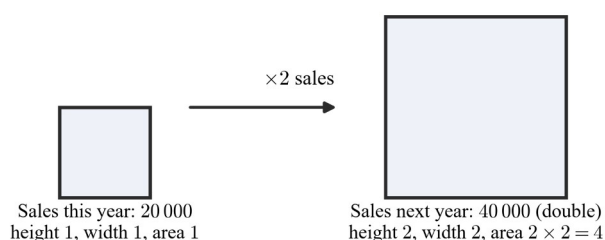
The numbers are almost equal (95 vs 90), but the cut axis makes Month 2 look about $3\frac{1}{2}$ times Month 1.

2. **A chosen window.** Only part of a time series is shown — the part that tells the story the author wants. The numbers outside the window are not wrong; they are simply hidden.



3. **A stretched vertical scale.** The same rise is drawn using fewer centimetres per unit, so the line looks steeper. Nothing has been changed except the drawing.
4. **A picture that grows too fast.** In a picture graph, a value that doubles is sometimes drawn as a picture twice as tall *and* twice as wide — and the picture's area then grows by $2 \times 2 = 4$ times. The eye reads the whole picture, so the growth looks four times as big as it is.

The value doubled, but the picture's area became 4 times as big



5. **A chosen statistic.** When one value in a data set sits far from the rest, the mean is pulled towards it and the median is not. Quoting the mean and quoting the median can tell two different stories about the same data.

One more technique needs no drawing at all: leaving information out. A claim with no sample size, no method, and no response rate cannot be checked — and a claim that cannot be checked is doing the reader's thinking for them.

None of this is about one side of any argument. Every technique above can promote any point of view, and the same critique applies whichever side uses it. In this subtopic, marks are given for the strength of the analysis, never for whether the reader agrees with the claim being analysed.

The analysis checklist

This checklist is the working tool of the subtopic. Every worked example and every question below runs some part of it.

1. Who is the population, and who was actually reached?
2. What method chose the sample — and who does that method over-represent or miss?
3. How large was the sample, and how much could a fair sample of that size vary?
4. What has the display changed — axis, window, scale, picture size, or choice of statistic?
5. What would a fair collection, or an honest display, show instead?
6. Whose point of view does the representation support?

Worked examples

Example 1 — scaffolded

SCENARIO. A school has 900 students. A student journalist wants to know whether the school day should start later. She stands outside the library at lunchtime and asks 30 Year-12 students; 21 of them say yes. She writes: “Students want a later start.” Analyse her survey.

Working	Comment
Population: all 900 students	Step 1 of the checklist: the group she wants to speak about.
Sample: 30 Year-12 students outside the library	Step 2: the group she actually reached.
$\frac{30}{900} = \frac{1}{30} \approx 3.3\%$ of the population	A small sample is not automatically bad — but its method must be fair, and this one is not.
Method: convenience sampling	She asked whoever was easiest to reach. There was no random step.
Over-represented: Year-12 students who use the library at lunch	Two groups at once: the oldest year level, and the students who spend lunch in the library.
Missed: every other year level, and students at sport, at work, or off site	They had no chance of being asked.
$\frac{21}{30} = 70\%$ of her sample said yes	The arithmetic is correct.
Conclusion: “students want a later start” is not supported	Step 6: the 70% describes Year-12 library users, not the school. The method chose a group with its own habits and timetable.

Example 2 — scaffolded

SCENARIO. A class of 20 students votes on an excursion; 12 of the 20 are in favour. Two groups each take a simple random sample of 5 students. Group A’s sample has 4 students in favour; group B’s sample has 2. Explain how both samples can be fair and both can miss the class value.

Working	Comment
Population value: $\frac{12}{20} = 60\%$ in favour	Step 3: the value the samples are estimating.
Sample A: $\frac{4}{5} = 80\%$	20 percentage points above the class value.
Sample B: $\frac{2}{5} = 40\%$	20 percentage points below the class value.
Neither group made a mistake	Both samples were chosen in a fair way. Chance gave group A a high sample and group B a low one.
A sample of 5 from 20 can sit 20 percentage points either side of the truth	Step 3 of the checklist: small samples vary a lot.
$\therefore$ One small sample alone cannot show whether the class is “for” or “against” — both results are inside the range that fair sampling produces	A result has to be read together with the variation a sample of that size can produce.

Example 3 — scaffolded

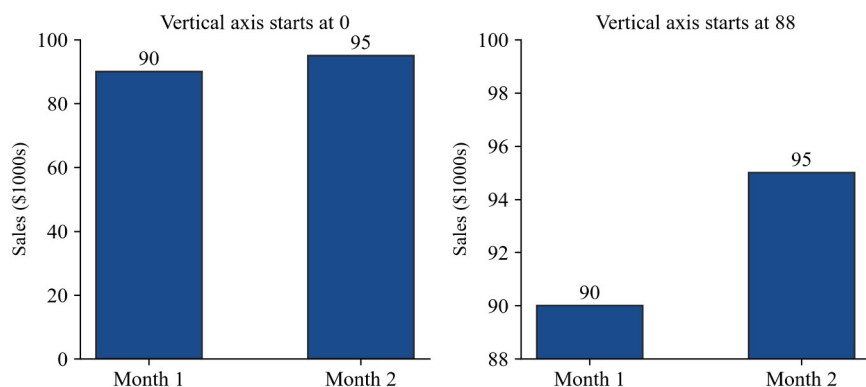
SCENARIO. A town has 200 households: 80 younger households and 120 older households. A survey company phones landline numbers during working hours. In this town, 16 of the 80 younger households have a landline and 84 of the 120 older households have one. What does a landline survey actually reach?

Working	Comment
Younger households with a landline: $\frac{16}{80} = 20\%$	The group the poll mostly misses.
Older households with a landline: $\frac{84}{120} = 70\%$	The group the poll mostly reaches.
All households with a landline: $16 + 84 = 100$, so $\frac{100}{200} = 50\%$	Half the town is reachable at all.
Younger households as a share of the town: $\frac{80}{200} = 40\%$	Step 1: who the town actually is.
Younger households as a share of landline households: $\frac{16}{100} = 16\%$	Step 2: who the poll can reach.
Conclusion: the poll's "town" is much older than the real town	Younger households are 40% of the town but only 16% of the reachable sample. Any result the poll gives is a result about landline households, not about the town.

The same mechanism has a second stage: even households that *have* a landline often do not answer calls from numbers they do not know. Each unanswered call is non-response, and it leans the sample a second time, towards people who answer unknown numbers.

Example 4 — scaffolded

SCENARIO. A small company's sales were \$90 000 in month 1 and \$95 000 in month 2. Two bar graphs are drawn from exactly these numbers: one with the vertical axis starting at 0, and one with it starting at 88. The figure shows both. Analyse the second graph.



The numbers are almost equal (95 vs 90), but the cut axis makes Month 2 look about $3\frac{1}{2}$ times Month 1.

Working	Comment
True comparison: $\frac{95}{90} = \frac{19}{18} \approx 1.06$	Step 5: what the numbers actually say — month 2 is about 1.06 times month 1.
Drawn bar lengths on the cut axis: $90 - 88 = 2$ and $95 - 88 = 7$	The eye reads the length of the bar above the axis start.
Comparison the picture shows: $\frac{7}{2} = 3.5$	Step 4: the cut axis makes month 2 look 3.5 times month 1.
3.5 against 1.06 — the display exaggerates the growth by more than three times	A difference of about 6% is drawn as if it were hundreds of percent.
$\therefore$ The cut axis promotes the point of view "sales are booming" without changing a single number; an honest version starts the axis at 0	Steps 4 and 6 of the checklist: name the change, then name whose story it tells.

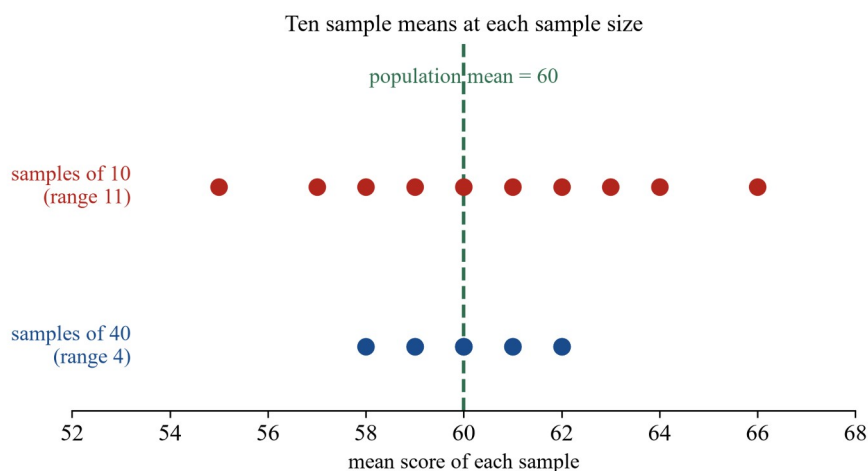
Example 5 — unscaffolded

SCENARIO. A streaming company wants to know whether viewers want a second series of one of its shows. It posts the question on the show’s own fan page, and the fan group’s page manager encourages members to vote. Analyse the poll.

The people who see the poll are already fans — the page belongs to the show. The people who vote choose themselves; viewers who gave up on the show, or never heard of it, never enter the data at all. This is self-selection stacked on top of a contact list that already leans towards fans. A vote that comes back overwhelmingly “yes, renew it” is exactly the result the method is built to produce. $\therefore$ The poll can describe the fans; it cannot say anything about all viewers, and a headline built on it would promote the show’s point of view, not measure the audience.

Example 6 — unscaffolded

SCENARIO. The mean score of a whole year level on a test is 60. Ten samples of 10 students give mean scores 61, 58, 64, 55, 66, 60, 57, 63, 59, 62. Ten samples of 40 students give mean scores 60, 59, 62, 58, 61, 60, 59, 61, 60, 62. Compare the two sets.



For the samples of 10, the range of the means is $66 - 55 = 11$, and the mean of the ten sample means is 60.5. For the samples of 40, the range of the means is $62 - 58 = 4$, and the mean of the ten sample means is 60.2. Both sets sit around the population mean of 60 — a fair method does not lean. But the samples of 10 spread over 11 points, while the samples of 40 stay within 4 points. $\therefore$ Larger samples vary less: an estimate from a sample of 40 is tighter than one from a sample of 10, and a result from a small sample should be read as one of a wide range of possible results, not as a number to rely on.

Example 7 — unscaffolded

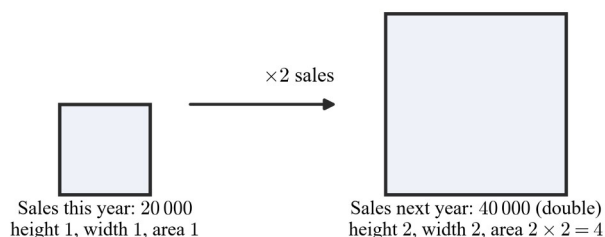
SCENARIO. Nine students record their weekly pocket money in dollars: 5, 5, 6, 6, 7, 8, 8, 9, 46. A student newspaper writes: “Students receive \$11 per week on average.” Analyse the choice of number.

The data is already ordered; the sum is $5 + 5 + 6 + 6 + 7 + 8 + 8 + 9 + 46 = 100$. The mean is $\frac{100}{9} \approx 11.11$ — the newspaper’s “\$11” is the mean, rounded down. The median is the 5th of the 9 ordered values: 7. The range is $46 - 5 = 41$. One value, 46, pulls the mean far above what 8 of the 9 students actually receive. If that one value grew to 96, the median would still be 7, while the mean would rise again. $\therefore$ Quoting the mean promotes the story “students receive a lot of pocket money”; the median describes what a typical student actually receives. Choosing which measure to quote is a choice of representation, and the reader deserves to see both.

Example 8 — unscaffolded

SCENARIO. A poster shows this year’s drink sales — \$20 000 — as a small carton, and next year’s forecast — \$40 000 — as a carton twice as tall and twice as wide. Analyse the picture.

The value doubled, but the picture's area became 4 times as big



The value doubles: $40 \div 20 = 2$. The picture's height doubles and its width doubles. Area grows by height $\times$ width: $2 \times 2 = 4$ times the area. The eye reads the whole carton, not just its height, so the picture says "four times" while the number says "twice". $\therefore$ The picture exaggerates the growth because it scales both dimensions; an honest version doubles one dimension only, or scales the area in proportion to the value.

Practice questions

Scaffolded

Q1. Name the sampling method used in each scenario, and the feature that tells you.

- Every student's name goes into a hat, and 50 names are drawn out.
- A school starts at a randomly chosen name on the roll, then selects every 15th student after that.
- A reporter asks the first 20 people who walk past the market.

Q2. SCENARIO. A town wants to know whether residents support a new walking track. Helpers stand outside the sports ground on Saturday morning and ask people as they arrive.

- Name the method.
- Which residents are over-represented, and which are missed?
- Predict the direction of the lean, and give your reason.

Q3. A Year-9 student says: "I asked 40 Year-9 students and 65% of them walk to school, so most students in the whole school walk to school."

- Why is the sample not representative of the whole school?
- Name a method that would be fairer, and say what it would fix.
- Is the mistake in the arithmetic or in the claim? Explain.

Q4. SCENARIO. A club has 10 members, and 6 of them support a proposal. Two simple random samples of 4 members are taken; sample 1 contains 3 supporters and sample 2 contains 2.

- What percentage of the whole club supports the proposal?
- Calculate the percentage of supporters in each sample.
- Both samples were chosen in a fair way. Explain how they can both miss the club value.

Q5. SCENARIO. A suburb has 300 households: 120 families with children and 180 other households. A landline list reaches 60 of the family households and 45 of the other households.

- How many households have a landline in all, and what percentage of all households is this?
- What percentage of all households are family households?
- What percentage of landline households are family households? What would a landline survey over-represent in this suburb?

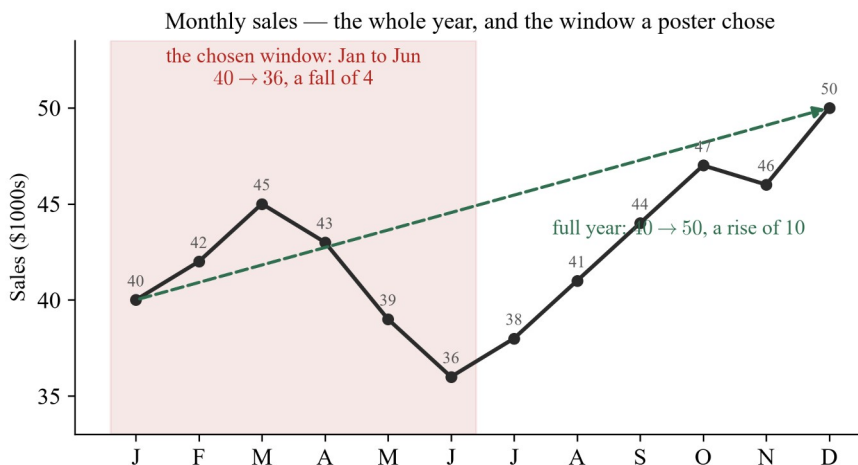
Q6. SCENARIO. A company phones 500 households about a new service; 120 households answer and give their opinion.

- Calculate the response rate as a percentage.
- How many households did not respond, and what percentage is this?
- Give one reason why the views of the 120 may not represent the views of all 500.

Q7. SCENARIO. A bar graph shows two values, 42 and 46, with the vertical axis starting at 40.

- Calculate the true ratio of the two values.
- Calculate the drawn length of each bar above the axis start.
- Calculate the ratio the picture shows. What does the graph make the reader believe?

Q8. SCENARIO. The figure shows a shop's monthly sales for a whole year. A poster about the shop displays January to June only.



- Read off the change in sales from January to December.
- Read off the change in sales from January to June.
- Explain why the poster's window tells one story and the full year tells the other.

Q9. SCENARIO. Eight students record their hours of sport per week: 2, 3, 3, 4, 4, 5, 6, 14.

- Calculate the mean.
- Calculate the median and the range.
- Which measure could be quoted to support "our students are very sporty"? Which describes the typical student?

Q10. SCENARIO. A picture graph shows this year's sales with a picture three times as tall and three times as wide as last year's, while sales only tripled.

- By what factor did the sales grow?
- By what factor did the picture's area grow?
- Explain why the picture overstates the growth.

Q11. SCENARIO. Five samples of size 5 from a population give mean scores 55, 63, 48, 67, 58. Five samples of size 30 from the same population give mean scores 59, 61, 60, 62, 58.

- Calculate the range of each set of sample means.
- Which sample size gives the more consistent results?
- Both sets average near 60. What does that say about the method, and what does the range say about trusting one small sample?

Q12. SCENARIO. A poster in the canteen reads: "Students love the new menu — 80% said yes!" In small print: 20 students were asked as they left the canteen on Monday.

- Checklist steps 1–2: name the population and the method.
- Checklist step 3: how large is the sample, and what can a sample this size not show?
- What would the poster need to add before “students” could be trusted?

Unscaffolded

Q13. For each situation, choose a sampling method (simple random, systematic, stratified, cluster, convenience, or judgement) and give your reason.

- The opinion of 300 Year-9 students spread across ten schools on a homework policy.
- Disease in 60 tomato plants growing in a single greenhouse.
- Support for a new policy across a whole city.

Q14. A school has 800 students, and a systematic sample of 40 students is needed.

- What interval should be used?
- How should the starting point be made random?
- Why is this method random but not simple random?

Q15. A town has a large older group and a smaller younger group, and a simple random sample might, by chance, contain few younger people.

- Explain how stratified sampling fixes this.
- Quota sampling also fills places from each group. What step does quota sampling leave out, and what can go wrong because of that?

Q16. SCENARIO. A coach wants to know whether students support buying new team uniforms. He asks his own team only; almost all of them say yes, and he reports that “students want new uniforms”.

- Name the method and the direction of the lean.
- Give a method that would support a claim about all students.
- What can the coach’s result honestly claim?

Q17. SCENARIO. Two groups each take a simple random sample of 25 students from the same school and ask the same question. Group A finds 48% in favour; group B finds 56%.

- What is the difference between the two results, in percentage points?
- Neither group made a mistake. Explain how two fair samples can differ.
- Would you expect a difference of this size if both samples had been 250 students? Give your reason.

Q18. SCENARIO. In a large population, exactly 55% support a proposal. Three simple random samples of size 10 give 40%, 60% and 50%.

- Calculate the range of the three sample results.
- None of the three results equals 55%. Explain why none of them is “wrong”.
- Predict what would change if the samples had been of size 100 instead, and give your reason.

Q19. SCENARIO. Five samples of size 5 give these percentages in favour: 40%, 60%, 40%, 80%, 60%. Five samples of size 25 give: 48%, 52%, 56%, 60%, 44%.

- Calculate the range of each set of results.
- Which sample size would you trust more for estimating the population value? Give your reason.
- The means of the two sets are 56% and 52%. What do these averages suggest about the population value?

Q20. SCENARIO. A survey in one region finds 52% support for an idea; a survey in another region, using the same method and the same sample size, finds 47%.

- What is the difference between the two results, in percentage points?
- Someone concludes that “the two regions really do think differently”. Explain why this conclusion is not supported yet.

- c. What extra evidence would support a real difference between the regions?

Q21. Design a fair survey to answer one question of your choice about your school (for example: start times, canteen menu, or sport options). Name the population, the method, the sample size, and what you would do about the people who do not respond. Give a one-line reason for each choice.

Q22. SCENARIO. A bar graph shows two values, 78 and 82, with the vertical axis starting at 76.

- Calculate the ratio the picture shows, using the drawn lengths only.
- Calculate the true ratio of the two values.
- Redraw the graph honestly without changing either number.

Q23. SCENARIO. A line graph shows a rise of 6 units over a horizontal run of 10 cm.

- At 1 unit per centimetre of height, how tall is the drawn rise?
- At 3 units per centimetre of height, how tall is the drawn rise?
- Calculate the drawn steepness in each case, and explain whose point of view the steeper version serves.

Q24. SCENARIO. A weather group records an average temperature in °C every five years: 12.8 (1990), 13.5 (1995), 13.1 (2000), 13.9 (2005), 13.4 (2010), 14.2 (2015), 14.0 (2020), 14.6 (2025).

- Calculate the change over the full record, from 1990 to 2025.
- Calculate the change over the window 1995 to 2010.
- A poster shows only the window 1995 to 2010 and says “temperatures stopped rising”. Evaluate this claim.
- Find a window in this data set that tells the opposite story, and calculate its change.

Q25. SCENARIO. Ten house prices on one street, in thousands of dollars, are: 580, 590, 600, 610, 615, 620, 625, 640, 650, 3000.

- Calculate the mean price in dollars.
- Calculate the median price in dollars.
- A real-estate flyer says “the average house on this street is worth \$853 000”. Whose point of view does the mean serve here, and which measure describes the typical house?

Q26. SCENARIO. Nine students record their hours of homework per week: 1, 2, 2, 3, 3, 3, 4, 4, 12.

- Calculate the mean, correct to two decimal places.
- Calculate the median and the range.
- Which measure would a student who wants to argue “we drown in homework” quote, and why? Which measure describes the typical student?

Q27. SCENARIO. A picture graph uses circles to show the sizes of two parks. The larger park has twice the area of the smaller one, so its circle is drawn with twice the radius.

- Show that doubling the radius multiplies the area of a circle by 4.
- Explain what the reader sees, and what is actually true.
- Describe an honest way to draw the two circles.

Q28. SCENARIO. An infographic about a proposed new road says “Survey: 68% of locals say the road is needed”. The survey asked 30 drivers at one intersection during the morning peak, and its bar graph starts the vertical axis at 60.

- Name three techniques from this subtopic that the infographic uses.
- Explain the effect of each one.
- Redesign the claim honestly: what should the infographic say and show instead?

Q29. SCENARIO. A newspaper emails a survey to 2000 of its readers; 60 reply, and most of the replies are against the proposal.

- Calculate the response rate as a percentage.

- b. Give two reasons why the 60 replies may lean in one direction.
- c. The newspaper writes “readers say no”. Evaluate the headline.

Q30. SCENARIO. An advertisement says “9 out of 10 customers recommend us”.

- a. What percentage is this?
- b. The advertisement does not say how the ten customers were chosen. Give two methods that would make the claim weak, and the direction of the lean each would create.
- c. What would the advertisement need to tell you before the claim could be trusted?

Q31. A radio show asks its listeners to phone in about a new parking rule; most of the calls are against the rule.

- a. Name the two leans that are stacked in this method.
- b. Why should the show not report the calls as a fair read of its listeners’ views?
- c. Describe a method that would give a fairer read.

Q32. You need a systematic sample of 40 students from a school of 800.

- a. Calculate the sampling interval.
- b. Explain how to choose the starting student in a fair way.
- c. A friend suggests “just take every 20th student who walks through the front gate”. What is wrong with that suggestion?

Q33. SCENARIO. The student council runs an online poll about the theme of the school dance; 25 of the school’s 400 students respond.

- a. Calculate the response rate as a percentage.
- b. Who is more likely to respond, and in which direction does that lean the result?
- c. The council writes “students have chosen”. Rewrite the claim honestly.

Q34. A student says: “I asked 5 friends and 80% of them support the proposal, so about 80% of the school must support it.” In Question 19, five fair samples of size 5 gave 40%, 60%, 40%, 80% and 60%.

- a. How wide is the spread of results for samples of size 5 in Question 19?
- b. Explain why 80% is a believable result from a fair sample of 5, even when far fewer than 80% of the population support the proposal.
- c. Give two changes that would make the student’s estimate more trustworthy.

Q35. SCENARIO. A band posts a poll on its followers’ page: “Should we tour next year?” Of the 200 replies, 196 say yes.

- a. What percentage of the replies say yes?
- b. Describe the population this poll actually reached, and give two reasons why it is not “everyone”.
- c. The band writes “98% want the tour”. Evaluate the claim using the analysis checklist.

Q36. SCENARIO. A media article says: “A new survey shows the town supports the shopping-centre plan. The survey asked 40 people outside the centre on Saturday, and 75% said yes; the result is shown in a bar graph with the vertical axis starting at 70.” Run the full analysis checklist on this claim. Your answer must include:

- the method, and who it over-represents;
- what a sample of 40 can support and what it cannot;
- what the axis choice does to the picture;
- one change to the collection and one change to the display that would make the claim fair;
- a final verdict on what the survey does show.

Marks are given for the strength of the reasoning, not for whether you happen to like shopping centres.

Answers

Q1 a Simple random — every sample of 50 is equally likely; b systematic — random start, fixed interval of 15; c convenience — whoever was easiest to reach. **Q2** a Convenience sampling; b over-represented: people who use the sports ground on Saturday mornings; missed: everyone else; c the result leans towards support — the people asked are the ones who already do sport there. **Q3** a Year-9 students only; b e.g. stratified sampling across year levels — it gives every year level its fair share; c the claim — the arithmetic is fine, the sample does not cover the school. **Q4** a 60%; b sample 1: 75%, sample 2: 50%; c sampling variation — fair samples of 4 vary around the population value. **Q5** a 105 households, 35%; b 40%; c about 57.1% — a landline survey over-represents family households. **Q6** a 24%; b 380 households, 76%; c e.g. the people who answer in working hours are more likely to be at home. **Q7** a $\frac{23}{21} \approx 1.10$; b 2 and 6; c ratio 3 — the graph makes a small difference look large. **Q8** a +10 (from 40 to 50); b -4 (from 40 to 36); c the window Jan–Jun happens to fall, while the whole year rises. **Q9** a 5.125 hours; b median 4 hours, range 12 hours; c the mean supports “very sporty”; the median describes the typical student. **Q10** a $\times 3$; b $\times 9$; c the area grows nine times while the value grows three times. **Q11** a 19 and 4; b samples of 30; c the method is fair — both sets centre near 60 — but one small sample can sit far from the population value. **Q12** a Population: students; method: convenience — students leaving the canteen on Monday; b 20 students — too small to rely on, and only canteen-users on one day; c the method, the sample size, the response rate, and a fair way of choosing students. **Q13** a e.g. stratified by school, so every school is heard in proportion; b e.g. systematic along the rows, or simple random — the plants are all in one place; c e.g. stratified across the city’s areas, or cluster sampling by suburb. **Q14** a Every 20th student; b choose the starting name at random from the first 20 names; c not every sample of 40 is possible — only samples spaced 20 apart. **Q15** a It takes a random sample from each group in proportion to the group’s size, so the younger group keeps its fair share; b the random selection — the collector chooses who fills each place, which lets convenience or judgement back in. **Q16** a Judgement sampling (his own team) — leans towards support; b e.g. a stratified or simple random sample of the whole school; c only that his team wants new uniforms. **Q17** a 8 percentage points; b sampling variation — fair samples differ from run to run; c no — larger samples vary less, so the difference would usually be smaller. **Q18** a 20 percentage points; b each is one fair draw — samples of 10 vary widely around the population value; c the results would cluster closer to 55%, because larger samples vary less. **Q19** a 40 and 16 percentage points; b 25 — the results are much tighter; c the population value sits somewhere near the middle — both averages sit in the 50s. **Q20** a 5 percentage points; b a difference this size can come from sampling variation alone; c larger samples, repeated surveys, or a difference too large for the variation to explain. **Q21** Answers vary: population named, a random method with a reason, a sensible sample size, and a plan for non-response (e.g. follow-up or a different contact method). **Q22** a 3; b $\frac{41}{39} \approx 1.05$; c start the axis at 0 (or clearly show the break) and keep the numbers unchanged. **Q23** a 6 cm; b 2 cm; c steepness $\frac{6}{10}$ then $\frac{2}{10}$ — the steeper drawing serves whoever wants the rise to look dramatic. **Q24** a +1.8 °C; b -0.1 °C; c the claim is true only inside the chosen window — over the full record the temperatures rise; d e.g. 2010 to 2025: +1.2 °C. **Q25** a \$853 000; b \$617 500; c the mean serves whoever wants the street to sound expensive; the median describes the typical house. **Q26** a about 3.78 hours; b median 3 hours, range 11 hours; c the mean — it is pulled up by the one 12-hour student; the median describes the typical student. **Q27** a $A = \pi r^2$, so radius $2r$ gives $\pi(2r)^2 = 4\pi r^2$; b the reader sees a picture four times as big for a park only twice as big; c scale the area, not the radius — draw the second circle’s area as twice the first’s. **Q28** a Convenience sampling (drivers at one intersection in the peak), a cut axis, and a missing sample size; b the sample leans towards commuters, the axis makes a modest majority look commanding, and the small sample is hidden; c report the method and the 30 people asked, draw the axis from 0, and describe the result as drivers at one intersection, not “locals”. **Q29** a 3%; b the people who open the email and care enough to reply lean one way, and most readers never enter the data; c the headline is not supported — 60 self-chosen readers are not “readers”. **Q30** a 90%; b e.g. customers chosen at the counter on a busy sale day (leans towards buyers), or customers who volunteered online (leans towards fans); c how the ten were chosen, when, and how many were asked in all. **Q31** a Self-selection, and reach limited to people listening at that time; b callers chose themselves — angry or keen listeners call, and the rest never enter the data; c e.g. a random sample of listeners contacted through the station’s records, with follow-up of non-responders. **Q32** a Every 20th student; b choose the first student at random from the first 20 names on the roll; c the gate picks by arrival time — latecomers, bus groups and walkers cluster, so the “every 20th” rule runs on a list that was never fair. **Q33** a 6.25%; b students who are keen on the dance — the result leans towards the tastes of enthusiasts; c “of the 25 students who responded, most preferred ...”. **Q34** a 40 percentage points (from 40% to 80%); b samples of 5 vary so widely that almost any result is possible

from one run; c use a larger sample and a random method instead of friends. **Q35** a 98%; b followers who saw the post and chose to reply — people already interested in the band, and only those who respond to its posts; c the arithmetic is honest but the sample is self-selected fans; the claim promotes the tour rather than measuring “everyone”. **Q36** Convenience sample of centre-shoppers on Saturday; a sample of 40 can describe that group on that day, not the town; the cut axis (from 70) makes a modest majority look commanding; fair collection: a random sample of the town; fair display: axis from 0; verdict: the survey shows that centre-shoppers there on Saturday lean towards the plan — nothing more.

Fully-worked solutions

Q1. a Names drawn from a hat give every group of 50 the same chance: simple random sampling. b A random start followed by every 15th name: systematic sampling — the fixed interval is the giveaway. c The first 20 people past the market were simply the easiest to reach: convenience sampling. $\therefore$ a simple random; b systematic; c convenience.

Q2. a The helpers asked whoever turned up at the sports ground: convenience sampling. b Over-represented: residents who use the sports ground on Saturday mornings. Missed: residents who are not there at that time — working, travelling, or simply not sporty. c People already at a sports ground are more likely to like a walking track, so the result leans towards support. $\therefore$ The survey describes Saturday-morning sports-ground users, not the town.

Q3. a The sample contains Year-9 students only; younger and older students walk different distances and have different habits. b Stratified sampling across the year levels would give each year level its fair share; a simple random sample of the whole school roll would also fix it. c The arithmetic — 65% of the 40 asked — is fine; the mistake is claiming the result for students the sample never included. $\therefore$ The claim is not supported, even though the calculation is correct.

Q4. a $\frac{6}{10} = 60\%$ of the club supports the proposal. b Sample 1: $\frac{3}{4} = 75\%$. Sample 2: $\frac{2}{4} = 50\%$. c Both samples were chosen in a fair way; chance gave one a high share and one a low share. Samples of 4 from 10 vary widely around the population value. $\therefore$ Both results are inside the range of fair sampling — that is sampling variation, not a mistake.

Q5. a $60 + 45 = 105$ households have a landline; $\frac{105}{300} = 35\%$ of all households. b $\frac{120}{300} = 40\%$ of all households are family households. c $\frac{60}{105} = \frac{4}{7} \approx 57.1\%$ of landline households are family households. $\therefore$ Families are 40% of the suburb but about 57% of who a landline survey reaches — a landline survey over-represents family households here.

Q6. a $\frac{120}{500} = 24\%$ responded. b $500 - 120 = 380$ households did not respond: $\frac{380}{500} = 76\%$. c For example, people who answer the phone in working hours are more likely to be at home — a group with its own habits and opinions. $\therefore$ With three-quarters of the sample silent, the 24% who answered speak only for themselves unless follow-up shows otherwise.

Q7. a True ratio: $\frac{46}{42} = \frac{23}{21} \approx 1.10$ — the larger value is about one-tenth bigger. b Drawn lengths above the axis start at 40: $42 - 40 = 2$ and $46 - 40 = 6$. c The picture shows the ratio $\frac{6}{2} = 3$. $\therefore$ The graph makes the reader believe the second value is three times the first, when it is really about one-tenth bigger.

Q8. a From the figure: January 40, December 50 — a rise of $50 - 40 = 10$. b January 40, June 36 — a change of $36 - 40 = -4$, a fall of 4. c The window January to June happens to fall, while the full year rises steadily overall. $\therefore$ The poster’s window tells a falling story that the whole year does not tell — the choice of window promotes the point of view.

Q9. a Sum: $2 + 3 + 3 + 4 + 4 + 5 + 6 + 14 = 41$; mean $= \frac{41}{8} = 5.125$ hours. b Ordered data: 2, 3, 3, 4, 4, 5, 6, 14 — the middle two values are 4 and 4, so the median is 4 hours. Range: $14 - 2 = 12$ hours. c The mean (5.125) is pulled up by

the one 14-hour student and supports “our students are very sporty”; the median (4) describes the typical student. $\therefore$ Quoting the mean or the median from the same data promotes different stories.

Q10. a Sales tripled: factor 3. b The picture is 3 times as tall and 3 times as wide, so its area grows by $3 \times 3 = 9$. c The eye reads the whole picture, so the growth looks nine times while the value grew only three times. $\therefore$ The picture overstates the growth by a factor of $9 \div 3 = 3$.

Q11. a Range for samples of 5: $67 - 48 = 19$. Range for samples of 30: $62 - 58 = 4$. b Samples of 30 — their means are far closer together. c Both sets average near 60, so the method is fair and does not lean; but the wide range for size 5 shows that one small sample can sit far from the population value. $\therefore$ Fair method, but small samples vary more — size 30 gives the tighter estimate.

Q12. a Step 1: the population is the school’s students. Step 2: the method is convenience — whoever left the canteen on Monday. b Step 3: 20 students is small, and the sample only contains students who use the canteen on a Monday — everyone else had no chance to be asked. c The poster would need the method, the sample size, the response rate, and a fair way of choosing students — otherwise “students” is a word the data does not support. $\therefore$ As printed, the poster reports canteen-users on one Monday, not the student body.

Q13. a Stratified sampling by school: ten schools means ten groups, and each gets its fair share of the 300. b The plants all live in one greenhouse, so systematic sampling along the rows (or a simple random sample of the 60 plants) is cheap and fair. c Stratified sampling across the city’s areas, or cluster sampling by suburb with a random choice of suburbs — both keep every part of the city in play. $\therefore$ Any answer that names a method and ties it to the population’s structure earns the mark.

Q14. a $800 \div 40 = 20$: take every 20th student. b Choose the starting name at random from the first 20 names on the roll, then count on in 20s. c In a simple random sample every group of 40 students is possible; here only groups spaced exactly 20 names apart can occur. $\therefore$ Every student still has a known chance of selection, which makes the method random — but it is systematic, not simple random.

Q15. a Stratified sampling takes a random sample from each group in proportion to its size, so the younger group keeps its fair share no matter what chance does. b Quota sampling leaves out the random selection: the collector chooses who fills each place, which lets convenience or judgement quietly back in — for example, asking the younger people who are easiest to find. $\therefore$ Stratified sampling is random within each group; quota sampling is not.

Q16. a He asked the people he chose himself — judgement sampling — and his own team is the group most likely to want team uniforms, so the result leans towards support. b A simple random or stratified sample of students across the whole school. c His survey shows that his team wants new uniforms. $\therefore$ “Students want new uniforms” claims a population the sample never reached.

Q17. a $56 - 48 = 8$ percentage points. b Each result is one fair draw from the same population; sampling variation means fair samples differ from run to run. c No — larger samples vary less, so two samples of 250 would usually sit closer to each other and to the population value. $\therefore$ A gap of 8 percentage points between small fair samples is ordinary, not evidence of a mistake or a real split.

Q18. a Range: $60 - 40 = 20$ percentage points. b Each sample was a fair draw from a population where 55% support the proposal; with samples of only 10, results scatter widely around 55% and rarely land on it exactly. c Samples of 100 vary less than samples of 10, so the results would cluster closer to 55%. $\therefore$ None of the three results is wrong — they are three fair draws, and small samples vary a lot.

Q19. a Range for size 5: $80 - 40 = 40$ percentage points. Range for size 25: $60 - 44 = 16$ percentage points. b Size 25 — its results are much tighter around their centre, so one run of it is a better estimate. c The means of the two sets are 56% and 52% — both sit in the 50s, so the population value is likely somewhere near there. $\therefore$ Bigger samples spread less; the averages of many fair samples point back at the population value.

Q20. a $52 - 47 = 5$ percentage points. b Both surveys use samples; sampling variation alone can move a result by 5 percentage points either way, so the two regions may in truth hold the same view. c Larger samples in both regions, repeated surveys that keep showing the gap, or a gap too large for the variation to explain. $\therefore$ The conclusion “the regions differ” needs more evidence than two sample results 5 points apart.

Q21. Marking guide — a full answer contains all four parts: 1. A named population (for example, all students at the school). 2. A random method with a reason — e.g. stratified sampling across year levels so every year level keeps its fair share. 3. A sample size with a reason — large enough to vary little, small enough to run (for example, 100 of a 1000-student school). 4. A plan for non-response — follow-up contact, or a second contact method, so the silent students are not simply lost.

Q22. a Drawn lengths: $78 - 76 = 2$ and $82 - 76 = 6$; ratio $\frac{6}{2} = 3$. b True ratio: $\frac{82}{78} = \frac{41}{39} \approx 1.05$. c Draw the same two bars with the vertical axis starting at 0 — or show the break clearly — and keep the numbers unchanged. $\therefore$ The cut axis makes a 5% difference look like a tripling; the honest graph shows two bars of almost equal height.

Q23. a At 1 unit per centimetre, a rise of 6 units is drawn 6 cm tall. b At 3 units per centimetre, the same rise is drawn $\frac{6}{3} = 2$ cm tall. c Steepness over the 10 cm run: $\frac{6}{10}$ then $\frac{2}{10}$. Nothing changed except the scale — but the first drawing looks three times as steep. $\therefore$ The steeper drawing serves whoever wants the rise to look dramatic; the numbers support no such drama.

Q24. a Full record: $14.6 - 12.8 = 1.8$ °C rise. b Window 1995 to 2010: $13.4 - 13.5 = -0.1$ °C — almost flat, slightly down. c Inside the chosen window the claim is true; over the full record the temperatures clearly rise. The window was chosen because it tells that story. d The window 2010 to 2025 gives $14.6 - 13.4 = 1.2$ °C — a rising story from the same data. $\therefore$ Windows can tell either story here; the honest display shows the whole record and lets the reader see both.

Q25. a Sum: $580 + 590 + 600 + 610 + 615 + 620 + 625 + 640 + 650 + 3000 = 8530$; mean = $\frac{8530}{10} = 853$ — \$853 000.

b Ordered prices: the middle two are 615 and 620; median = $\frac{615 + 620}{2} = 617.5$ — \$617 500. c The one \$3 000 000 house pulls the mean far above nine of the ten prices; the mean serves whoever wants the street to sound expensive. $\therefore$ The median, \$617 500, describes the typical house on the street.

Q26. a Sum: $1 + 2 + 2 + 3 + 3 + 3 + 4 + 4 + 12 = 34$; mean = $\frac{34}{9} \approx 3.78$ hours. b Ordered data: 1, 2, 2, 3, 3, 3, 4, 4, 12 — the 5th value is 3, so the median is 3 hours. Range: $12 - 1 = 11$ hours. c The mean (about 3.78) is pulled up by the one 12-hour student, so it serves the “we drown in homework” story; the median (3) describes the typical student. $\therefore$ Same data, two measures, two stories — the choice of statistic is a choice of representation.

Q27. a Area of a circle: $A = \pi r^2$. With radius $2r$: $\pi(2r)^2 = 4\pi r^2 = 4A$. b The bigger circle covers four times the area on the page, so the reader sees “four times” while the park is only twice as big. c Scale the area, not the radius: draw the second circle with area $2A$ — that is, radius $r\sqrt{2}$ — or use a bar, where only the height grows. $\therefore$ Picture graphs must scale area in proportion to value, or the picture does the exaggerating.

Q28. a Three techniques: convenience sampling (30 drivers at one intersection in the morning peak), a cut axis (starting at 60), and a missing sample size. b The sample leans towards people who drive through that intersection at peak time — exactly the commuters the road would serve; the cut axis makes a modest majority look commanding; and the reader cannot weigh a result of 30 because the number is hidden. c State the method and the 30 people asked, draw the axis from 0, and describe the result as “68% of the 30 drivers surveyed at one intersection”, not “locals”. $\therefore$ The infographic as printed promotes the road; the redesigned version lets the reader judge.

Q29. a $\frac{60}{2000} = 3\%$ responded. b First, the people who open the email and care enough to reply lean towards strong feelings about the proposal. Second, the other 97% never enter the data at all, so their views are simply absent. c “Readers say no” speaks for 2000 readers on the evidence of 60 self-chosen ones. $\therefore$ The headline is not supported; an honest version reads “of the 60 readers who responded, most said no”.

Q30. a $\frac{9}{10} = 90\%$. b For example: the ten were customers at the counter on a busy sale day — buyers only, which leans towards praise; or the ten volunteered online — fans only, which leans the same way. Either method picks

people who already like the product. c How the ten were chosen, when they were asked, and how many customers were asked in all. $\therefore$ “9 out of 10” is honest arithmetic; without a method it cannot support a claim about customers in general.

Q31. a Self-selection — callers choose themselves; and reach limited to whoever is listening at that hour. b People who phone in about a parking rule are the ones with strong feelings about it; the quiet majority of listeners never enters the data, so the calls are a loud sample, not a fair one. c Contact a random sample of listeners through the station’s records and follow up the ones who do not respond. $\therefore$ The calls measure the callers’ feelings, not the listeners’ views.

Q32. a $800 \div 40 = 20$: every 20th student. b Choose the first student at random from the first 20 names on the roll, then take every 20th name after that. c The gate orders students by arrival time, not by a fair list: latecomers, bus groups and walkers arrive in clusters, so “every 20th through the gate” is systematic sampling run on an order that already leans. $\therefore$ A systematic sample needs a random start on a fair list — the roll, not the gate.

Q33. a $\frac{25}{400} = 6.25\%$ responded. b Students who are keen on the dance are the ones who bother to respond; the result leans towards enthusiasts’ tastes. c “Of the 25 students who responded, most preferred ...” $\therefore$ With 93.75% silent, the poll describes the responders, not the students.

Q34. a In Question 19 the five samples of size 5 gave 40%, 60%, 40%, 80% and 60% — a spread of $80 - 40 = 40$ percentage points. b Results from samples of 5 vary so widely that 80% is one ordinary draw — it can come up even when the true support in the population is far lower. c Use a much larger sample, and choose the people randomly instead of asking friends (friends share the student’s views on purpose). $\therefore$ One sample of 5 supports almost nothing; size and method together are what make an estimate trustworthy.

Q35. a $\frac{196}{200} = 98\%$. b The poll reached followers who saw the post and chose to reply: people already interested in the band, and only the ones who respond to its posts. Non-followers, and followers who scroll past, never enter the data. c Checklist: the population claimed is “everyone who might attend”; the population reached is self-selected fans. 200 replies is a decent count, but self-selection means the 98% was predictable before the poll ran. $\therefore$ The arithmetic is honest, but “98% want the tour” promotes the tour; it does not measure a population beyond the fans.

Q36. Step 1 — the population claimed is the town; the people reached were shoppers outside the centre on Saturday. Step 2 — the method is convenience sampling; it over-represents people who already shop there on weekends — the very group the plan would please. Step 3 — 40 people is small enough to vary widely, and a convenience sample of 40 can describe that group on that day; it cannot support a claim about the town. Step 4 — the bar graph’s axis starts at 70, so the “yes” bar towers over 70 and the majority looks commanding; $\frac{30}{40} = 75\%$ yes leaves 25% — a quarter — against, which the picture hides. Step 5 — fair collection: a random sample of the town’s residents; fair display: the axis drawn from 0, with the sample size and method printed. Step 6 — the representation serves whoever wants the centre plan approved. $\therefore$ Verdict: the survey shows that shoppers outside the centre on that Saturday leaned towards the plan — and nothing more than that.

Planning and conducting a statistical investigation

6E · Chapter 6 — Statistics · Level 9

Content description VC2M9ST05 (word for word): *plan and conduct statistical investigations involving the collection and analysis of different kinds of data; report findings and discuss the strength of evidence to support any conclusions*

Achievement standard anchor (AS 17): *They explain how sampling techniques and representation can be used to support or question conclusions or to promote a point of view.*

Elaboration ids for linkage: VC2M9ST05.E01, VC2M9ST05.E02, VC2M9ST05.E03, VC2M9ST05.E04

Prerequisite pointer: Level 8 Statistics is assumed — population vs sample and collection techniques (VC2M8ST01), distributions from primary and secondary sources with random and non-random sampling (VC2M8ST02), variation between random samples of the same size and the effect of sample size (VC2M8ST03), and investigations with samples and fair inference (VC2M8ST04) — and 6A–6D of this chapter: the investigation cycle runs on the whole chapter’s tools. The method names are used here, not re-taught.

Note on VC2M9ST05.E04 (exclusion recorded): the Aboriginal and Torres Strait Islander material of this elaboration is excluded from this book under CONTENT_POLICY Rule 1 (D10; recorded as ACCEPTED here, since TOC.md must not be edited from inside a section). The content description is delivered in full via the remaining elaborations and the complete investigation cycle below.

Note on sources: the weather data in this section is real — Bureau of Meteorology daily observations for July 2026, retrieved into this tree on 20 August 2026 — and carries the source lines R1 requires; every other quantitative context is a clearly-labelled SCENARIO with invented parameters (R2).

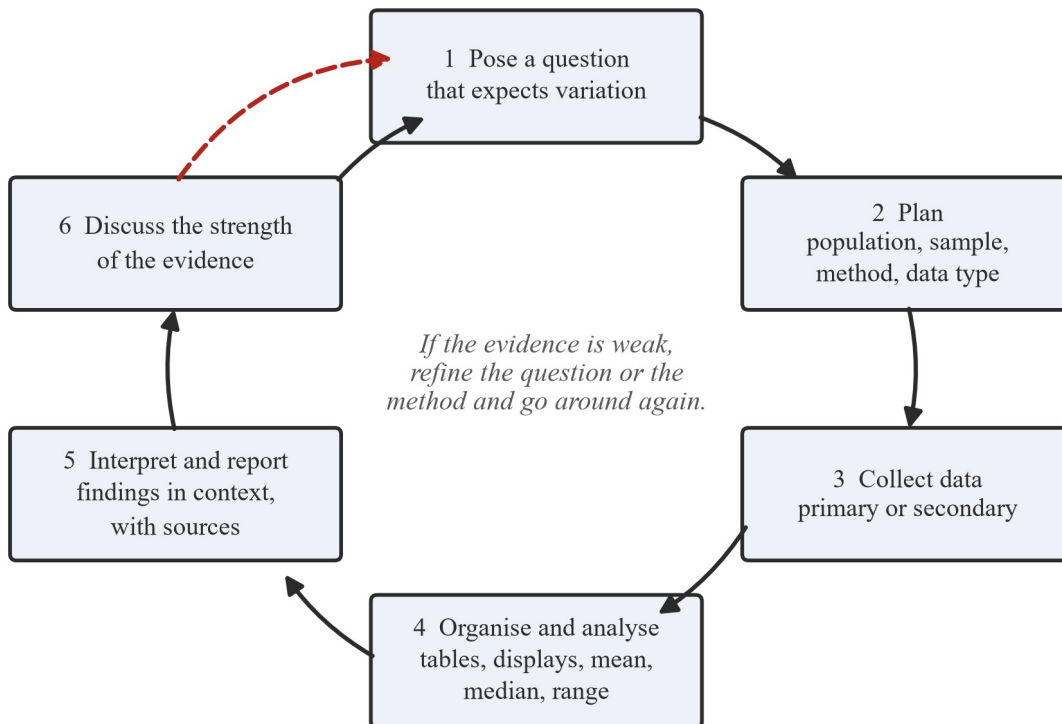
Theory

The investigation cycle

A statistical investigation is not a straight line; it is a loop you go around as many times as the evidence needs. The six steps, in order, are:

1. **Pose a question** that expects variation — the answers will differ, and that variation is what you study.
2. **Plan** — name the population, the sample, the method, and the kind of data the question needs.
3. **Collect** — primary data you gather yourself, or secondary data someone else has published.
4. **Organise and analyse** — tables and displays, then the mean, median and range.
5. **Interpret and report** — say what the numbers show, in context, with the sources named.
6. **Discuss the strength of the evidence** — how far the findings can be trusted, and whether the conclusion should go that far.

The statistical investigation cycle



Six boxes joined in a loop: pose a question, plan, collect, organise and analyse, interpret and report, discuss the strength of the evidence, with a red dashed arrow looping from the last box back to the first

The statistical investigation cycle. If the evidence is weak, refine the question or the method and go around again.

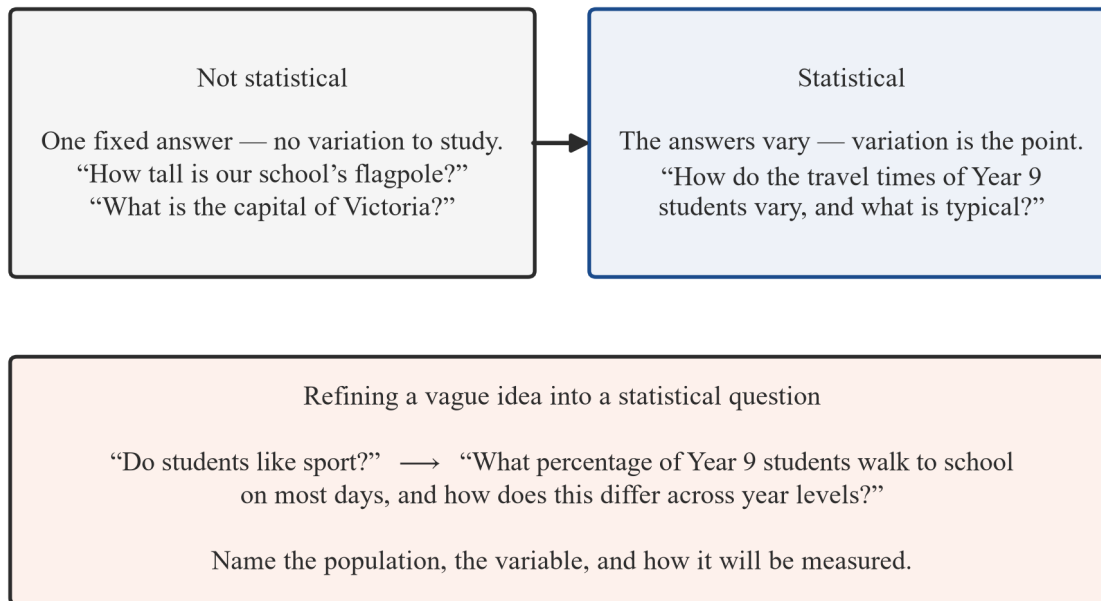
Step 6 is what separates a statistical report from a story. The rest of this section walks once around the loop.

Posing a statistical question

A question is statistical when its answers vary across the people or things you measure, and you care about the pattern in that variation. “What is the capital of Victoria?” is not statistical — one fixed answer. “How do the travel times of Year 9 students vary, and what is typical?” is statistical — the answers differ, and the variation is the point.

A vague idea is not yet a question. To refine it, name three things: the **population** (who or what you want to say something about), the **variable** (what you will record for each member), and **how it will be measured** (the exact question or measurement, so two people collecting would get the same kind of answer).

Statistical questions expect variation



Two panels comparing a non-statistical question with one fixed answer to a statistical question whose answers vary, above a panel refining the vague idea “do students like sport?” into a question that names the population, the variable and how it is measured

Statistical questions expect variation.

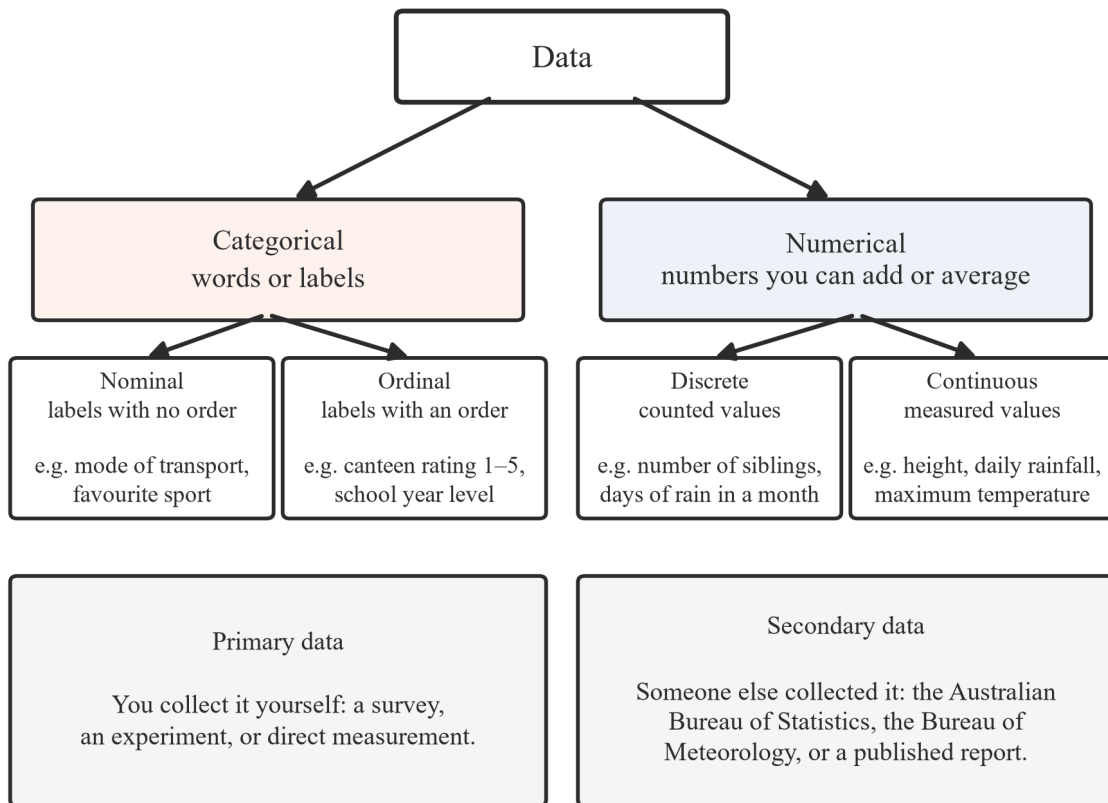
The kinds of data a question needs

The variable you choose decides the kind of data you collect, and the kind of data decides the displays and statistics that are honest to use.

Categorical data are words or labels. Labels with no order are **nominal** (mode of transport, favourite sport); labels with an order are **ordinal** (a canteen rating of 1–5, year level). **Numerical** data are numbers you can add or average. Counted values are **discrete** (number of siblings, days of rain in a month); measured values are **continuous** (height, daily rainfall, maximum temperature).

Data are also **primary** (you collect them yourself — a survey, an experiment, direct measurement) or **secondary** (someone else collected them and you reuse them). This book’s default secondary sources are the free online databases of the Australian Bureau of Statistics and the Bureau of Meteorology: large, carefully collected, and free to cite.

The kinds of data an investigation can use



A tree: data splits into categorical (nominal, ordinal) and numerical (discrete, continuous), with two panels below for primary and secondary data

The kinds of data an investigation can use.

Planning the collection

The plan is where most investigations are won or lost. Name the **population** (the whole group the question is about), then decide between a **census** (every member — only possible for small populations) and a **sample** (a part). If a sample, name the **method** — random methods (simple random, systematic, stratified, cluster) support claims about the population; non-random methods (quota, convenience, judgement) are cheap and quick but lean, and the lean must be named. Decide the **sample size** — bigger samples vary less. And write down, before collecting, what could go wrong: who the method cannot reach, who might not answer, and what might make people answer untruthfully.

Collecting data from people carries duties. Participation must be optional; names are not needed for a class or school survey, so do not ask for them; and the results should be reported as a group, never as individuals. A question that people cannot answer honestly is not worth asking.

A **planning frame** holds the plan on one page:

A planning frame, worked for one question

Question	Are teenagers prepared to spend more on technology than on clothing in a typical month?
Population	All Year 9 students at the school (about 240).
Sample and method	Stratified sample of 60 students: 15 chosen at random from each Year 9 class.
What is recorded	Usual monthly spend on technology and on clothing, in dollars (numerical, continuous).
Display and statistics	Back-to-back stem-and-leaf plot; mean, median and range of each set.
Bias to watch	Students may round their spending up, or answer the way they think the surveyor wants.

A worked planning frame in six rows: question, population, sample and method, what is recorded, display and statistics, bias to watch, filled in for a survey of Year 9 monthly spending

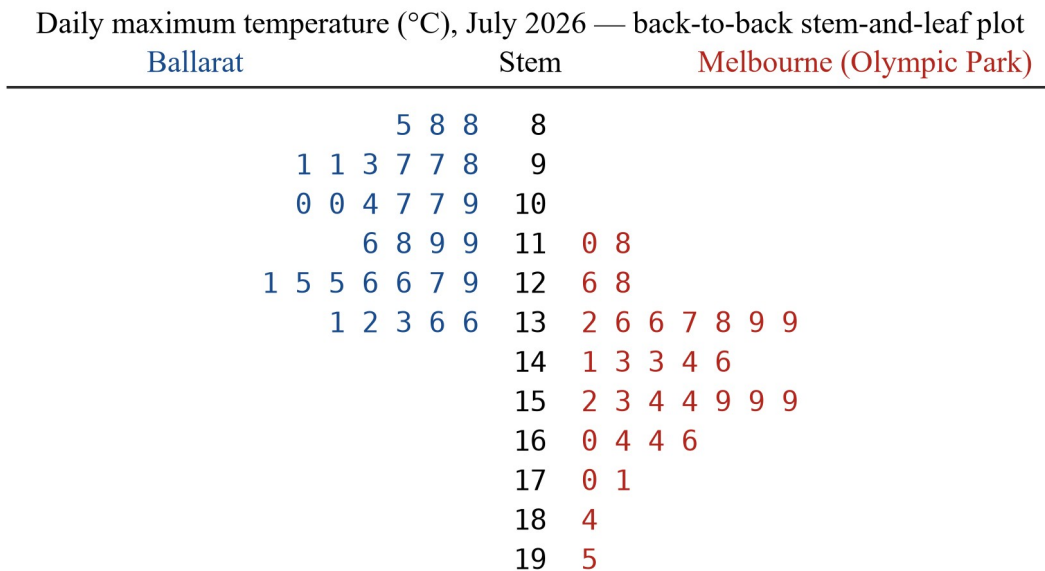
A planning frame, worked for one question.

Organising and analysing

Match the display to the data: a bar graph for categorical data; a dot plot, stem-and-leaf plot or similar for numerical data. Then summarise with the three statistics this level uses — the **mean** (the total shared out), the **median** (the middle value in order), and the **range** (largest minus smallest).

The weather data below are real. For every day of July 2026, the Bureau of Meteorology recorded the daily maximum temperature at Ballarat Aerodrome and at Melbourne Olympic Park. A back-to-back stem-and-leaf plot keeps every value and shows the shape of each set side by side.

Source: Bureau of Meteorology, *Daily Weather Observations, Ballarat Aerodrome, station 089002 (product IDCJDW3005.202607), July 2026*, prepared 14 August 2026. Retrieved 20 August 2026. Source: Bureau of Meteorology, *Daily Weather Observations, Melbourne Olympic Park, station 086338 (product IDCJDW3033.202607), July 2026*, prepared 14 August 2026. Retrieved 20 August 2026.



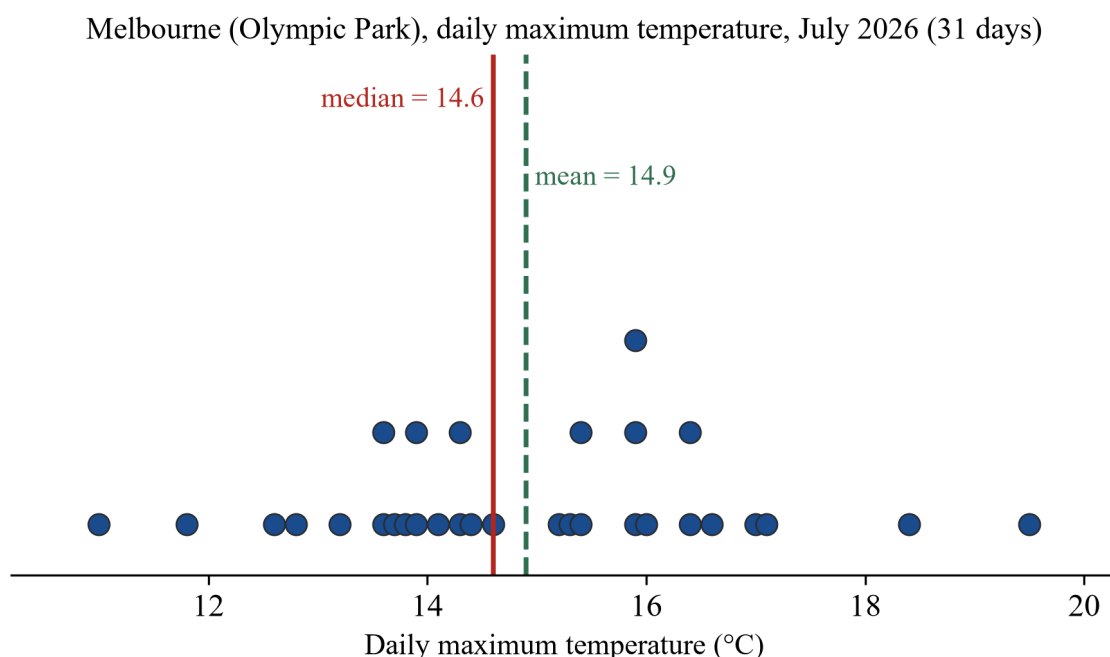
Leaf = tenths of 1°C . Key: $11 | 6$ means 11.6°C .

Source: Bureau of Meteorology daily weather observations, July 2026.

Back-to-back stem-and-leaf plot of daily maximum temperatures for July 2026, Ballarat leaves on the left and Melbourne leaves on the right, stems 8 to 19

Daily maximum temperature ($^{\circ}\text{C}$), July 2026. Leaf unit: 0.1°C ; the stem 12 with leaf 5 is 12.5°C .

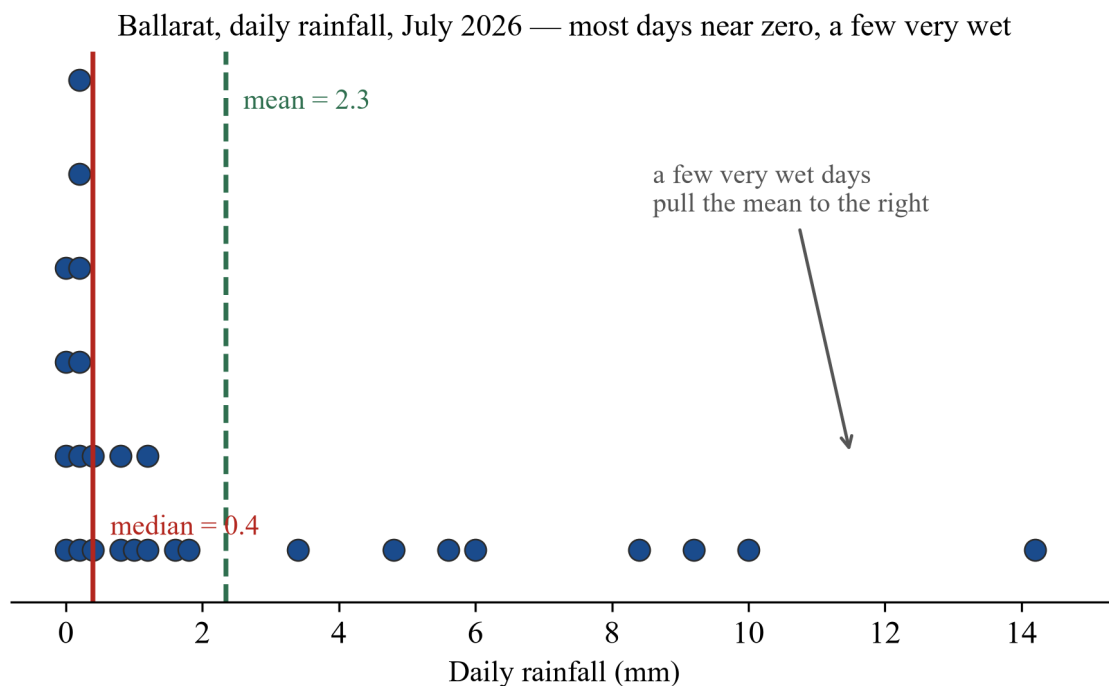
The Melbourne set sits higher — its middle is at 14.6°C against Ballarat's 11.6°C — and it also spreads more (range 8.5°C against 5.1°C). A dot plot with the median and mean marked shows the same story for one city:



Dot plot of Melbourne daily maximum temperatures for July 2026 with the median at 14.6 and the mean at 14.9 marked by vertical lines

Melbourne (Olympic Park), daily maximum temperature, July 2026. The mean sits a little right of the median.

Watch what shape does to the mean. Ballarat’s daily rainfall for the month is mostly near zero, with a few very wet days:



Dot plot of Ballarat daily rainfall for July 2026, most days stacked near zero with a long tail to 14.2 mm; the median line at 0.4 sits far left of the mean line at 2.3

Ballarat, daily rainfall, July 2026 — most days near zero, a few very wet. The wet days pull the mean to the right of the median.

The median day was 0.4 mm, yet the mean is 2.3 mm — nearly six times the median — because the few wet days pull the mean their way. When data lean like this, the median is the safer “typical” value to report, and saying so is part of the analysis.

Reporting findings

A statistical report has seven parts, and a reader should be able to repeat the investigation from what you write.

The parts of a statistical report

Title and question	the statistical question, worded exactly
Method and data source	population, sample, method, and where the data came from, with dates
Analysis	displays, and the mean, median and range
Findings	what the numbers show, in context, with units
Limitations	what the method could not capture
Strength of the evidence	how far the findings can be trusted
Conclusion	the answer to the question — no further than the evidence allows

A reader should be able to repeat the investigation from your method.

Table of the seven parts of a statistical report: title and question; method and data source; analysis; findings; limitations; strength of the evidence; conclusion

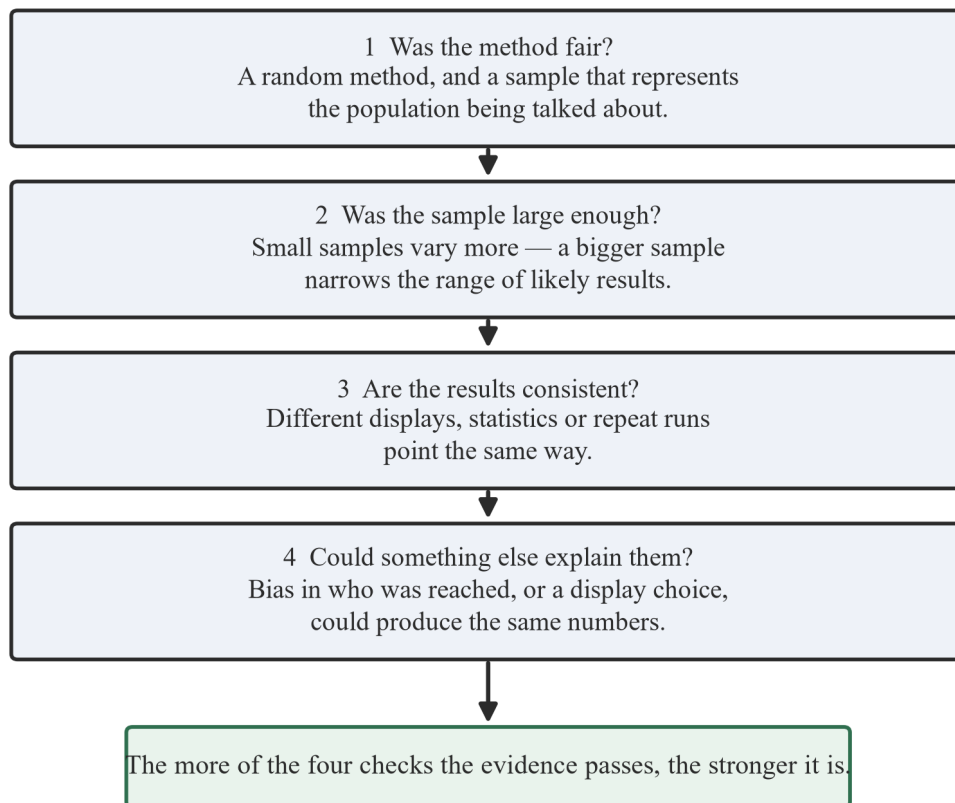
The parts of a statistical report.

Two of those parts are where reports most often fail. **Findings** must stay in context, with units — not “the mean is 11.2” but “the mean daily maximum temperature at Ballarat in July 2026 was 11.2 °C”. And the **method and data source** part must name where secondary data came from, with dates, the way the source lines above this section do.

The strength of the evidence

The last step of the cycle is the content description’s verb *discuss*. Every finding comes with a strength of evidence — how far it can be trusted — and the conclusion must not run further than that strength allows. Four questions test it:

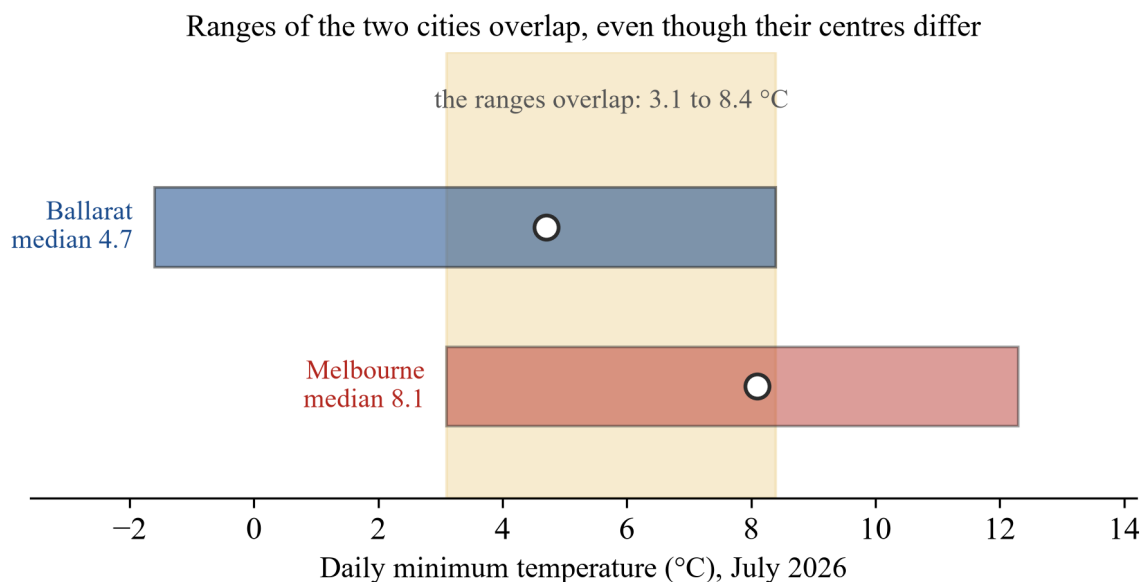
Four questions that test the strength of evidence



Four stacked boxes: was the method fair? was the sample large enough? are the results consistent? could something else explain them? above a green strip reading “the more of the four checks the evidence passes, the stronger it is”

Four questions that test the strength of evidence.

Sample size is the second check, and sampling variation is why it matters: small samples swing widely, big samples settle down. It also shapes comparisons. The two cities’ daily **minimum** temperatures for July 2026 differ in the middle (medians 4.7 °C against 8.1 °C) — yet their ranges overlap from 3.1 °C to 8.4 °C, so a cold Melbourne morning and a mild Ballarat morning can look identical. A difference between centres is a fact about the groups; it does not mean every member of one group beats every member of the other.



Two horizontal range bars for daily minimum temperature in July 2026, Ballarat from minus 1.6 to 8.4 and Melbourne from 3.1 to 12.3, with the overlap from 3.1 to 8.4 shaded and each median marked

Ranges of the two cities overlap, even though their centres differ.

That is exactly the thinking the achievement standard anchor asks for: sampling techniques and representation can support **or question** a conclusion. A fair method, a big enough sample, consistent results and no other likely explanation make evidence strong; a convenience sample of ten, a voluntary response, or a display chosen to flatter makes it weak — and naming which is the work of step 6.

Evaluate and refine

If the evidence is weak, the loop closes by going around again: refine the question, fix the method, take a bigger or fairer sample, or choose an honest display. A report that says “the evidence is weak, and here is why” is not a failed investigation — it is a complete one.

Worked examples

WE1 (scaffolded) — from a vague idea to a statistical question

Refine the idea “students and sport” into a statistical question.

Working	Comment
“Do students like sport?”	As written: no population named, “like” cannot be recorded, and the answer yes/no would hide the variation.
Population: all Year 9 students at this school. Variable: whether the student plays or watches sport outside school on most weeks, and which code.	Naming the population and the variable turns an idea into something measurable.
“What percentage of Year 9 students at this school play a sport outside school on most weeks, and which code is most common?”	Now statistical: the answers vary across students, and the two summaries answer the question.

∴ A statistical question names the population, the variable and how it is measured — and expects variation.

WE2 (scaffolded) — filling in a planning frame

Plan the investigation of the planning-frame question shown in the theory: monthly spending on technology and on clothing by Year 9 students.

Working	Comment
Population: all Year 9 students at the school (about 240).	The question is about Year 9, so the population is Year 9 — not the whole school.
Sample and method: stratified sample of 60 students — 15 chosen at random from each Year 9 class.	Four classes of about 60; 15 from each keeps every class in proportion (60 of 240 is one quarter). Random choice inside each class keeps the method fair.
What is recorded: usual monthly spend on technology and on clothing, in dollars.	Numerical, continuous — so means, medians and ranges will make sense.
Display and statistics: back-to-back stem-and-leaf plot; mean, median and range of each set.	One display compares the two spends; the three statistics summarise each.
Bias to watch: students may round their spending up, or answer the way they think the surveyor wants.	Names the lean before collecting, so the report can discuss it.
$\therefore$ The frame fixes the population, the fair method, the data type, the analysis and the risks before a single value is collected.	

WE3 (scaffolded) — comparing two real data sets

Use the July 2026 maximum-temperature data (sources in the theory) to compare the two cities.

Working	Comment
$n = 31$ for each city — every day of the month.	A full month is a census of July 2026 days, not a sample.
Ballarat: mean 11.2°C , median 11.6°C , range 5.1°C .	Read from the stem-and-leaf plot's left side.
Melbourne: mean 14.9°C , median 14.6°C , range 8.5°C .	Read from the right side.
Mean difference $14.9 - 11.2 \approx 3.7^{\circ}\text{C}$; median difference $14.6 - 11.6 = 3.0^{\circ}\text{C}$.	Both centres agree on the direction: Melbourne ran about $3\text{--}4^{\circ}\text{C}$ warmer.
$\therefore$ In July 2026, Melbourne's daily maximum temperatures ran about $3\text{--}4^{\circ}\text{C}$ warmer than Ballarat's on average, and varied more (range 8.5°C against 5.1°C).	

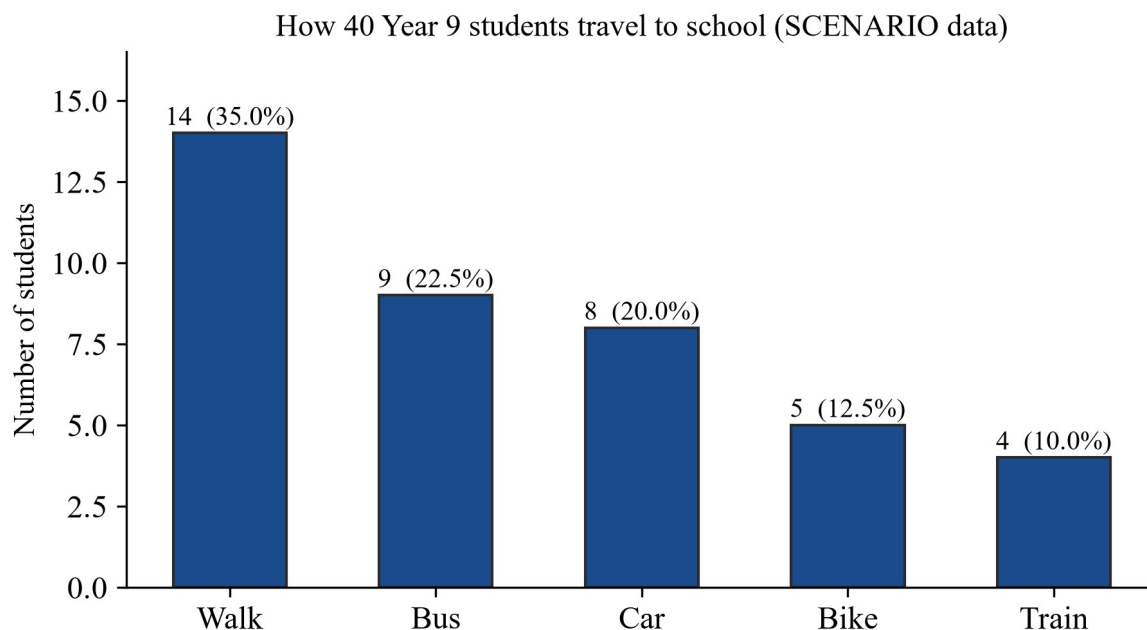
WE4 (scaffolded) — two surveys, one question, different strength

Two groups ask the same question — “Should the school day start half an hour later?” — by different methods. Group A asks 10 friends at the gate; 4 say yes. Group B takes a stratified sample of 60 across all year levels; 36 say yes. Compare the evidence.

Working	Comment
Group A: $\frac{4}{10} = 40\%$. Group B: $\frac{36}{60} = 60\%$.	Same question, results 20 percentage points apart.
Check 1, method: A is convenience (friends at one gate, one morning); B is stratified and random.	A's method leans toward one friendship group; B's represents the school.
Check 2, size: 10 against 60.	Small samples vary more; 10 is far too small to settle a school-wide question.
Checks 3 and 4: the results are not consistent with each other, and A's method gives an obvious other explanation (who was asked).	The 20-point gap is explained by method, not by a change of minds.
$\therefore$ Group B's 60% is the stronger evidence; Group A's 40% passes almost none of the four checks, so it supports no claim about the school.	

WE5 — reading a categorical display

A Year 9 class surveys 40 students: “How do you usually travel to school?” The results (SCENARIO data) are 14 walk, 9 bus, 8 car, 5 bike, 4 train.



Bar graph of how 40 Year 9 students travel to school: walk 14 (35 percent), bus 9 (22.5 percent), car 8 (20 percent), bike 5 (12.5 percent), train 4 (10 percent)

How 40 Year 9 students travel to school (SCENARIO data).

Find the share of each mode and state what the data show.

Walk: $\frac{14}{40} = 35\%$; bus: $\frac{9}{40} = 22.5\%$; car: $\frac{8}{40} = 20\%$; bike: $\frac{5}{40} = 12.5\%$; train: $\frac{4}{40} = 10\%$ — the shares add to 100%, as they must.

∴ Walking is the most common mode (35%), and the two active-and-public modes walk and bus together carry most of the students in this sample.

WE6 — judging a weak claim

A canteen committee wants to know what students spend at the canteen each week. It puts a sheet by the door and 12 volunteers write down their spend. The mean is \$45. The committee reports: “Students spend \$45 a week at the canteen.” The total the 12 recorded is $12 \times 45 = \$540$. Discuss the strength of the evidence.

The method is voluntary response — people who choose to record their spend are likely the bigger spenders, and the sheet sits where canteen-users see it. Twelve is a small number, so the mean will swing widely from one group of 12 to the next; and one big spender among 12 pulls the mean hard. The report names no limitation.

∴ The \$45 mean is weak evidence about all students: voluntary response leans it upward, the sample is small, and the claim runs further than the evidence allows.

WE7 — a mini-investigation with an outlier

A class tests one paper-plane design with 15 throws in the hall (SCENARIO data). Distances in metres: 3.2, 3.5, 3.8, 4.0, 4.1, 4.2, 4.4, 4.5, 4.6, 4.8, 5.0, 5.2, 5.5, 6.1, 9.7. Summarise with and without the last throw, and discuss.

With all 15: mean ≈ 4.84 m, median 4.5 m, range 6.5 m. The 9.7 m throw is far beyond the rest — an outlier.
Without it ($n = 14$): mean ≈ 4.49 m, median 4.45 m, range 2.9 m.

∴ One odd throw moved the mean by about 0.35 m and more than doubled the range, while the median moved only 0.05 m — so the report should give the median, name the outlier, and treat 4.4–4.5 m as the design’s typical distance.

WE8 — rewriting a conclusion that runs too far

A group surveys 40 students at one gate and finds 23 walk to school. Its report ends: “Most students at our school walk to school.” Rewrite the conclusion so it does not run further than the evidence.

The sample is a convenience sample at one gate, at one time — students who walk may pass that gate more than students who take the bus from other streets. 23 of 40 is 57.5%... but the honest move is to shrink the claim to what the method can carry.

∴ “In this sample of 40 students at one gate, 23 walked to school. Because the sample was taken at one gate at one time, it may not represent the school, so the result is a first estimate — a stratified sample across all gates and year levels would test it.”

Practice questions

Q1 (scaffolded) Statistical or not? For each, say which and why. (a) What is the capital of Victoria? (b) How did the daily maximum temperatures of July 2026 vary between Ballarat and Melbourne? (c) How tall is the school flagpole?

Q2 (scaffolded) Refine the idea “is the canteen any good?” into a statistical question. (a) Name a population and a variable. (b) Write the refined question, including how the variable will be measured.

Q3 (scaffolded) Classify each variable as categorical (nominal or ordinal) or numerical (discrete or continuous). (a) Favourite sport. (b) Canteen rating from 1 to 5. (c) Number of siblings. (d) Daily rainfall in millimetres.

Q4 (scaffolded) Primary or secondary? Say which, and name the kind of data. (a) You measure the reaction times of your classmates with a ruler drop. (b) You download the Bureau of Meteorology’s daily rainfall for July 2026.

Q5 (scaffolded) SCENARIO: a swimming club surveys 80 of its 340 members about training times. (a) Name the population and its size. (b) Name the sample and its size, and say whether this is a census.

Q6 (scaffolded) A school has 240 Year 9 students in four classes. A group wants 60 students, with every class represented, choosing 15 at random from each class list. (a) Name the sampling method. (b) Say whether it is random or non-random, and why that matters.

Q7 (scaffolded) Ballarat’s daily rainfall for July 2026 (real data; source in the theory) totals 72.8 mm over 31 days. (a) Find the mean, and confirm the median 0.4 mm and range 14.2 mm from the dot plot in the theory. (b) Four days recorded 0.0 mm. What percentage of the month was dry?

Q8 (scaffolded) Melbourne’s daily maximum temperatures for July 2026 have mean 14.9 °C, median 14.6 °C and range 8.5 °C; Ballarat’s are 11.2 °C, 11.6 °C and 5.1 °C. (a) Find the difference between the two medians. (b) Which city varied more, and by how much range?

Q9 (scaffolded) Ballarat’s July 2026 rainfall has mean 2.3 mm but median 0.4 mm. (a) Which statistic is pulled by the few very wet days? (b) Which is the safer “typical day” value to report, and why?

Q10 (scaffolded) SCENARIO: three samples of 10 students give results 40%, 70%, 50% on a question; three samples of 200 give 54%, 58%, 56%. (a) Find the range of the three small-sample results and of the three large-sample results. (b) What does the comparison show about sample size?

Q11 (scaffolded) A group plans to survey students about how much sleep they get on school nights. (a) Name two things the plan must include so that taking part is okay. (b) Explain why each matters.

Q12 (scaffolded) Choose an honest display, and say why it suits the data. (a) Favourite sport of 50 students (categorical, nominal). (b) The 31 daily maximum temperatures of a month (numerical, continuous).

Q13 Write two different statistical questions about how students in your year get to school. Each must name the population and the variable.

Q14 Nominal or ordinal? (a) The colour of a car. (b) A medal: gold, silver or bronze. (c) A satisfaction rating from 1 to 5. (d) The suburb a student lives in.

- Q15** Discrete or continuous? (a) Goals kicked in a match. (b) A student's height. (c) The number of wet days in a month. (d) Walking time to school, in minutes.
- Q16** Primary or secondary, and why? (a) The Australian Bureau of Statistics' online table of population by age. (b) Hand spans you measure across your class. (c) Bureau of Meteorology daily observations for last month. (d) A survey of lunch orders you run in your class.
- Q17** A group asks its own friends whether the uniform should change. Explain why the result will lean, and name the sampling method.
- Q18** SCENARIO: a sports website runs an open poll — “Should the grand final weekend be a public holiday?” — and 87% of 12 000 votes say yes. Give two reasons this is weak evidence about Australians generally.
- Q19** Melbourne's daily rainfall for July 2026 (real data; source in the theory): find the mean (1 d.p.), and confirm the median 0.0 mm and range 14.6 mm. Sixteen days recorded 0.0 mm — what percentage of the month was dry?
- Q20** Compare the two cities' July 2026 rainfall totals: Ballarat 72.8 mm, Melbourne 44.4 mm. Find the difference and the ratio (1 d.p.), and state in one sentence which city had the wetter month.
- Q21** A dot plot shows 12 values: 2, 3, 3, 4, 4, 4, 5, 5, 6, 6, 7, 8. Find the mean, median and range.
- Q22** Ten students report hours spent gaming in a week: 4, 5, 5, 6, 6, 6, 7, 7, 8, 40. Find the mean, median and range; then the same three without the 40. Which statistics did the outlier not change?
- Q23** SCENARIO: 50 students name their favourite sport: AFL 18, soccer 12, basketball 9, cricket 6, tennis 5. What percentage chose AFL?
- Q24** SCENARIO: 50 students rate the canteen from 1 to 5: the counts are 2, 5, 14, 17, 12. Find the most common rating and the percentage who gave it.
- Q25** SCENARIO: a headline reads “80% of students hate the canteen”, based on 20 replies to an angry post. Give two reasons the evidence is weak.
- Q26** SCENARIO: two fair random samples of the same size ask the same question and get 48% and 52%. One group calls this a contradiction. Explain why it is not, using sampling variation.
- Q27** Compare the two cities' daily **minimum** temperatures for July 2026 (real data): Ballarat mean 3.7 °C, median 4.7 °C, range 10.0 °C; Melbourne mean 7.7 °C, median 8.1 °C, range 9.2 °C. Find the difference between the means (1 d.p.), and use the range-bar figure in the theory to explain why the two cities can still have identical mornings.
- Q28** For the minimum-temperature comparison of Q27, say which display the range-bar figure is doing the work of, and why a bar graph of the two means alone would be a poorer choice.
- Q29** Write the R1 source line for the Ballarat July 2026 temperature data, including the station number, the product code, the preparation date and the retrieval date.
- Q30** SCENARIO: a year level of 250 students is surveyed by email; 55 reply. Find the response rate as a percentage, and name the risk the non-responders bring.
- Q31** SCENARIO: a survey sheet asks for a student's name, home address and answers. Say what should change before the survey runs, and why.
- Q32** Read from the back-to-back stem-and-leaf plot in the theory: (a) how many Ballarat days had a maximum of 13 °C or more; (b) how many Melbourne days reached 18 °C or more; (c) which city's maximums varied more, and by how much range.
- Q33** The mean difference between the cities' July 2026 maximum temperatures is about 3.7 °C. Discuss the strength of the evidence for the claim “Melbourne July days are warmer than Ballarat's”: what does a full month of data make strong, and what can it not support?

Q34 Match each sentence to the report part it belongs to (findings, limitations, method): (i) “Fifty-five of the 250 students replied.” (ii) “The sample was chosen at random from each class list.” (iii) “Students who ride to school may have been left out, as the survey ran in winter.”

Q35 SCENARIO: a report on 40 students surveyed at one gate ends “the whole town walks to school”. Rewrite the conclusion so it does not run further than the evidence.

Q36 Plan a full investigation of the question “How well do Year 9 students sleep on the night before a test?” Name: the population; the sample and method (with a size); the kind of data; one display; the statistics; and one limitation you would name in the report.

Answers

Q1 (a) Not statistical — one fixed answer. (b) Statistical — the daily values vary and the variation is the point. (c) Not statistical — one fixed value. **Q2** Example: (a) population: all students at the school; variable: a rating of the canteen food from 1 to 5. (b) “What is the median rating, from 1 to 5, that students at this school give the canteen food?” **Q3** (a) Categorical, nominal. (b) Categorical, ordinal. (c) Numerical, discrete. (d) Numerical, continuous. **Q4** (a) Primary, numerical continuous. (b) Secondary, numerical continuous. **Q5** (a) All 340 members. (b) The 80 surveyed; not a census — a census asks all 340. **Q6** (a) Stratified sampling. (b) Random — 15 chosen at random from each class; random methods support claims about the population. **Q7** (a) Mean $72.8 \div 31 \approx 2.3$ mm; median 0.4 mm; range 14.2 mm. (b) $4 \div 31 \approx 12.9\%$. **Q8** (a) $14.6 - 11.6 = 3.0$ °C. (b) Melbourne, by $8.5 - 5.1 = 3.4$ °C of range. **Q9** (a) The mean. (b) The median — the wet days pull the mean far above a typical day. **Q10** (a) Small: $70 - 40 = 30$ points; large: $58 - 54 = 4$ points. (b) Bigger samples vary less, so they give a tighter estimate. **Q11** (a) Example: no names asked; taking part optional. (b) Privacy is protected; no one is pressed into sharing personal data. **Q12** (a) Bar graph — counts of labels need separate bars. (b) Dot plot or stem-and-leaf plot — keeps every one of the 31 values and shows the shape. **Q13** Example: “What percentage of Year 9 students at this school travel by bus on most days?” and “How do walking times to school vary across Year 9, and what is the median?” **Q14** (a) Nominal. (b) Ordinal. (c) Ordinal. (d) Nominal. **Q15** (a) Discrete. (b) Continuous. (c) Discrete. (d) Continuous. **Q16** (a) Secondary — the ABS collected it. (b) Primary — you measure it. (c) Secondary — the Bureau collected it. (d) Primary — you run it. **Q17** Convenience sampling: friends share tastes and opinions with the group, so the result leans toward the group’s view. **Q18** Voluntary response (people with strong feelings vote), and the visitors to one sports website are not a representative sample of Australians. **Q19** Mean $44.4 \div 31 \approx 1.4$ mm; median 0.0 mm; range 14.6 mm; $16 \div 31 \approx 51.6\%$ dry. **Q20** Difference $72.8 - 44.4 = 28.4$ mm; ratio $72.8 \div 44.4 \approx 1.6$; Ballarat had the wetter July. **Q21** Mean 4.75, median 4.5, range 6. **Q22** With 40: mean 9.4, median 6, range 36. Without: mean 6, median 6, range 4. The median did not change. **Q23** $18 \div 50 = 36\%$. **Q24** Rating 4; $17 \div 50 = 34\%$. **Q25** Voluntary response to an angry post (strong feelings only), and 20 replies is far too small. **Q26** Fair samples of the same size differ from run to run — sampling variation; a 4-point gap at that size is ordinary variation, not a contradiction. **Q27** $7.7 - 3.7 \approx 4.1$ °C; the ranges overlap from 3.1 °C to 8.4 °C, so a Melbourne morning inside the overlap can equal a Ballarat one. **Q28** It does the work of two dot plots drawn on one axis; a bar graph of two means would hide the overlap and the spread. **Q29** Source: Bureau of Meteorology, *Daily Weather Observations, Ballarat Aerodrome, station 089002* (product IDCJDW3005.202607), July 2026, prepared 14 August 2026. Retrieved 20 August 2026. **Q30** $55 \div 250 = 22\%$; non-response — if the non-responders differ (busier, later sleepers), the result leans. **Q31** Drop the name and address — they are not needed, and asking them may stop honest answers or expose individuals. **Q32** (a) 5 days. (b) 2 days. (c) Melbourne, by $8.5 - 5.1 = 3.4$ °C. **Q33** Strong for July 2026 itself — all 31 days, a census of the month; it cannot support “every July” or “always”, other months and years may differ. **Q34** (i) Findings. (ii) Method. (iii) Limitations. **Q35** Example: “Of the 40 students passing one gate, most walked; a sample from one gate at one time may not represent the town, so this is a first estimate only.” **Q36** Example: population — all Year 9 students; sample — 60, stratified by class, random within class; data — numerical discrete (hours of sleep) plus categorical (slept well: yes/no); display — dot plot; statistics — mean, median, range; limitation — students may round or report the sleep they think is expected.

Fully-worked solutions

Q1 A question is statistical when its answers vary and the variation is studied. (a) and (c) each have one fixed value, so they are not statistical. (b) has 31 varying answers per city, and the pattern of variation is exactly what is studied. ∴ Only (b) is statistical.

Q2 (a) The population must be a named group — all students at the school — and the variable must be recordable — a rating of the canteen food from 1 to 5 (categorical, ordinal). (b) The refined question names population, variable and measure: “What is the median rating, from 1 to 5, that students at this school give the canteen food?” ∴ The idea becomes statistical once population, variable and measure are named.

Q3 (a) Favourite sport is a label with no order — categorical, nominal. (b) A rating 1–5 is a label with an order — categorical, ordinal. (c) Siblings are counted — numerical, discrete. (d) Rainfall is measured — numerical, continuous. ∴ Nominal, ordinal, discrete, continuous respectively.

Q4 (a) You collect the measurements yourself, so the data are primary; reaction time is measured, so numerical continuous. (b) The Bureau collected the rainfall and you reuse it — secondary; also numerical continuous. ∴ (a) Primary; (b) secondary; both numerical continuous.

Q5 (a) The population is the whole group the question is about: all 340 members. (b) The sample is the 80 who are surveyed. A census asks every member, so 80 of 340 is a sample, not a census. ∴ Population 340; sample 80; not a census.

Q6 (a) Splitting the population into the four classes and sampling randomly within each is stratified sampling. (b) It is random — every student in a class has a known chance of being chosen — and only random methods support a claim about the whole population, so the 60 can stand for the 240. ∴ Stratified, random, and fit to support a claim about all 240.

Q7 (a) Mean = total ÷ days = $72.8 \div 31 \approx 2.3$ mm. The dot plot shows the middle value 0.4 mm and the spread $14.2 - 0 = 14.2$ mm. (b) Dry days 4 of 31: $4 \div 31 \approx 0.129 = 12.9\%$. ∴ Mean ≈ 2.3 mm, median 0.4 mm, range 14.2 mm; about 12.9% of days were dry.

Q8 (a) $14.6 - 11.6 = 3.0$ °C — Melbourne's middle day was 3.0 °C warmer. (b) Spread is measured by the range: Melbourne 8.5 °C against Ballarat 5.1 °C, a difference of 3.4 °C. ∴ Median difference 3.0 °C; Melbourne varied more, by 3.4 °C of range.

Q9 (a) The mean: the few very wet days add large totals, pulling the mean to 2.3 mm. (b) The median, 0.4 mm: half the days were at or below it, so it describes a typical day; the mean describes the total shared out, not a typical day. ∴ The mean is pulled; report the median for a typical day.

Q10 (a) Small samples: $70 - 40 = 30$ percentage points. Large samples: $58 - 54 = 4$ points. (b) The large samples cluster close together while the small ones swing — bigger samples vary less, giving a tighter estimate of the population value. ∴ Ranges 30 vs 4 points; sample size narrows variation.

Q11 (a) Two duties: ask no names, and make taking part optional. (b) Sleep is personal; without names the data cannot expose anyone, and optional participation means no one is pressed into sharing. (A third duty — reporting the group, never individuals — would also earn the mark.) ∴ No names and optional participation, each protecting the students asked.

Q12 (a) A bar graph: the data are labels with counts, and separate bars show each count honestly (bars start at zero). (b) A dot plot or stem-and-leaf plot: with 31 numerical values, either keeps every value and shows the shape — a single number or a pie would hide both. ∴ Bar graph for (a); dot plot or stem-and-leaf for (b).

Q13 Each question must name a population and a variable and expect variation. Example 1: “What percentage of Year 9 students at this school travel by bus on most days?” (population: Year 9 here; variable: usual mode). Example 2: “How do walking times to school vary across Year 9, and what is the median?” (variable: walking time in minutes, numerical continuous). ∴ Any two questions that name population and variable and expect variation.

Q14 (a) Car colour — labels, no order — nominal. (b) Medals have an order — ordinal. (c) Ratings 1–5 have an order — ordinal. (d) Suburbs are labels with no order — nominal. ∴ Nominal, ordinal, ordinal, nominal.

Q15 (a) Goals are counted — discrete. (b) Height is measured — continuous. (c) Wet days are counted — discrete. (d) Time is measured — continuous. ∴ Discrete, continuous, discrete, continuous.

Q16 (a) Secondary: the ABS collected and published the table; you reuse it. (b) Primary: you take the measurements yourself. (c) Secondary: the Bureau's observations were collected by the Bureau. (d) Primary: your survey, your data. ∴ (a) and (c) secondary; (b) and (d) primary.

Q17 Asking only friends is convenience sampling — a non-random method. Friends tend to share opinions, and the group's view is over-represented, so the result leans toward “change” or “keep” whichever way the group leans — every time the method is run. ∴ Convenience sampling; the lean is built into who is reached.

Q18 First, the poll is voluntary response: people with strong feelings vote, and the rest stay silent. Second, the people on one sports website are not representative of Australians — sports fans, already online, already on that site. 12 000

votes do not fix either lean: a big biased sample is still biased. $\therefore$ Voluntary response plus an unrepresentative pool make it weak evidence.

Q19 Mean = $44.4 \div 31 \approx 1.4$ mm. The dot-plot shape (most days at zero) gives median 0.0 mm and range $14.6 - 0 = 14.6$ mm. Dry days: $16 \div 31 \approx 0.516 = 51.6\%$. $\therefore$ Mean ≈ 1.4 mm, median 0.0 mm, range 14.6 mm; about 51.6% of days dry.

Q20 Difference: $72.8 - 44.4 = 28.4$ mm. Ratio: $72.8 \div 44.4 \approx 1.6$. One sentence: Ballarat collected about 1.6 times Melbourne's July 2026 rainfall. $\therefore$ 28.4 mm more at Ballarat; ratio ≈ 1.6 ; Ballarat wetter.

Q21 Sum = $2 + 3 + 3 + 4 + 4 + 4 + 5 + 5 + 6 + 6 + 7 + 8 = 57$; mean = $57 \div 12 = 4.75$. Ordered, the middle pair is 4 and 5, so median = 4.5. Range = $8 - 2 = 6$. $\therefore$ Mean 4.75, median 4.5, range 6.

Q22 With the 40: sum 94, mean 9.4; median (middle pair 6, 6) = 6; range $40 - 4 = 36$. Without: sum 54, $n = 9$, mean 6; median 6; range $8 - 4 = 4$. The median was 6 both times and the mean fell from 9.4 to 6 — the outlier moved the mean and the range, not the median. $\therefore$ Mean $9.4 \rightarrow 6$, median $6 \rightarrow 6$, range $36 \rightarrow 4$; the median did not move at all.

Q23 AFL share = $18 \div 50 = 0.36 = 36\%$. $\therefore$ 36% chose AFL.

Q24 The largest count is 17, for rating 4 — the most common (modal) rating. Share = $17 \div 50 = 34\%$. $\therefore$ Rating 4, given by 34%.

Q25 First, replies to an angry post are voluntary response — anger replies, mild views do not. Second, 20 replies is far too small to carry an “80% of students” claim. The headline also runs from “20 replies” to “students” generally — further than the evidence allows. $\therefore$ Voluntary response and a tiny sample.

Q26 Repeated fair samples of the same size give different results — that is sampling variation, not error. At that sample size a 4-point gap sits well inside ordinary variation (compare Q10: samples of 200 still differed by up to 4 points). $\therefore$ Not a contradiction — expected sampling variation.

Q27 Mean difference: $7.7 - 3.7 \approx 4.1$ °C — Melbourne's nights ran warmer on average. The range bars overlap from 3.1 °C to 8.4 °C: a Ballarat morning of 6 °C and a Melbourne morning of 6 °C are both ordinary. $\therefore$ Means differ by ≈ 4.1 °C, yet overlapping ranges mean individual mornings can match.

Q28 The range-bar figure puts both full data ranges and both medians on one axis — the work of two dot plots drawn together. A bar graph of the two means alone would show two bars 4.1 °C apart and hide the overlap entirely, letting a reader conclude every Melbourne morning beats every Ballarat one. $\therefore$ The range bars show spread and overlap; two mean-bars would hide both.

Q29 R1 wants publisher, title, year and retrieval date; the product details make it repeatable: Source: Bureau of Meteorology, *Daily Weather Observations, Ballarat Aerodrome, station 089002 (product IDCJDW3005.202607)*, July 2026, prepared 14 August 2026. Retrieved 20 August 2026. $\therefore$ A reader can fetch the same data from that line alone.

Q30 Response rate = $55 \div 250 = 0.22 = 22\%$. The 78% who did not reply produce non-response: if they differ from responders (later sleepers may ignore a morning email), the result leans toward the responders' habits. $\therefore$ 22% response; non-response risk of a leaned result.

Q31 The name and home address are not needed to analyse the answers, and collecting them exposes individuals and may stop honest answers. Drop both before the survey runs; report results as a group only. $\therefore$ Remove the name and address; keep the survey to the variables the question needs.

Q32 (a) Ballarat leaves on stem 13: 1, 2, 3, 6, 6 — 5 days. (b) Melbourne: stem 18 gives 1 day (18.4) and stem 19 gives 1 day (19.5) — 2 days. (c) Ranges: Melbourne 8.5 °C, Ballarat 5.1 °C — Melbourne varied more, by 3.4 °C. $\therefore$ 5 days; 2 days; Melbourne, by 3.4 °C of range.

Q33 The data cover all 31 days of July 2026 — a census of that month, so there is no sampling variation at all for the month itself: the 3.7 °C mean difference is a fact about July 2026, not an estimate. What it cannot support is “every July” or “always” — other months and other years vary, and one month is one run of the weather. $\therefore$ Strong for July 2026; no support beyond that month.

Q34 (i) reports a result in context — findings. (ii) says how the sample was chosen — method. (iii) names who the method could not capture — limitations. ∴ (i) Findings, (ii) method, (iii) limitations.

Q35 The sample — 40 students, one gate, one time — is a convenience sample, so the conclusion must stay inside it. Rewrite: “Of the 40 students passing one gate, most walked. Because one gate at one time may not represent the town, this is a first estimate; a stratified sample across gates and times would test it.” ∴ The claim shrinks to the sample, and names the next fairer step.

Q36 A defensible plan: population — all Year 9 students at the school; sample — 60 students, stratified by class with random choice within each class; data — numerical discrete (hours of sleep) plus a categorical yes/no (slept well?); display — dot plot of hours; statistics — mean, median and range; limitation named in the report — students may round their hours or report what they think is expected. ∴ A complete plan names population, fair sample, data kind, display, statistics and one limitation.

Extended task X1 — plan, conduct and report your own statistical investigation

This task is the main assessment of 6E (VC2M9ST05). You will plan and conduct a statistical investigation of your own, report the findings, and discuss the strength of your evidence. Time: two to three weeks including class time. The investigation may use primary data you collect, secondary data from a source you cite (the ABS and BOM online databases are the right default), or both.

The brief

1. Pose one statistical question. It must name the population and the variable, and it must expect variation.
2. Plan the collection with the planning frame below. Name the sampling method and its limitations before you collect.
3. Collect the data — at least 30 values if numerical. Keep the duties: optional participation, no names, results as a group.
4. Organise and analyse: at least one honest display, and the mean, median and range where they suit the data.
5. Report in the seven parts shown in the theory, with a source line for any secondary data.
6. Discuss the strength of your evidence with the four checks, and write a conclusion that does not run further than the evidence allows.

Planning frame

Part	Write
Question	the statistical question, worded exactly
Population	the whole group the question is about, with a size if knowable
Sample and method	the size, the method by name, and how the choice is made fair
What is recorded	each variable, and its kind (categorical / numerical; nominal / ordinal; discrete / continuous)
Display and statistics	the display(s) and the mean, median, range where they suit
Bias to watch	who the method cannot reach, who may not answer, what may lean the answers
Sources	for secondary data: publisher, title, year, retrieval date

Data-quality checklist — run it before you report

- The method is named and fair, or its lean is stated.
- The sample reaches the population the question names.
- The size is big enough for the claim (30+ numerical values; a small range across repeat samples).

- Displays start at zero where they should, and show the data honestly.
- If the data lean, the median — not just the mean — is reported.
- Every secondary datum carries a source line.
- One limitation at least is named in the report.

Rubric — 20 marks (5 criteria × 4)

Criterion	The 4 marks
Formulate	a statistical question that names the population, the variable and the measure, and expects variation
Solve	a fair or honestly-named method, data collected as planned, and the analysis carried through with an honest display and correct mean, median and range
Interpret	findings stated in context with units, and the displays and statistics read correctly
Evaluate	the strength of the evidence discussed with the four checks, limitations named, and the conclusion kept inside the evidence
Report	all seven parts present, sources cited, and a reader able to repeat the investigation from the method

EXEMPLAR — practice, not the task. The report below is worked on the Bureau of Meteorology’s July 2026 data so you can see the seven parts in action. Your task is your own question and your own data — this exemplar is practice, not the task.

Title and question. How did the daily maximum temperatures of July 2026 compare between Ballarat and Melbourne?

Method and data source. Secondary data: all 31 days of July 2026 at each station — a census of the month. Source: Bureau of Meteorology, *Daily Weather Observations, Ballarat Aerodrome, station 089002 (product IDCJDW3005.202607), July 2026*, prepared 14 August 2026. Retrieved 20 August 2026; and the matching Melbourne Olympic Park product (station 086338, IDCJDW3033.202607).

Analysis. Back-to-back stem-and-leaf plot (theory figure). Ballarat: mean 11.2 °C, median 11.6 °C, range 5.1 °C. Melbourne: mean 14.9 °C, median 14.6 °C, range 8.5 °C.

Findings. Melbourne’s July 2026 days ran about 3–4 °C warmer than Ballarat’s (mean difference ≈ 3.7 °C), and varied more (range 8.5 °C against 5.1 °C).

Limitations. One month, two stations; temperatures say nothing about wind or rain, and the stations are aerodrome and park sites, not street level.

Strength of the evidence. Check 1: no sampling at all — every day of the month is included, so the method is fair by design. Check 2: 31 days each is the whole population of interest, so size is not a limit. Check 3: mean, median and the stem-and-leaf plot all point the same way. Check 4: no display trick — both sets on one axis, same leaf unit. The evidence is strong for July 2026.

Conclusion. In July 2026, Melbourne was the warmer and more variable of the two. Because one month is one run of the weather, the result says nothing certain about other Julys — repeating the comparison in another month or year is the natural next loop of the cycle.

Test

Test T6 · Chapter 6 — Statistics · Level 9

This is a school-designed paper. It is not a VCAA instrument. Achievement in the Victorian Curriculum is reported against the Level 9 achievement standard, not against a mark on this paper.

Total marks	36 (Part A: 10 marks · Part B: 26 marks)
Timing	5 minutes reading time, then 36 minutes writing time (1 minute per mark). Total time: 41 minutes.
Equipment — Part A	Technology-free: no calculator and no digital tool.
Equipment — Part B	Technology-active: a scientific calculator and a spreadsheet, or another tool nominated by your teacher (a graphing tool, dynamic geometry software or a general-purpose programming language). No CAS (computer algebra system).
Cumulative	This test is cumulative over Chapters 1–5: earlier number skills (decimals, fractions, percentages) are used inside the Chapter 6 questions.
Reasoning	Reasoning is marked separately from the answer in this paper. Marks for ‘explain’ and ‘justify’ parts are awarded for the stated reason; a correct number with no reason does not score full marks.

Marks by subtopic

Subtopic	Content description	Marks	Achievement-standard anchor (§4)
6A — Comparing distributions	VC2M9ST03	10	AS 16
6B — Choosing and justifying a data display	VC2M9ST04	6	AS 16
6C — Reading survey reports	VC2M9ST01	6	AS 17
6D — Sampling methods and representation	VC2M9ST02	10	AS 17
6E — Planning and conducting a statistical investigation	VC2M9ST05	4	AS 17
Total		36	

Attempt all questions. Show clear working wherever a part carries more than 1 mark.

Part A — Technology-free (10 marks)

No calculator and no digital tool in this part.

Q1. SCENARIO. The ages, in years, of the eleven volunteers restoring a local creek reserve are shown in the stem-and-leaf plot.

Stem	Leaf
2	3 5 8

Stem	Leaf
3	1 4 6 8
4	1 4 7
5	2

Key: 3 / 4 means 34 years.

- Find the median age. **[1 mark]**
- Find the range of the ages. **[1 mark]**
- A 74-year-old joins the group. Without recalculating every measure, state which of the mean, the median and the range changes the most. **[1 mark]**

Q2. SCENARIO. Thirty Year 9 students each record their travel time to school, in minutes. The times are numerical (continuous) data. The students want one display that shows the shape of the times — the centre, the spread, and any outlier.

- Which display is the most appropriate choice? **[1 mark]**

A. a pie chart B. a histogram C. a line graph D. a stacked bar chart

- Justify your choice in part a: name what the display shows that the question asks for. **[1 mark]**

Q3. SCENARIO. A school newsletter reports:

We surveyed 120 Year 9 students at our school. Each student named their favourite canteen lunch and recorded the number of days per week they buy lunch at the canteen. The mean number of days was 2.4.

- Name the numerical variable and the categorical variable in this survey. **[1 mark]**
- State the population this survey is trying to say something about, and explain why the figure 2.4 is an estimate rather than an exact value. **[1 mark]**

Q4. SCENARIO. A student council wants to estimate how long students take to travel to school. They decide to survey 20 students from each year level. The surveyor stands at the school gate and chooses students until 20 from each year level have answered.

- Name the sampling method. **[1 mark]**
- Explain one way this method can lean the result. **[1 mark]**

Q5. SCENARIO. A student asks, “What is the typical travel time to school for Year 9 students at our school?” They survey only their own friendship group and conclude that the median travel time for all Year 9 students is 12 minutes. Explain one reason why the evidence is too weak to support a conclusion about all Year 9 students. **[1 mark]**

Part B — Technology-active (26 marks)

You may use a scientific calculator and a spreadsheet, or the tool your teacher nominates. Show your method — marks are awarded for method, not just for the final number. No CAS.

Q6. The daily rainfall, in millimetres, recorded through July 2026 at two Victorian stations. The data is real; the values are in date order (1 July to 31 July).

Melbourne (Olympic Park): 6.6, 4.0, 7.0, 0.2, 1.6, 0.0, 0.0, 0.0, 0.2, 0.0, 0.0, 14.6, 1.4, 0.4, 0.2, 0.2, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 3.0, 0.0, 0.0, 0.6, 0.0, 0.0, 2.2, 2.2, 0.0

Ballarat: 9.2, 10.0, 6.0, 1.6, 1.2, 0.2, 0.2, 0.2, 0.4, 0.2, 0.2, 14.2, 3.4, 5.6, 1.8, 0.2, 0.0, 0.4, 0.0, 0.2, 0.2, 0.2, 0.8, 1.2, 0.8, 4.8, 0.0, 0.0, 8.4, 1.0, 0.2

Source: Australian Bureau of Meteorology, “Daily Weather Observations for Melbourne (Olympic Park), Victoria, July 2026” (station 086338) and “Daily Weather Observations for Ballarat, Victoria, July 2026” (station 089002), prepared 14 August 2026. Retrieved 20 August 2026.

- Enter the Melbourne data into a spreadsheet or a calculator list. Find the mean and the median daily rainfall. **[2 marks]**
- Repeat for the Ballarat data: find the mean and the median daily rainfall. **[2 marks]**
- Compare the two distributions: write one comparison sentence about centre and one about spread, in context. **[2 marks]**
- At both stations the mean is well above the median. Name the shape of the distributions and explain what creates it. **[1 mark]**

Q7. SCENARIO. A student records the temperature of a classroom every hour from 9:00 am to 3:00 pm on one day, and wants to show how the temperature changed through the day.

- Name the most appropriate display. **[1 mark]**
- Justify your choice: name the feature of the data that makes this display the right one. **[1 mark]**

Q8. SCENARIO. Thirty students each name their favourite sport. Another student draws the results as a line graph, joining the five category points with line segments.

- Explain what is wrong with using a line graph for this data. **[1 mark]**
- Name a display that is appropriate. **[1 mark]**

Q9. SCENARIO. A sporting organisation publishes a report:

We surveyed 40 Year 9 students from 12 Victorian schools. Each student recorded the hours of organised sport they do each week and their main sport. The mean was 4.5 hours per week; 30% named football as their main sport.

- Calculate how many of the 40 students named football as their main sport. **[1 mark]**
- Use the sample result to estimate the mean hours of organised sport per week for all Victorian Year 9 students, and name the population value you are estimating. **[1 mark]**
- Suppose all 40 students were recruited through football clubs. Explain what effect this has on the estimate in part b. **[1 mark]**
- The report also states that the median was 3.5 hours. With a mean of 4.5 and a median of 3.5, what does this suggest about the shape of the distribution? **[1 mark]**

Q10. SCENARIO. The 20 students in a Year 9 form group each record the number of brothers and sisters they have:

0, 0, 0, 1, 1, 1, 1, 2, 2, 2, 2, 2, 2, 2, 3, 3, 3, 4, 4, 5

Two students each take a simple random sample of 5 values from this list.

Sample 1: 1, 2, 2, 3, 4 Sample 2: 0, 0, 1, 3, 3

- Calculate the mean of Sample 1 and the mean of Sample 2. **[1 mark]**
- The mean of the full list of 20 is 2.0. Explain why the two sample means differ from 2.0 and from each other, and state one way to reduce the effect. **[1 mark]**

Q11. SCENARIO. A shopping centre's management wants to know whether shoppers support extending Thursday night trading. They set up a stall near the food court between 6:00 pm and 8:00 pm on a Thursday and survey the people who walk past.

- Name the sampling method. **[1 mark]**
- Explain one way this method can lean the result. **[1 mark]**

Q12. SCENARIO. A company's brochure shows its profit for two years as a bar graph: \$1.00 million in 2023 and \$1.20 million in 2024. The vertical axis starts at \$0.95 million, so the 2024 bar looks many times the height of the 2023 bar.

- Explain how the choice of axis promotes the point of view that profits "grew strongly". **[1 mark]**
- State one change that would make the graph a fairer representation. **[1 mark]**

Q13. SCENARIO. A real-estate agent's flyer reads: "Average house price in our suburb: \$1,150,000." The figure is the mean of the suburb's last 20 sales: 18 of the sales were between \$600,000 and \$750,000 (averaging \$650,000), one was \$5,100,000 and one was \$6,200,000.

- a. Explain why quoting the mean of \$1,150,000 promotes the point of view that the suburb is expensive, and name the shape of the price distribution. **[1 mark]**
- b. Give a fairer single-number summary of a typical sale and justify your choice. **[1 mark]**

Q14. SCENARIO. A student investigates: "Do students at our school drink more water on sports days?" They survey 30 students from one Year 9 class, asking each to remember how many glasses of water they drank on the last sports day, and conclude: "Students drink more water on sports days."

- a. State one weakness in how the data was collected. **[1 mark]**
 - b. Explain why the conclusion does not follow from this evidence. **[1 mark]**
 - c. Suggest one change that would strengthen the investigation. **[1 mark]**
-

Solutions

For the teacher. The answer key below gives final values only; the fully-worked solutions that follow show the complete method for every item.

Answer key

Q1 a 36 years; b 29 years; c the range. **Q2** a B (a histogram); b the data is numerical (continuous) and the histogram shows its shape, centre and spread and any outliers. **Q3** a numerical: the number of days per week the student buys lunch at the canteen; categorical: the favourite canteen lunch; b the population is all Year 9 students at the school — only the sample of 120 answered, so 2.4 is an estimate of the population mean. **Q4** a quota sampling; b the surveyor chooses who fills each quota place, so the sample is not random and can lean the result. **Q5** The sample is a convenience sample of friends, who may live near each other and travel alike — it is not representative of all Year 9 students. **Q6** a mean ≈ 1.4 mm, median 0.0 mm; b mean ≈ 2.3 mm, median 0.4 mm; c Ballarat was typically wetter (higher mean and median) and the two months were similarly spread (ranges 14.2 mm against 14.6 mm); d positive skew — many near-dry days and a few very wet days pull each mean above the median. **Q7** a a line graph; b the data is numerical and recorded in order over time; a line graph shows the trend. **Q8** a line segments invent an order and values between categories, but favourite sport is nominal data; b a bar chart. **Q9** a 12 students; b about 4.5 hours per week — an estimate of the population mean; c the estimate is biased high, because club players typically do more sport; d positive skew — a few high-hour students pull the mean above the median. **Q10** a Sample 1: 2.4; Sample 2: 1.4; b each random sample holds different members, so sample means vary around the population mean; larger samples reduce the effect. **Q11** a convenience sampling; b people shopping on a Thursday night already shop late, so they are more likely to support the extension — others cannot be chosen. **Q12** a the axis starting at \$0.95 million makes a 20% rise look about five times as tall; b start the axis at zero (or state the growth as 20%). **Q13** a the two extreme sales pull the mean far above a typical sale — the distribution is positively skewed; b the median, which lies between \$600,000 and \$750,000, describes a typical sale better. **Q14** a e.g. the glasses are remembered, not measured; b the sample (one Year 9 class) cannot speak for the whole school, and no non-sports-day data was collected to compare against; c e.g. sample students from every year level, measure the water actually drunk, or collect data on both sports days and non-sports days.

Fully-worked solutions

Q1.

Stem	Leaf
2	3 5 8
3	1 4 6 8
4	1 4 7
5	2

Key: 3 / 4 means 34 years. The plot holds eleven values: 23, 25, 28, 31, 34, 36, 38, 41, 44, 47, 52.

a Eleven values: the median is the 6th value in order — 36 years. b Range = $52 - 23 = 29$ years. c With the 74-year-old included: the range becomes $74 - 23 = 51$ years (a change of 22), the median moves from 36 to $\frac{36+38}{2} = 37$ (a change of 1), and the mean moves from $\frac{399}{11} \approx 36.3$ to $\frac{473}{12} \approx 39.4$ (a change of about 3.1). $\therefore$ a median 36 years; b range 29 years; c the range changes the most — one unusual value stretches the range completely and barely moves the median.

Q2.

a B — a histogram. b The data is numerical (continuous) with 30 values, and the question asks for shape, centre, spread and outliers. A histogram groups continuous data into classes and shows exactly those features; a pie chart

needs categories that are parts of a whole, a line graph needs values recorded in order over time, and a stacked bar chart needs totals made of parts.

Option	Verdict
A pie chart	needs categories as shares of a whole — not this data
B histogram	shows shape, centre, spread and outliers of numerical data ✓
C line graph	needs values recorded in order over time
D stacked bar chart	needs totals made of parts across categories

∴ a B; b the histogram is the display that shows the shape, centre and spread of numerical data, and any outliers.

Q3.

a The numerical variable is the number of days per week the student buys lunch at the canteen (its values are numbers that can be averaged). The categorical variable is the favourite canteen lunch (its values are labels). b The population is all Year 9 students at the school. Only the 120 students in the sample answered, so 2.4 is an estimate of the population mean — a different sample of 120 students would give a slightly different mean. ∴ a numerical: days per week buying lunch; categorical: favourite canteen lunch; b population: all Year 9 students at the school; 2.4 is an estimate because it comes from a sample, not from every student.

Q4.

a Quota sampling — a set number of places (20) is fixed for each year level, and the person collecting the data chooses who fills them. b The surveyor's choice is not random: students who look likely to answer quickly are picked, and the others have no chance of being chosen. If, for example, the surveyor favours students arriving by car, the recorded travel times lean shorter than the school's true spread of times. ∴ a quota sampling; b the collector chooses who fills each place, so the sample can lean the result in a direction the method does not control.

Q5.

The friendship group is a convenience sample. Friends often live near each other and travel to school the same way, so their times can all sit in one narrow band, and no student outside the group had any chance of being chosen. ∴ The sample is not representative of all Year 9 students, so its median cannot be reported as the median for the whole year level.

Q6.

The two distributions as frequency tables (a value of 0.0 is a day with no recorded rain):

Class (mm)	Melbourne frequency	Ballarat frequency
0.0	16	4
0.1–2.0	8	19
2.1–4.0	4	1
4.1–6.0	0	3
6.1–8.0	2	0
8.1–10.0	0	3
10.1–12.0	0	0
12.1–14.0	0	0
14.1–16.0	1	1

a Melbourne: the 31 values sum to 44.4, so the mean is $\frac{44.4}{31} \approx 1.4$ mm. Sixteen of the 31 days recorded 0.0 mm, so the 16th of the 31 ordered values is 0.0 — the median is 0.0 mm. b Ballarat: the 31 values sum to 72.8, so the mean is

$\frac{72.8}{31} \approx 2.3$ mm. In order, the values run 0.0 (4 days), then 0.2 (10 days), then 0.4, ... — the 16th value is 0.4 mm, so the median is 0.4 mm. c Centre: Ballarat was typically wetter — its median (0.4 mm against 0.0 mm) and its mean (2.3 mm against 1.4 mm) both sit higher, and its month totalled 72.8 mm against Melbourne's 44.4 mm. Spread: the months were similarly variable — the ranges are 14.2 mm and 14.6 mm. d Both distributions are positively skewed. Most days recorded little or no rain (16 of Melbourne's 31 days recorded 0.0 mm), while a few very wet days — the wettest was 14.6 mm at Melbourne and 14.2 mm at Ballarat, both on 12 July — pull each mean well above the median and stretch each range. $\therefore$ a Melbourne: mean ≈ 1.4 mm, median 0.0 mm; b Ballarat: mean ≈ 2.3 mm, median 0.4 mm; c Ballarat typically wetter on both centre measures, with similar spread; d positive skew at both stations — a few large values pull the mean above the median.

Q7.

a A line graph. b The temperatures are numerical data recorded in order over time (hourly through the day). A line graph joins the points in order and shows the trend — when the temperature rises, when it falls, and its highest and lowest moments. $\therefore$ a line graph; b numerical data recorded in order over time is the data type a line graph is built for.

Q8.

a Favourite sport is nominal data — labels with no natural order. The line segments of a line graph suggest an order and values between the categories, inventing a trend that is not there: a line graph is only honest when the order means something, almost always time. b A bar chart — one bar per sport, with the bars free to stand in any order (sorting them by height makes the comparison easiest to read).

Feature	Line graph	Bar chart
Suggests an order between categories	yes — wrongly	no
Shows how the category counts compare	weakly	clearly

$\therefore$ a the line segments invent an order and a trend for data that has neither; b a bar chart.

Q9.

a $\frac{30}{100} \times 40 = 12$ students. b The estimate is about 4.5 hours per week. It is an estimate of the population mean — the mean hours of organised sport per week for all Victorian Year 9 students. c Students recruited through football clubs already play organised football, so they typically do more sport than students who do not join clubs. The sample leans towards high-hour students, and the estimate of 4.5 hours is biased high. d The mean (4.5) sits above the median (3.5), so a tail of high values pulls the mean up: the distribution is positively skewed, with a few students doing many hours of sport each week. $\therefore$ a 12 students; b ≈ 4.5 hours per week, an estimate of the population mean; c the estimate is biased high; d positive skew.

Q10.

The full list as a frequency table:

Brothers and sisters	0	1	2	3	4	5
Frequency	3	4	7	3	2	1

a Sample 1: $\frac{1+2+2+3+4}{5} = \frac{12}{5} = 2.4$. Sample 2: $\frac{0+0+1+3+3}{5} = \frac{7}{5} = 1.4$. b The population mean is $\frac{40}{20} = 2.0$.

Each simple random sample happens to hold different members of the form group, so each sample mean lands in a different place around 2.0 — this is sample-to-sample variation, and it is expected, not a mistake. Drawing larger

samples reduces the effect: larger samples sit closer to the population mean on average. $\therefore$ a Sample 1 mean 2.4, Sample 2 mean 1.4; b random samples vary from sample to sample; use larger samples to reduce the effect.

Q11.

a Convenience sampling — the management surveys whoever is easiest to reach at that stall. b Anyone at the food court on a Thursday night already shops late, so they are more likely to support extending Thursday trading; people who work at night, shop at other times or never visit the centre cannot be chosen at all. The result leans towards support. $\therefore$ a convenience sampling; b the sample over-represents people who already shop on Thursday nights, so support is overstated.

Q12.

Year	Profit (\$ million)	Height drawn above the axis start at \$0.95 million
2023	1.00	$1.00 - 0.95 = 0.05$
2024	1.20	$1.20 - 0.95 = 0.25$

a The drawn heights are 0.05 and 0.25, so the 2024 bar looks $\frac{0.25}{0.05} = 5$ times as tall as the 2023 bar — while the true growth is only $\frac{1.20 - 1.00}{1.00} = 20\%$. Cutting off the bottom of the axis exaggerates the difference and promotes the “grew strongly” point of view. b Start the vertical axis at zero — the 2024 bar is then only 1.2 times the height of the 2023 bar — or state the growth plainly as 20%. $\therefore$ a the truncated axis makes a 20% rise look five times as tall; b start the axis at zero, or report the growth as 20%.

Q13.

The mean checks: total of the 20 sales
 $= 18 \times \$650,000 + \$5,100,000 + \$6,200,000 = \$11,700,000 + \$11,300,000 = \$23,000,000$, and
 $\$23,000,000 \div 20 = \$1,150,000$.

Ordered sales	Positions	Values
the 18 ordinary sales	1st to 18th	between \$600,000 and \$750,000
the two large sales	19th and 20th	\$5,100,000 and \$6,200,000

a Eighteen of the twenty sales sit between \$600,000 and \$750,000, but the two extreme sales pull the mean up to \$1,150,000 — far above what a typical house sold for. Quoting the mean promotes the “expensive suburb” view, and the two high values give the price distribution a positive skew. b The median is a fairer single number: it is the mean of the 10th and 11th ordered sales, and both of those are among the 18 ordinary sales, so the median lies between \$600,000 and \$750,000 — close to what a typical house actually sold for. $\therefore$ a the mean is pulled up by two extreme sales — the distribution is positively skewed; b the median (between \$600,000 and \$750,000) describes a typical sale.

Q14.

a Any one weakness: the glasses are remembered and self-reported, not measured; only one Year 9 class was surveyed; only sports-day drinking was recorded. b The conclusion compares sports days with other days, but no non-sports-day data was collected, so there is nothing to compare against; and one Year 9 class cannot speak for the students in Years 7 to 12. The evidence does not reach as far as the conclusion. c Any one strengthening change: sample students from every year level, not one class; measure the water actually drunk (for example, count filled bottles) instead of asking for memories; collect data on both sports days and non-sports days so the two can be compared. $\therefore$ a e.g. self-reported recall; b no comparison data and an unrepresentative sample; c e.g. sample across all year levels and measure both day types.

Marking guide

This paper is a classroom instrument for the school's own reporting (D12). Marks map to the corrected achievement-standard anchors of §4:

AS 16 items — Q1, Q2, Q6, Q7, Q8 (16 marks): credit is given where the student compares and analyses the distributions of numerical data sets using mean, median and range, describes shape (positive skew) and the effect of unusual values, and chooses and justifies representations for the data type.

AS 17 items — Q3, Q4, Q5, Q9, Q10, Q11, Q12, Q13, Q14 (20 marks): credit is given where the student explains how the way data was obtained and the choice of representation can lean a result, support or question a conclusion, or promote a point of view, and interprets findings against the strength of the evidence.

Question	Subtopic	CD	Marks	Notes
Q1 a–c	6A	VC2M9ST03	3	median, range, effect of one unusual value
Q2 a–b	6B	VC2M9ST04	2	1 for the selection, 1 for the justification
Q3 a–b	6C	VC2M9ST01	2	variables; population and estimate
Q4 a–b	6D	VC2M9ST02	2	method named; 1 mark for the stated lean
Q5	6E	VC2M9ST05	1	strength of evidence
Q6 a–d	6A	VC2M9ST03	7	real Victorian data; the multi-part context question
Q7 a–b	6B	VC2M9ST04	2	display chosen and justified
Q8 a–b	6B	VC2M9ST04	2	1 for the stated reason
Q9 a–d	6C	VC2M9ST01	4	estimating a population mean; bias; shape
Q10 a–b	6D	VC2M9ST02	2	sample-to-sample variation
Q11 a–b	6D	VC2M9ST02	2	method named; 1 mark for the stated lean
Q12 a–b	6D	VC2M9ST02	2	representation promoting a point of view
Q13 a–b	6D	VC2M9ST02	2	choice of mean vs median; 1 mark for the justification
Q14 a–c	6E	VC2M9ST05	3	interpret the finding against the evidence
Total			36	Part A: 10 · Part B:

Question	Subtopic	CD	Marks	Notes
				26

Reasoning marks in this paper are the ‘explain’, ‘justify’ and ‘state one reason’ parts — each is awarded only for a stated reason; a correct number with no reason scores 0 for that part. Accept any mathematically correct reason, not only the one printed in the solutions.