

Science 7–8

Chapter 5 — Evaluating

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
Interrogating a method and a conclusion: assumptions, sources of error, conflicting evidence and the questions still unanswered	2
Building an evidence-based argument: source credibility, evaluating claims, and the ethics and cultural protocols of using secondary data	20

Interrogating a method and a conclusion: assumptions, sources of error, conflicting evidence and the questions still unanswered

What we teach here — VC2S8I06 (Victorian Curriculum F–10 Version 2.0, Science, Levels 7 and 8):

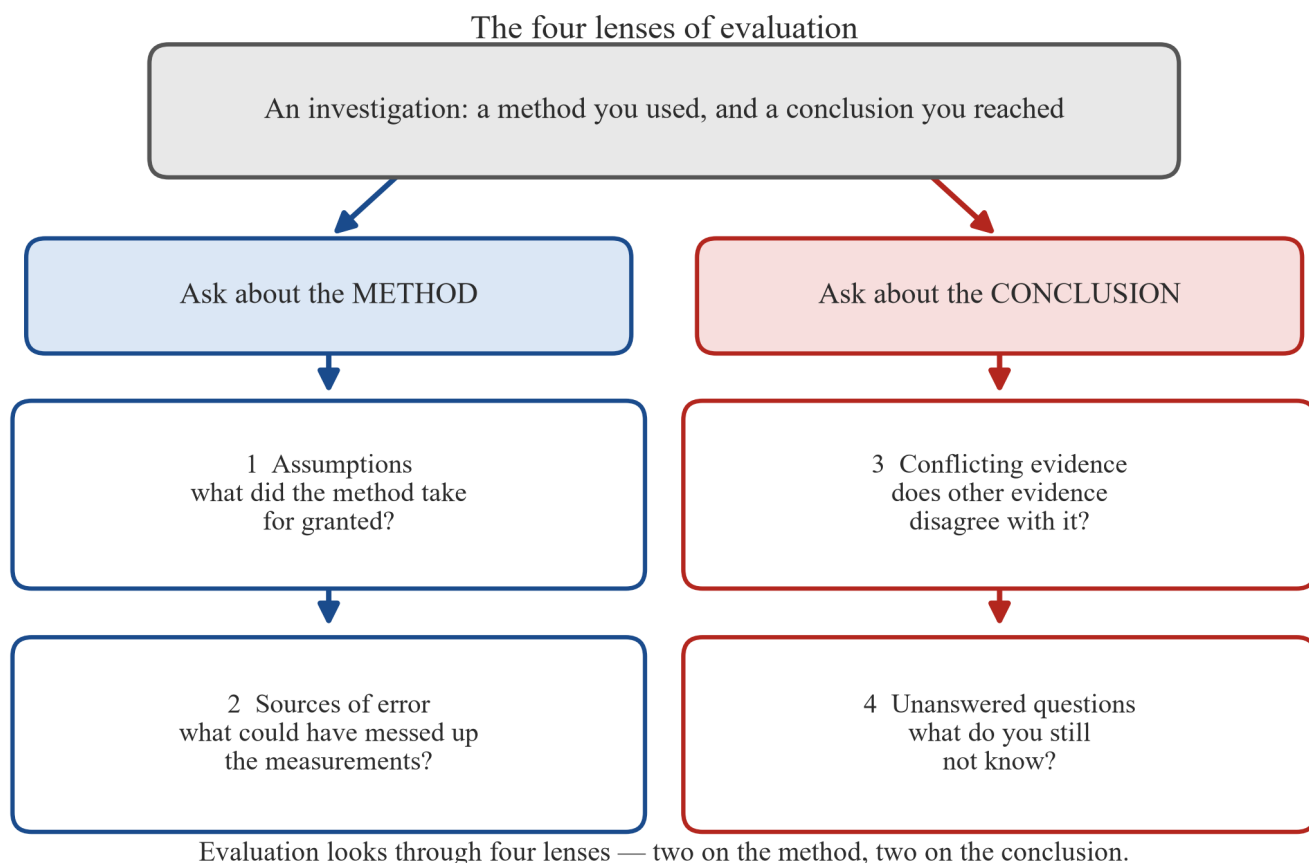
“scientific methods, conclusions and claims can be analysed to identify assumptions, possible sources of error, conflicting evidence and unanswered questions” **This subtopic owns the whole CD:** there is no partner subtopic (single delivery, D6) — 5A carries every clause of the CD, and every one of its seven elaborations is drawn on below as context. This is the evaluation subtopic of the book: where 2A taught you to ask a good question and 3A will teach you to plan a good test, this one teaches you to look hard at a method and a conclusion after the measuring is done. **In our words:** every investigation takes things for granted and lets a little luck into the numbers, and every conclusion is only as strong as the evidence standing behind it. To interrogate a method and a conclusion is to look through four lenses — *what did the method take for granted? what could have messed up the measurements? does other evidence disagree with the conclusion? what is still unanswered?* — and none of it is about calling anyone stupid. It is how science keeps score of what it actually knows. **Achievement-standard extract served here (D13):** *“They identify assumptions and sources of error in methods and analyse conclusions and claims with reference to conflicting evidence and unanswered questions.”* 5A owns VC2S8I06 outright, so the extract appears here in full. **Elaborations drawn on:** VC2S8I06.E01 (sources of error and suggestions for improvement), .E02 (evaluating a parachute-drop method), .E03 (naming assumptions that are so widely accepted nobody checks them, here around physical and chemical change), .E04 (a claim about what is in a mixture, weighed against conflicting evidence), .E05 (the bamboo-in-an-earthquake -area claim and the evidence it would take), .E06 (comparing your results with other groups and with secondary sources), .E07 (asking the further questions a conclusion leaves open). **Scope notes:** the ceiling is held (D9). *Why* a bigger canopy falls more slowly is the forces story of Ch15; the gas behind the fizzing powder is named and tested in 11C; earthquakes themselves are Ch12. Here those contexts are settings for evaluation, never things we explain. The maths we lean on is the D11 interlock: percentage error (VC2M8N05), the spread of repeated readings and what more trials do to it (VC2M8ST02, VC2M8ST03), and rounding (VC2M7N04). Planning a fair test from scratch is 3A’s job; graphing conventions are 4A/4C’s. Here we evaluate methods and conclusions that already exist.

Word watch — *assumption, error, random, systematic, claim, evidence*. An **assumption** is anything a method takes for granted without checking. Assumptions are not cheating — every method on Earth has them, because you cannot check everything at once. The habit science builds is naming them and asking how reasonable each one is. An **error**, in science, is not a slip or a mistake; it is any gap between a measurement and the true value, and it comes in two kinds. A **random error** is chance in the measurement — readings scatter on both sides of the true value, a little high here, a little low there. A **systematic error** is a built-in shift — every reading moves the same way, by roughly the same amount, because something in the setup or the instrument is off. (Those two words are the curriculum’s own from Year 7 and Year 4 science respectively; we define them anyway because they carry the whole subject.) A **claim** is a conclusion or statement someone asks you to believe; **evidence** is the observations and measurements standing behind it. Evaluation is what happens when claims and evidence are made to face each other.

Theory

The four lenses of evaluation

An investigation ends on paper: a method, a table of results, a conclusion. The temptation is to call it done. Science does not. Everything on that page now gets looked at through four lenses — two pointed at the method, two pointed at the conclusion.



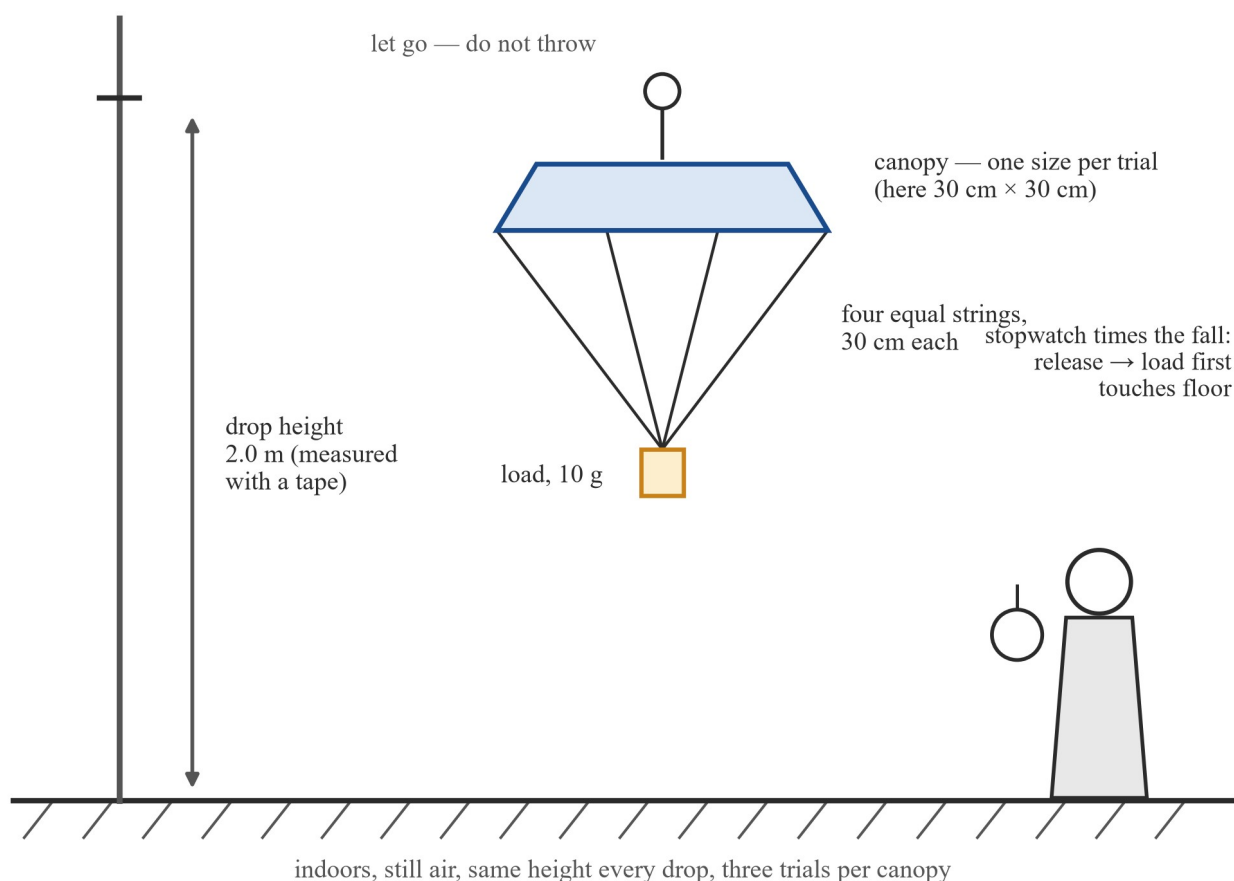
Flow chart titled “The four lenses of evaluation”: an investigation — a method you used and a conclusion you reached — is examined through two blue lenses on the method (1 Assumptions: what did the method take for granted? 2 Sources of error: what could have messed up the measurements?) and two red lenses on the conclusion (3 Conflicting evidence: does other evidence disagree with it? 4 Unanswered questions: what do you still not know?); caption: evaluation looks through four lenses — two on the method, two on the conclusion

The lenses are questions, not accusations. A good method still assumes things and still collects error; a good conclusion still leaves questions open. Writing an evaluation means saying *which* assumptions, *which* errors, *which* conflicting evidence, *which* open questions — and how much each one matters. One investigation will carry us through all four lenses: a parachute drop.

Lens 1 — assumptions: what did the method take for granted? (VC2S8I06.E02)

Group A’s method, in full: cut three square canopies — 20 cm, 30 cm and 40 cm — from the same sheet of plastic; attach four equal 30 cm strings and the same 10 g load to each; drop each parachute indoors from a 2.0 m mark measured with a tape; time the fall from release until the load first touches the floor; three trials per canopy, and take the mean.

The parachute-drop investigation, set up to be fair



Labelled sketch of the parachute-drop investigation set up to be fair: a wall on the left carries a 2.0 m drop-height mark measured with a tape; a square canopy (here 30 cm × 30 cm) with four equal 30 cm strings hangs above a 10 g load; a hand releases it — let go, do not throw; a student at the right times the fall with a stopwatch from release until the load first touches the floor; caption: indoors, still air, same height every drop, three trials per canopy

Their means: 1.3 s for the 20 cm canopy, 1.8 s for 30 cm, 2.5 s for 40 cm. Conclusion: *for these three canopies, the bigger the canopy, the longer the fall time.* Now the lens. Some of that method is **checked** — written down and controlled on purpose. Some of it is **assumed** — taken for granted, never written, quietly doing work in the background.

What the parachute method checked — and what it assumed

Checked in the method	Quietly assumed
✓ canopy sizes cut and measured	the air in the room stays still
✓ drop height measured with a tape: 2.0 m, marked on the wall	the stopwatch counts seconds correctly
✓ same 10 g load on every parachute	each release is the same: let go, do not throw
✓ three trials for each canopy	fall time runs from release until the load first touches the floor

Assumptions are allowed in science — as long as you name them and ask how reasonable they are.

Two-column chart titled “What the parachute method checked — and what it assumed”: ticked green rows on the left — canopy sizes cut and measured; drop height measured with a tape at 2.0 m and marked on the wall; the same 10 g load on every parachute; three trials for each canopy — beside amber rows on the right — the air in the room stays still; the stopwatch counts seconds correctly; each release is the same, let go and do not throw; fall time runs from release until the load first touches the floor; caption: assumptions are allowed in science as long as you name them and ask how reasonable they are

Read the right-hand column again and notice something important: **none of it is a mistake**. Assuming the stopwatch counts correctly is sensible — you have to assume *something*, or no investigation could ever start. The skill is not avoiding assumptions; it is naming them and asking how reasonable each one is, in this setting. “The air stays still” is reasonable in a closed classroom with the doors shut and no one walking briskly past — and unreasonable outside on a windy day, which will matter in Lens 3. “Each release is the same” is the shakiest one: *let go, do not throw* is an instruction, and people are people. A good evaluation says exactly that — *this assumption is fine here, this one is shakier, and here is what it could do to the results*.

Try it 1 — name the hidden assumptions (VC2S8I06.E02)

Here is another group’s method, written exactly as they wrote it.

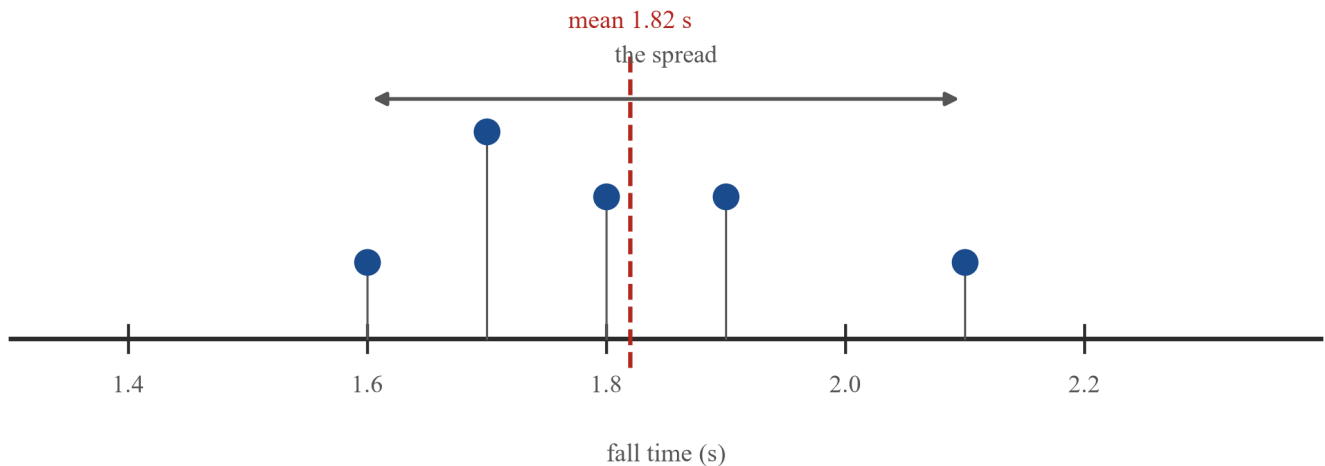
1. Pour 50 mL of water into a tray.
2. Lay one sheet of paper towel on the water for 10 s.
3. Lift the sheet out and let it drip for a moment.
4. Measure the water left in the tray.
5. Repeat three times for each of three brands; the brand that leaves the least water soaks up the most.

Copy the two-column idea from the parachute chart: write one thing this method **checks** and three things it **assumes**. (Hints live in the lifting, the dripping, the timing and the sheets themselves.) Then pick your three assumptions and grade each one *reasonable* or *shaky* for a classroom, with a half-line of reason.

Lens 2 — sources of error: what could have messed up the measurements? (VC2S8I06.E01)

Even a method with every assumption named still measures with wobbling hands. One student times the very same 30 cm drop five times and gets 1.6 s, 1.9 s, 1.7 s, 2.1 s and 1.8 s. Same drop, same student, same stopwatch — five different numbers. The mean is 1.82 s, and no single reading is “the” answer.

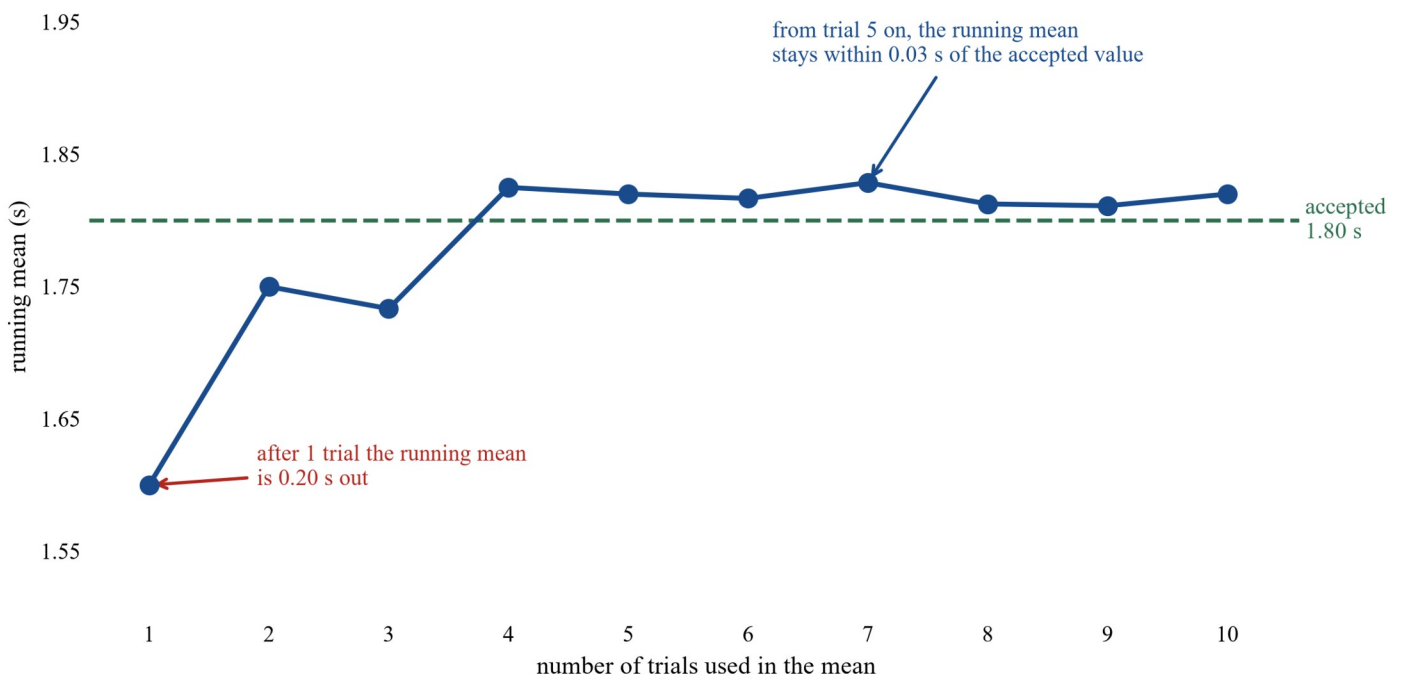
One student times the very same drop five times (simulated data for practice)



Dot plot titled “One student times the very same drop five times (simulated data for practice)”: five dots at 1.6, 1.7, 1.8, 1.9 and 2.1 s on a fall-time axis, a dashed line at the mean 1.82 s, and a double arrow marking the spread; the scatter on both sides of the mean is the fingerprint of chance (random) error

That scatter is **random error** — chance in the measurement. Your thumb is a little early on one drop, a little late on the next; the readings land on both sides of the true value. Random error is why one trial is never enough: a single reading can be luck. The good news is that random error can be squeezed. Average more trials and the high-lows cancel; the mean settles towards the accepted value, which is exactly what the next figure shows for ten trials of the same drop.

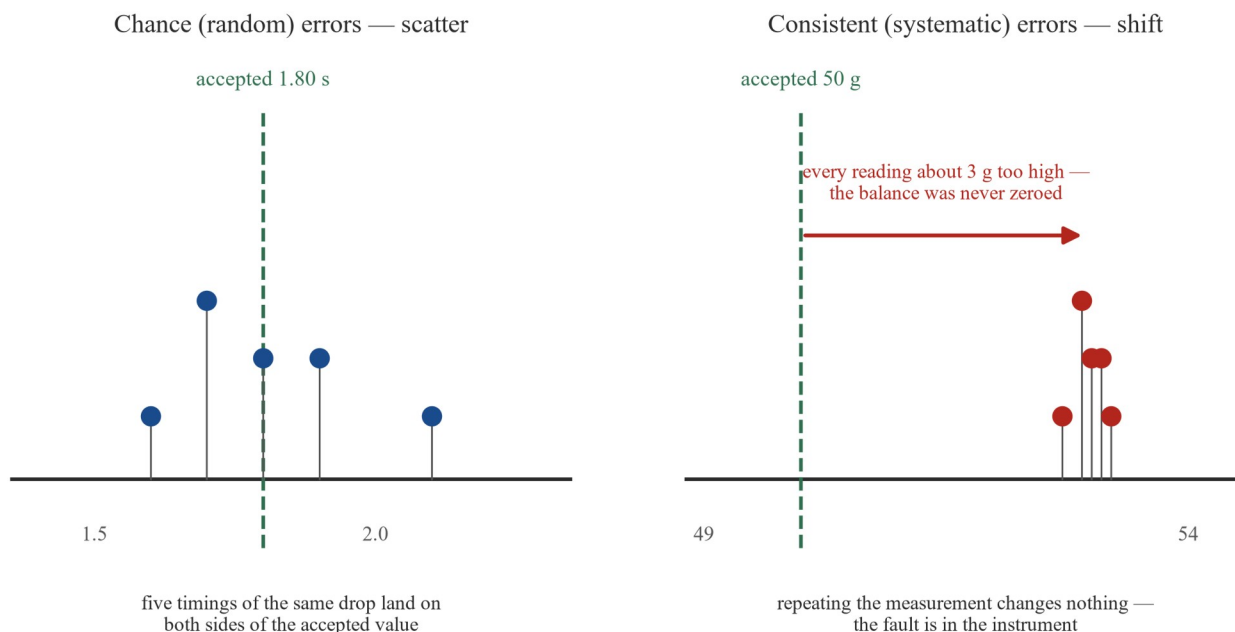
The more trials in the mean, the closer it settles to the accepted value (simulated data for practice)



Line graph titled “The more trials in the mean, the closer it settles to the accepted value (simulated data for practice)”: the running mean starts at 1.60 s after one trial (0.20 s out), wobbles through trials 2 to 4, and from about

trial 5 stays within 0.03 s of the accepted value 1.80 s, shown as a dashed green line; labels mark the one-trial mean and the settling point

Systematic error is the other kind, and it is sneakier because nothing wobbles. A balance was never zeroed: it reads 52.7 g, 53.1 g, 52.9 g, 53.2 g and 53.0 g for a mass whose accepted value is 50 g. Look at the two patterns side by side.



simulated data for practice

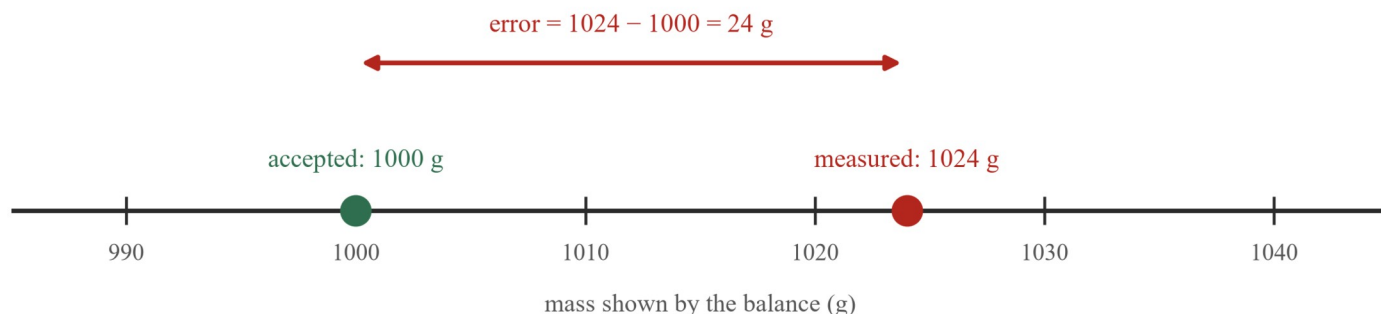
Two dot plots side by side. Left, “Chance (random) errors — scatter”: five timings of the same drop from 1.6 s to 2.1 s land on both sides of the accepted value 1.80 s, captioned “five timings of the same drop land on both sides of the accepted value”. Right, “Consistent (systematic) errors — shift”: five weighings of a 50 g mass cluster close together between 52.7 g and 53.2 g, every reading about 3 g too high because the balance was never zeroed, captioned “repeating the measurement changes nothing — the fault is in the instrument”. Labelled “simulated data for practice”

The difference is the shape. Random error **scatters** — readings on both sides of the accepted value. Systematic error **shifts** — every reading wrong in the same direction by roughly the same amount. And here is the trap the picture shows: repeating a systematically wrong measurement changes nothing. The mean of five shifted readings is a shifted mean. More trials fixes random error; only fixing the setup — zero the balance, check the stopwatch against a known time, shut the door on the draft — fixes systematic error. That pairing, *name the source of error → suggest the improvement that actually matches it*, is the heart of VC2S8I06.E01.

One more tool for Lens 2, borrowed from the D11 maths interlock (VC2M8N05): a single number that says how far off a measurement is. The **percentage error** is the error divided by the accepted value, times 100.

The balance reads 2.4% high — every mass it measures will be 2.4% too big

$$\begin{aligned}\text{percentage error} &= \text{error} \div \text{accepted value} \times 100 \\ &= 24 \div 1000 \times 100 = 2.4\%\end{aligned}$$



Number line titled “The balance reads 2.4% high — every mass it measures will be 2.4% too big”: a green dot at the accepted 1000 g and a red dot at the measured 1024 g, a red double arrow spanning the error, $1024 - 1000 = 24$ g, and a box showing percentage error = $\text{error} \div \text{accepted value} \times 100 = 24 \div 1000 \times 100 = 2.4\%$

Percentage error lets you compare errors of different sizes side by side: 24 g off a 1000 g mass and 24 g off a 24 g mass are very different stories, and the percentage says so. Rounding follows the usual rule from VC2M7N04 — keep as many figures as the data earns, and no more.

Evaluation written out looks like the next figure: a method on the left, and on the right what each lens finds when it looks.

A flawed method, pulled apart

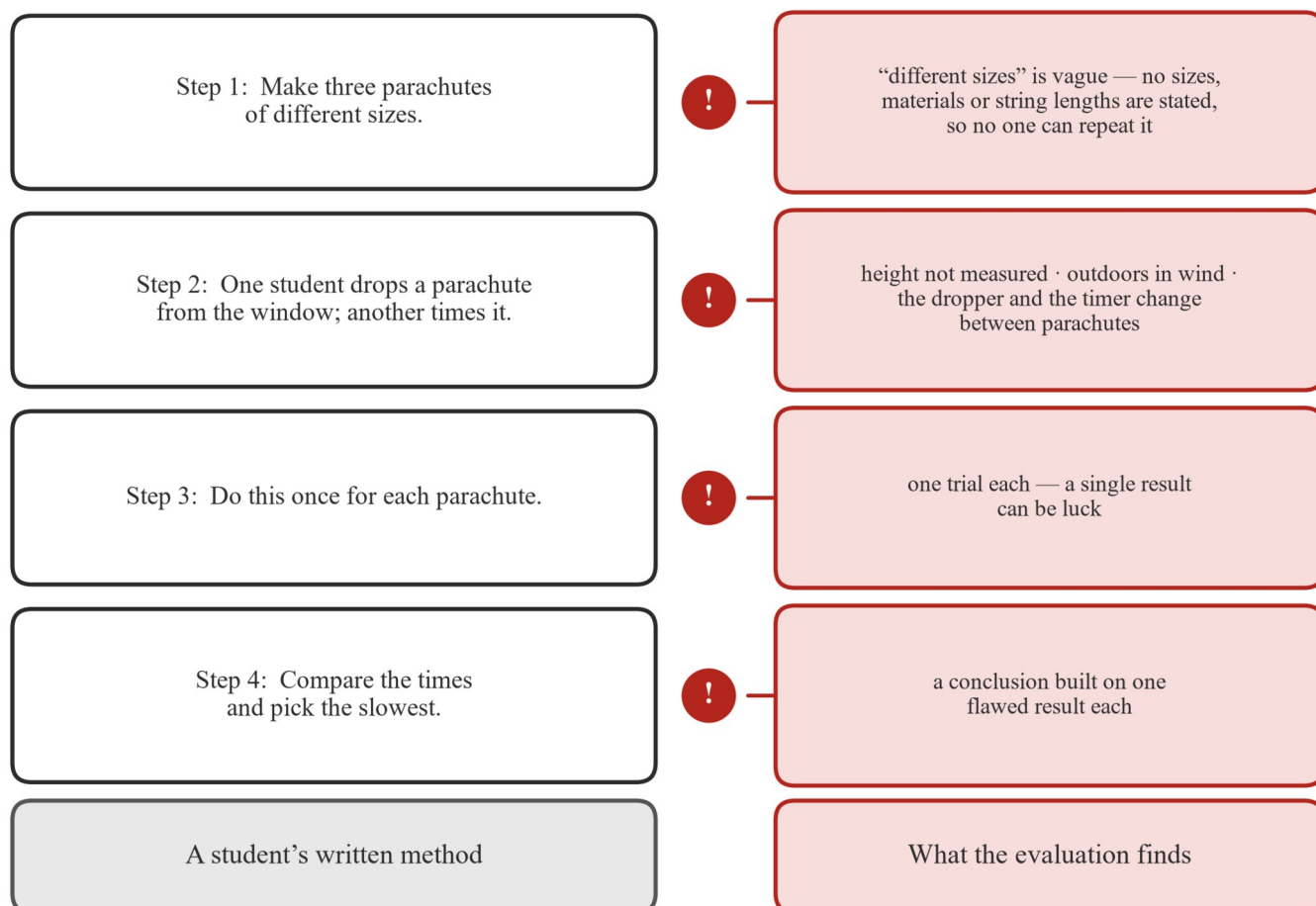


Diagram titled "A flawed method, pulled apart": four steps of a student's written method on the left — make three parachutes of different sizes; one student drops a parachute from the window while another times it; do this once for each parachute; compare the times and pick the slowest — each step carrying a red flag that points to what the evaluation finds on the right: the sizes, materials and string lengths are not stated so no one can repeat it; the height is not measured, it is outdoors in wind, and the dropper and the timer change between parachutes; one trial each, and a single result can be luck; so the conclusion is built on one flawed result each

Try it 2 — sort the errors, match the fix (VC2S8I06.E01)

For each source of error below, say *random* or *systematic* with a half-line of reason, then write the improvement that actually matches it (repeat-and-average, or fix the setup).

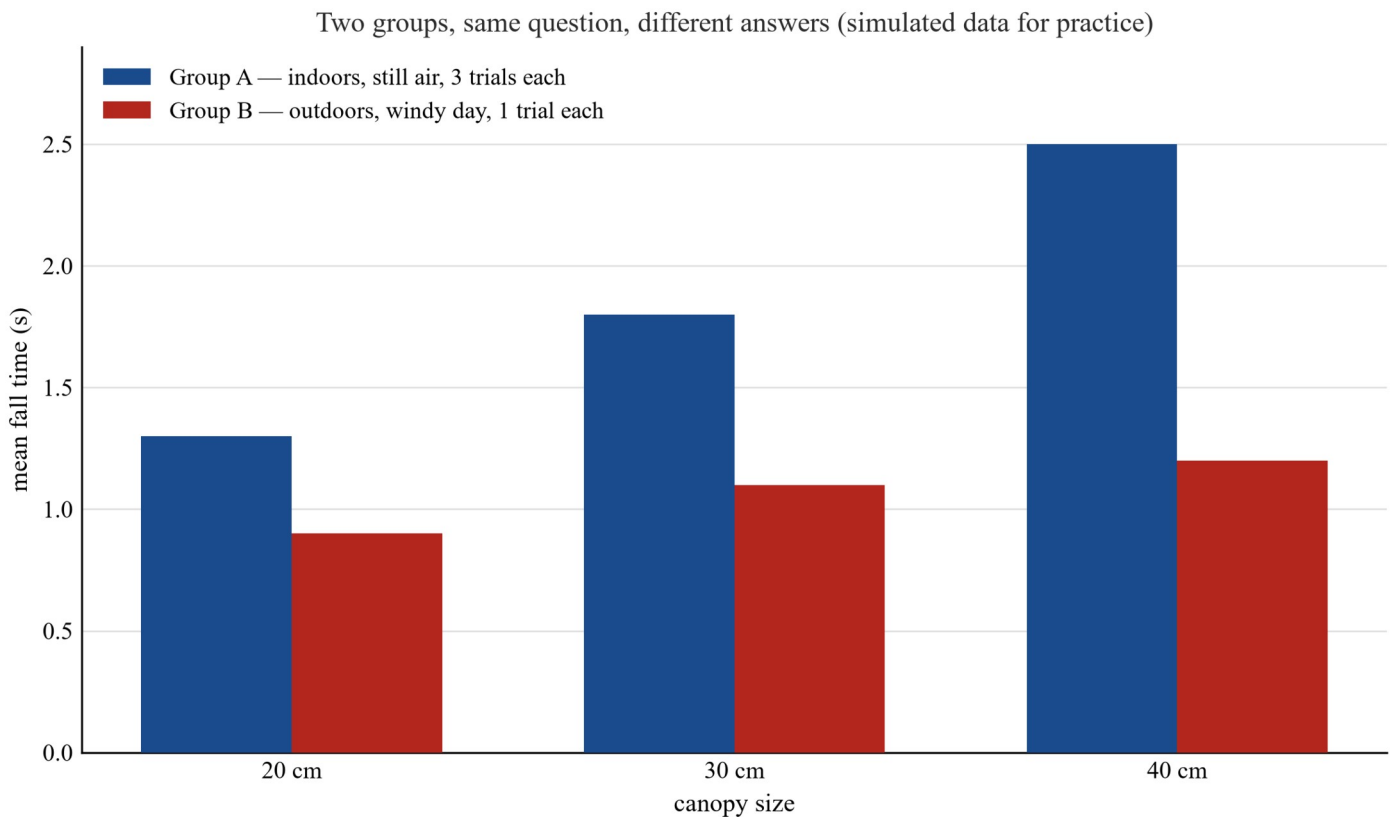
1. The timer's thumb is early on some drops and late on others.
2. The balance was never zeroed and reads about 3 g high, every time.
3. A door keeps opening; drafts cross the room now and then.
4. A stopwatch's battery is low and it runs slow all day.
5. The end of the ruler is worn, so every length measured from the end comes out 1 cm short.

Try it 3 — size the error up (VC2S8I06.E01, VC2M8N05)

- (a) A group measures the school oval with a trundle wheel and gets 24.5 m for a stretch whose accepted length is 25.0 m. Calculate the percentage error.
- (b) Which is the more serious error: 2 g off a 50 g mass, or 2 g off a 1000 g mass? Show both percentage errors to justify your answer.

Lens 3 — conflicting evidence: does other evidence disagree? (VC2S8I06.E04, .E06)

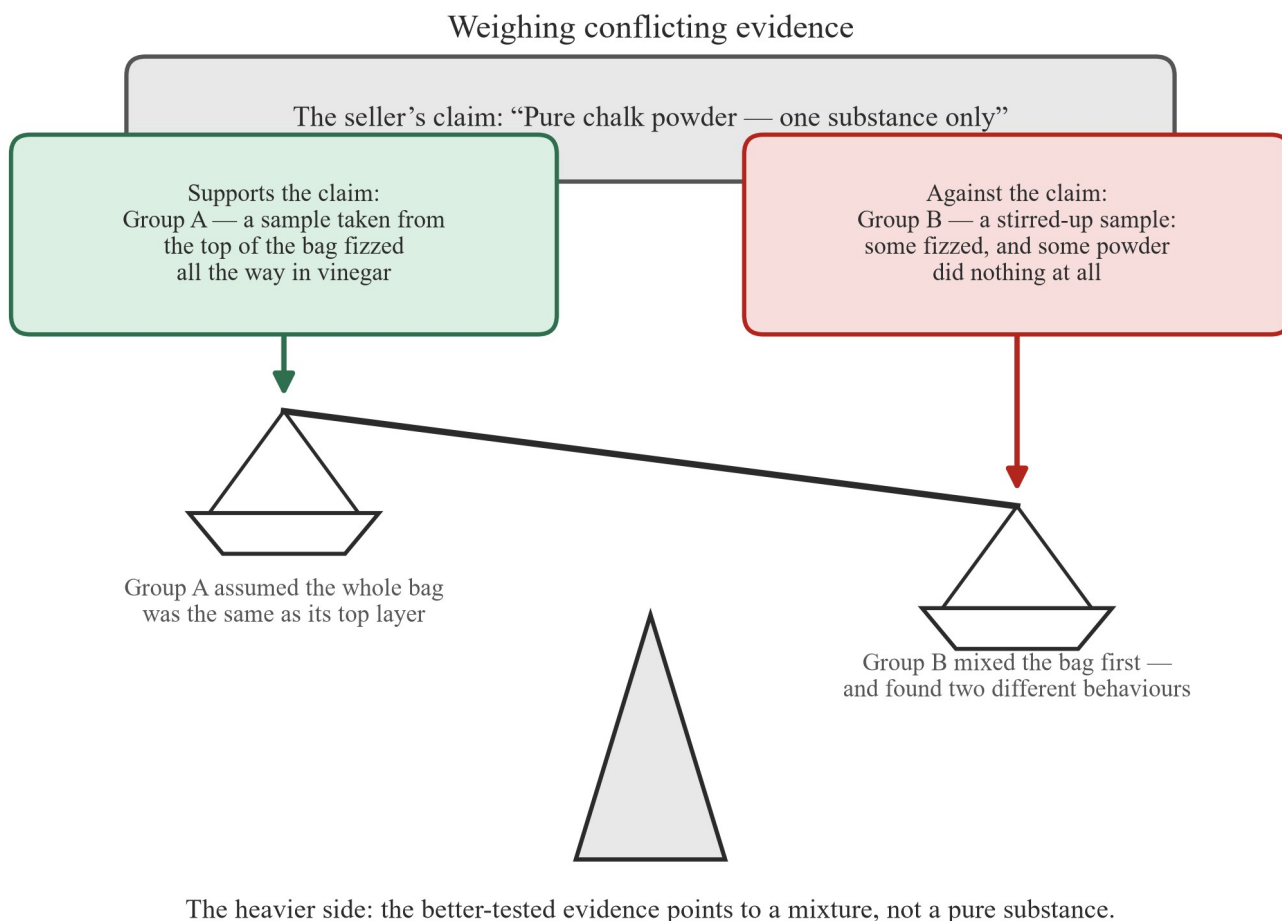
Back to the parachutes. Group B repeats “the same” investigation and gets means of 0.9 s, 1.1 s and 1.2 s — far below Group A’s 1.3 s, 1.8 s and 2.5 s. Two groups, same question, different answers. This is **conflicting evidence**, and the first move is never “somebody lied.” The first move (VC2S8I06.E06) is to line the methods up side by side and look for differences — and to check what other groups and secondary sources (results other people publish) found.



Bar chart titled “Two groups, same question, different answers (simulated data for practice)”: mean fall times for 20 cm, 30 cm and 40 cm canopies; Group A in blue — indoors, still air, 3 trials each — gets 1.3 s, 1.8 s and 2.5 s; Group B in red — outdoors, windy day, 1 trial each — gets 0.9 s, 1.1 s and 1.2 s; every Group B bar sits well below its Group A partner

The methods were not the same investigation at all. Group B worked outdoors on a windy day — the assumption “the air stays still” that was reasonable indoors is broken outdoors — and took one trial per canopy, so no averaging was done at all. (Why moving air changes a fall is the forces story of Ch15; Lens 3 does not need the *why* to do its job — the method difference alone tells you the two results cannot be averaged into one answer.) When evidence conflicts, the evaluation names the differences in method and says which evidence is better tested — more trials, steadier conditions, assumptions named. *Conflicting evidence does not mean no evidence; it means the methods have to be compared before the numbers are.*

The same lens works on claims, not just investigations. A market stall sells “pure chalk powder — one substance only.” Group A scoops powder from the top of the bag; every bit of it fizzes in vinegar; they conclude the seller is right. Group B stirs the bag first, then scoops; some of the powder fizzes and some does nothing at all.



Balance diagram titled “Weighing conflicting evidence”: the seller’s claim “Pure chalk powder — one substance only” hangs above two pans; the green pan — supports the claim: Group A, a sample taken from the top of the bag fizzed all the way in vinegar — rides high, while the red pan — against the claim: Group B, a stirred-up sample: some fizzed, and some powder did nothing at all — rides low; caption: the heavier side: the better-tested evidence points to a mixture, not a pure substance

Group A’s evidence is real — their sample did fizz. What their method assumed was that *the top of the bag is the same as the whole bag*, and a stirred sample is the better test of that assumption, because it gives every part of the bag a chance to be picked. Two behaviours in one stirred sample means the bag holds more than one substance: a **mixture**, not the pure substance the claim promised. The claim loses — not because Group A did nothing, but because conflicting evidence arrived, and the better-tested side outweighs it.

One assumption deeper still sits under this whole case (VC2S8I06.E03): that *fizzing in vinegar means a chemical change is happening — a new gas being made*. Nobody in the room tested the gas; everyone simply accepted it. That is how science usually works — most assumptions are so widely accepted, and so well tested in other places, that checking them again would be a waste of the afternoon. The point of naming them is not doubt; it is knowing what your conclusion stands on. (The gas itself, and the test for it, arrive in 11C.)

Lens 4 — unanswered questions: what is still open? (VC2S8I06.E05, .E07)

The last lens turns from the evidence you have to the evidence you don’t. A conclusion is always about what was tested — these three canopies, this height, this room — and the words of the conclusion decide how much it claims. Read a claim’s biggest words (*all, always, every, proves*) and ask what evidence they would take to hold up.

What evidence would this conclusion take?

“All buildings in an earthquake area should be made of bamboo.”

each step
is a kind
of evidence
the claim
would have
to earn

Material tests — bamboo shaken the same way as brick,
timber, steel and concrete, on a shaking table

Records from many earthquakes, in many towns,
over many years — not one building in one quake

Buildings of every kind and size — homes, schools
and taller buildings, not just small huts

The other dangers too — fire, rot, insects and storms,
not just earthquakes

Cost, supply and the skills to build with it — a material
nobody can afford or work saves no one

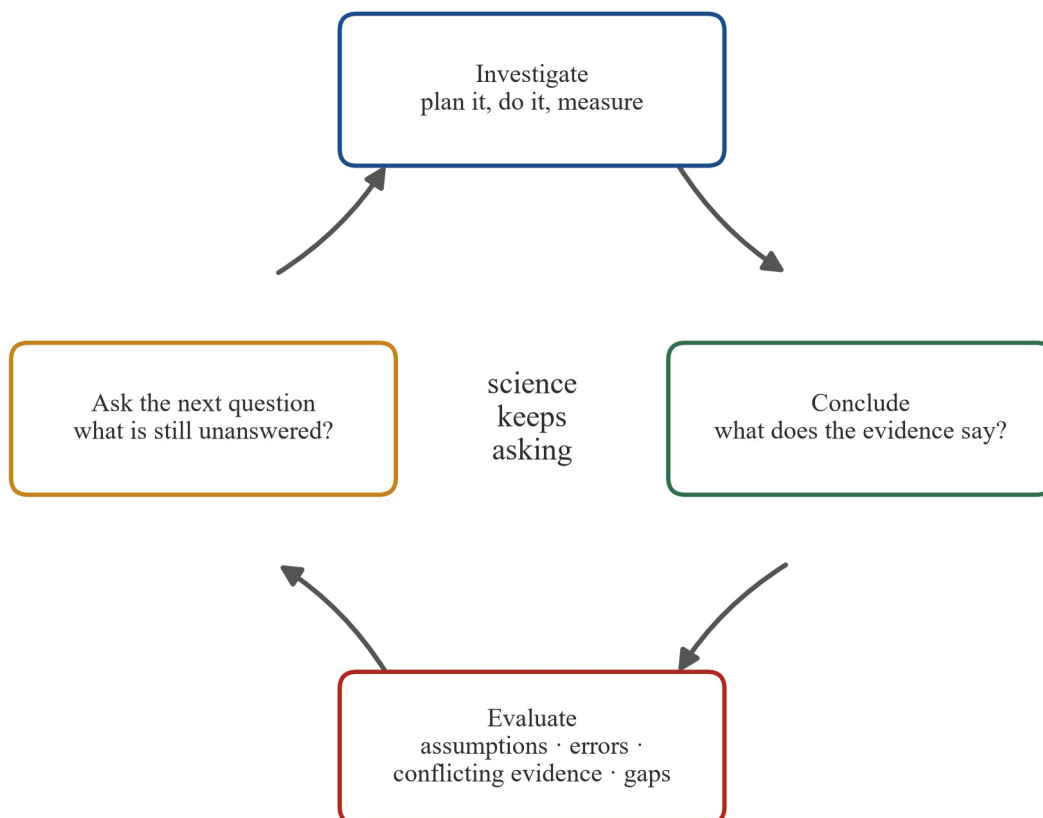
The bigger the words in a claim — all, always, proves — the more evidence it takes to hold it up.

Evidence-ladder diagram titled “What evidence would this conclusion take?”: the claim “All buildings in an earthquake area should be made of bamboo” stands above five narrowing steps, each a kind of evidence the claim would have to earn — material tests shaking bamboo the same way as brick, timber, steel and concrete on a shaking table; records from many earthquakes in many towns over many years, not one building in one quake; buildings of every kind and size, not just small huts; the other dangers too — fire, rot, insects and storms, not just earthquakes; and cost, supply and the skills to build with it; caption: the bigger the words in a claim — all, always, proves — the more evidence it takes to hold it up

The claim says *all buildings*, so it needs evidence about every kind of building, every kind of danger, and the real world of cost and skills — and a one-quake, one-town result is a first rung at best. Lens 4 does not say the claim is false; it says the claim is **bigger than its evidence**, and it lists exactly what is still unanswered. (VC2S8I06.E05 is precisely this habit: the bamboo claim is evaluated by asking what the evidence would have to look like, not by believing or dismissing it. Earthquakes themselves stay Ch12’s territory.)

And unanswered questions are not a failure of science — they are its fuel. A conclusion, finished the right way, is already pointing at the next investigation.

A conclusion is not the end — it is where the next question comes from



Cycle diagram titled “A conclusion is not the end — it is where the next question comes from”: four boxes in a loop — Investigate (plan it, do it, measure) leads to Conclude (what does the evidence say?), which leads to Evaluate (assumptions, errors, conflicting evidence, gaps), which leads to Ask the next question (what is still unanswered?), which leads back to Investigate — with “science keeps asking” at the centre

Group A’s conclusion — *bigger canopy, longer fall time, for these canopies in this room* — is a good conclusion, and it is full of next questions: bigger than 40 cm? smaller than 20? a different shape? a heavier load? a higher drop? Lens 4 writes those questions down on purpose, because the next investigation starts here (VC2S8I06.E07).

Try it 4 — weigh the claim (VC2S8I06.E04, .E06)

A fertiliser ad claims: “*Our feed makes every plant grow taller — tested by a class who grew one pot with the feed and one without, and the fed plant grew 4 cm more.*” Use Lens 3 and Lens 4: (a) name one assumption the one-plant-each test makes about pots and plants; (b) explain why a second class’s result — no difference with the feed — is not “somebody lied”, and what the two classes should compare first (VC2S8I06.E06); (c) write one unanswered question the ad’s conclusion leaves open.

Try it 5 — ask the next questions (VC2S8I06.E05, .E07)

- Write three unanswered questions left open by Group A’s parachute conclusion. At least one must be about a range the method never tested, and at least one about an assumption from the amber column of the assumption chart.
- Take one rung of the bamboo evidence ladder and write it as the question it answers. (Example: the second rung answers “*Do bamboo buildings hold up in many earthquakes, in many towns, over many years — or in just the one quake?*”)

The whole story in four lenses

Evaluation reads a method and a conclusion through four lenses. Lens 1 lists what the method took for granted and grades each assumption for reasonableness, because assumptions are allowed once named. Lens 2 finds the sources of error — random error scatters and is squeezed by more trials and averaging; systematic error shifts and is fixed only by fixing the setup — and can size an error up as a percentage error. Lens 3 brings in conflicting evidence and compares methods before trusting numbers, weighing the better-tested evidence more heavily and naming even the widely accepted assumptions underneath. Lens 4 lists the questions a conclusion still leaves open, because a conclusion only ever covers what was tested, and bigger words need bigger evidence. None of the four lenses calls anyone stupid; together they answer the only question evaluation asks — *how much does this conclusion actually earn?*

Worked examples

Example 1 — scaffolded

A student writes up a sugar-cube investigation: “*We put sugar cubes into hot, warm and cold water. We dropped each cube from above the surface. I stirred the hot one until the sugar was gone. We wrote down how long each cube took to disappear. Conclusion: heat makes sugar disappear faster.*” Evaluate the method.

- (a) Name one assumption the method makes.
- (b) Name one source of error or unfairness, and say whether it is random or systematic.
- (c) Suggest one improvement that matches the problem you named in (b).

Working	Comment
(a) The three cubes are the same size, so each one is the same amount of sugar to make disappear.	An assumption is anything taken for granted and not checked — the method never says the cubes were matched. “Stirring does not matter” and “the water stays at its starting temperature” are also acceptable.
(b) Only the hot cup was stirred — stirring is a second change riding along with the heat, so the method is not comparing like with like. That is a built-in, same-direction problem: the hot cube always gets the extra stirring help, every time. That is the shape of a systematic error in the method.	Systematic = a built-in shift, the same every time. Random error would be timing by thumb, a little early, a little late.
(c) Treat all three cups the same: stir all of them the same way (or none at all), and keep every other thing — cube size, water volume, dropping height — the same.	The improvement must match the problem: fix the setup. Repeating and averaging would not help here, because the unfairness repeats too.

Example 2 — scaffolded

Use the chalk-powder case and the weighing diagram.

- (a) State the assumption Group A’s sampling made about the bag.
- (b) Explain why Group B’s stirred sample is the better test of “what is in the whole bag”.
- (c) Given both results, what conclusion about the seller’s claim is now supported — and name one question that stays unanswered.

Working	Comment
(a) That the top layer of the bag is the same as the rest of it — a scoop from the top stands for the whole bag.	Lens 1: the assumption was never checked, and it is exactly the thing Group B’s stirring checks.
(b) Stirring mixes the whole bag before sampling, so every part of the bag can be picked; a scoop from an unstirred bag only ever tests one layer.	Better-tested evidence = the method that actually tests what the claim is about (the whole bag).
(c) The claim “one substance only” is not supported: one stirred sample showed two behaviours, so the bag is a mixture. Unanswered: <i>what are the substances in the</i>	Lens 3 weighs the claim against conflicting evidence; Lens 4 writes down what the conclusion still does not

mixture? — naming the non-fizzing powder needs the chemical-change tools of 11C.

say.

Example 3 — unscaffolded

A balance reads 52.7 g, 53.1 g, 52.9 g, 53.2 g and 53.0 g for a 50 g mass.

- Show, from the shape of the readings, whether this is random or systematic error.
- Calculate the mean of the readings and the percentage error of that mean against the accepted 50 g (VC2M8N05). Round sensibly (VC2M7N04).
- A student says “let’s take twenty more trials and the mean will fix it.” Is that right? Explain.

The readings all sit between 52.7 g and 53.2 g — a tight cluster, but every reading is on the same side of the accepted value, about 3 g high. Nothing scatters around 50 g. Same direction, same rough size, every time: that is the shift shape, so the error is **systematic** (a balance that was never zeroed fits perfectly). **(b)** Mean = $(52.7 + 53.1 + 52.9 + 53.2 + 53.0) \div 5 = 264.9 \div 5 = 52.98$ g, so about 53.0 g. Error = $52.98 - 50 = 2.98$ g. Percentage error = $2.98 \div 50 \times 100 = 5.96\%$, which rounds to **6.0%**. **(c)** No. More trials squeezes *random* error; a systematic shift repeats in every extra trial, so the mean of twenty shifted readings is still shifted. The fix is in the setup: zero the balance before weighing.

Example 4 — unscaffolded

After one earthquake in one town, a report notes that the bamboo buildings in the town stood while some others fell, and concludes: “*All buildings in an earthquake area should be made of bamboo.*” Evaluate the conclusion using Lens 3 (conflicting evidence) and Lens 4 (unanswered questions).

Lens 3: the conclusion rests on one quake, one town, one set of buildings, and it checked nothing against anything — no records from other earthquakes or other towns, no comparison of bamboo against other materials shaken the same way, no look at what other groups or secondary sources have published. Any of those could hold conflicting evidence; none was sought, so the evidence side of the scales is a single observation. Lens 4: the claim’s words are *all buildings*, so the unanswered questions are the rungs of the ladder — does bamboo hold up in many quakes over many years, in buildings of every kind and size, against the other dangers too (fire, rot, insects, storms), and at a cost and skill real towns can afford? The one-town result does not answer any of them. Verdict: the observation is real and worth a next question, but the claim is bigger than its evidence — evaluate it as a hypothesis to test, not a conclusion to build with.

Practice questions

Scaffolded

Q1. About the four lenses. (a) Name the two lenses that look at the method. (b) Name the two lenses that look at the conclusion. (c) Which lens asks “what did the method take for granted?”

Q2. About the parachute assumption chart. (a) Name one thing the parachute method checked. (b) Name one thing it quietly assumed. (c) Is an assumption a mistake? Answer in one sentence.

Q3. About random and systematic error. (a) Which kind of error scatters readings on both sides of the true value? (b) Which kind shifts every reading the same way? (c) Repeating and averaging squeezes which kind?

Q4. About the five stopwatch readings (1.6 s, 1.9 s, 1.7 s, 2.1 s, 1.8 s). (a) Calculate the mean. (b) The spread of these readings is which kind of error? (c) Suggest one improvement that would shrink it.

Q5. About percentage error. (a) Write the percentage-error rule in words. (b) Calculate it for a measured 1024 g against an accepted 1000 g. (c) If a balance reads 2.4% high, what does that mean for every mass it measures?

Q6. About the two parachute groups. (a) Describe, in one sentence, how Group B’s results differ from Group A’s. (b) Name two differences between the groups’ methods. (c) What is the first move when two groups’ evidence conflicts (VC2S8I06.E06)?

Un scaffolded

Q7. Use the flawed-method diagram. Pick two of the four red flags. For each, say which lens it belongs to (assumptions or sources of error) and one line on how it weakens the conclusion. (2 marks)

Q8. Back to Try it 1's paper-towel method. Name one assumption it makes and one source of error in it, and give the improvement that matches your error. (3 marks)

Q9. A balance reads 205 g for a 200 g mass. (a) State the error. (b) Calculate the percentage error. (c) If the balance reads 5 g high for *every* mass, is that random or systematic? One line of reason. (3 marks)

Q10. Use the running-mean graph. (a) After 1 trial, how far out was the running mean? (b) From roughly which trial does it stay within 0.03 s of the accepted value? (c) In one sentence, what does this say about trial number and random error (VC2M8ST03)? (3 marks)

Q11. Use the chalk-powder case. (a) State Group A's sampling assumption. (b) Which pan of the weighing diagram carries the better-tested evidence, and why? (c) Name the widely accepted assumption about fizzing that both groups stood on (VC2S8I06.E03). (3 marks)

Q12. Use the bamboo evidence ladder. Write two of the rungs as unanswered questions — the questions those rungs of evidence would answer. (2 marks)

Q13. A class grows one pot of seedlings with music playing and one in silence, measures each plant once, and concludes "*music makes plants grow taller.*" Evaluate the conclusion through any two lenses you choose, naming each lens you use. (3 marks)

Q14. Explain, in two or three sentences, why "the mean of three trials beats one trial" is true for random error but no help at all for systematic error. (2 marks)

Answers

Q1. (a) Assumptions; sources of error. (b) Conflicting evidence; unanswered questions. (c) The assumptions lens (Lens 1). **Q2.** (a) Any one of: canopy sizes cut and measured; drop height 2.0 m measured with a tape; same 10 g load; three trials per canopy. (b) Any one of: the air stays still; the stopwatch counts correctly; each release is the same; timing runs to first floor touch. (c) No — it is something taken for granted without checking, and every method has them; it only becomes a problem when it is unreasonable and unnamed. **Q3.** (a) Random. (b) Systematic. (c) Random. **Q4.** (a) 1.82 s. (b) Random error. (c) Any matching one: more trials and the mean; the same student timing every drop; timing from a video of the fall. **Q5.** (a) Error divided by the accepted value, times 100. (b) $24 \div 1000 \times 100 = 2.4\%$. (c) Every mass it shows is 2.4% too big — a systematic shift. **Q6.** (a) Group B's fall times are much shorter for every canopy size. (b) Any two of: outdoors in wind versus indoors in still air; one trial each versus three trials; (also acceptable) no averaging done. (c) Compare the methods (and other groups' and secondary sources' results) before trusting either set of numbers. **Q7.** Example pair: "do this once for each parachute" → sources of error (one trial = a single result can be luck, nothing averaged); "outdoors in wind, height not measured" → assumptions (the method takes for granted that conditions are the same for each parachute, and they are not). **Q8.** Assumption: every sheet lifts out and drips the same way / the sheets of each brand are the same size and thickness. Error: "let it drip for a moment" is un-timed, so drip time changes from sheet to sheet — random. Improvement: time the drip (say 5 s each) — or lift and hold each sheet the same way. **Q9.** (a) 5 g. (b) $5 \div 200 \times 100 = 2.5\%$. (c) Systematic — same direction, same size, every time. **Q10.** (a) 0.20 s out (1.60 s against 1.80 s). (b) From about trial 5. (c) More trials in the mean squeeze random error, so the mean settles closer to the accepted value. **Q11.** (a) The top layer of the bag is the same as the whole bag. (b) The red (against) pan — Group B stirred first, so the sample could come from the whole bag, not just one layer. (c) That fizzing in vinegar means a chemical change — a new gas being made. **Q12.** Any two rung-shaped questions, e.g. "Does bamboo hold up in many earthquakes in many towns over many years, or just in one quake?" and "Does it hold up in buildings of every kind and size, not just small huts?" **Q13.** Any two lenses well used. E.g. Lens 1: one pot each — pot, soil, light, seed differences are all unnamed assumptions; Lens 2: one plant per condition with one measurement each, nothing repeated or averaged, so the 4 cm-style difference could be luck; Lens 4: the conclusion says "music makes plants grow taller" generally, far beyond two pots measured once. **Q14.** Random readings scatter high and low, so averaging cancels the scatter; a systematic shift moves every reading the same way, so the mean of shifted readings is just as shifted.

Fully-worked solutions

Q1. (a) The two method lenses are **assumptions** and **sources of error**. (b) The two conclusion lenses are **conflicting evidence** and **unanswered questions**. (c) **Lens 1, the assumptions lens** — "what did the method take for granted?"

Q2. (a) Any checked row, e.g. **the drop height was measured with a tape: 2.0 m, marked on the wall**. (b) Any amber row, e.g. **the air in the room stays still**. (c) **No — an assumption is something taken for granted without checking, not a slip; every method has them, and they only weaken a conclusion when they are unreasonable or left unnamed**.

Q3. (a) **Random error** scatters on both sides of the true value. (b) **Systematic error** shifts every reading the same way by roughly the same amount. (c) Repeating and averaging squeezes **random** error; systematic error repeats in every trial and must be fixed in the setup.

Q4. (a) $\text{Mean} = (1.6 + 1.9 + 1.7 + 2.1 + 1.8) \div 5 = 9.1 \div 5 = \mathbf{1.82 \text{ s}}$. (b) The scatter on both sides of the mean is **random error** — chance in the timing. (c) A matching improvement: **more trials, then the mean** (or the same student timing every drop, or timing from a video of the fall).

Q5. (a) $\text{Percentage error} = \text{error} \div \text{accepted value} \times 100$. (b) $\text{Error} = 1024 - 1000 = 24 \text{ g}$; $24 \div 1000 \times 100 = \mathbf{2.4\%}$. (c) **Every mass the balance shows comes out 2.4% too big** — the same-direction, same-size signature of a systematic shift.

Q6. (a) **Group B's mean fall times (0.9 s, 1.1 s, 1.2 s) sit far below Group A's (1.3 s, 1.8 s, 2.5 s) at every canopy size**. (b) Any two: **Group B worked outdoors on a windy day (breaking the still-air assumption); Group B took one trial per canopy (no averaging)**. (c) **Compare the methods first — and check other groups' results and secondary sources — before trusting either set of numbers** (VC2S8I06.E06).

Q7. One mark per flag correctly placed with its lens and a reason. Example: “do this once for each parachute” → **Lens 2, sources of error — one trial means a single result can be luck and nothing is averaged, so each time in the conclusion could simply be wrong (1 mark).** “the height is not measured, outdoors in wind, dropper and timer change between parachutes” → **Lens 1, assumptions — the method takes for granted that every parachute gets the same drop, and it does not, so the times are not comparing like with like (1 mark).**

Q8. Assumption (1 mark): e.g. **the three brands’ sheets are the same size and thickness, so one sheet stands for one brand** (or: each sheet is lifted and dripped the same way). Error (1 mark): “let it drip for a moment” is not timed — **drip time changes from sheet to sheet by chance, a random error** (or: lifting speed differs, random). Matching improvement (1 mark): **hold the drip time fixed — e.g. 5 s over the tray for every sheet — or lift every sheet the same slow way**; note that repeat-and-average alone would not remove an un-timed drip done a different way each time.

Q9. (a) Error = $205 - 200 = 5$ g (1 mark). (b) $5 \div 200 \times 100 = 2.5\%$ (1 mark). (c) **Systematic — every reading is high by the same 5 g, same direction, same size; nothing scatters around 200 g (1 mark).**

Q10. (a) **0.20 s out** — the one-trial mean is 1.60 s against the accepted 1.80 s (1 mark). (b) **From about trial 5 the running mean stays within 0.03 s of 1.80 s (1 mark).** (c) **The more trials that go into a mean, the more the random high-lows cancel, so the mean settles closer to the accepted value** — one trial is the most luck you can have (1 mark).

Q11. (a) **That the top layer of the bag is the same as the whole bag** — a scoop from the top stands for everything (1 mark). (b) **The red pan (against the claim)** — Group B stirred the bag first, so their sample could come from any part of the bag; two behaviours in one stirred sample is evidence about the whole bag, which Group A’s top-layer scoop never was (1 mark). (c) **That fizzing in vinegar means a chemical change — a new gas being made** — accepted by everyone in the room without being tested (the gas and its test arrive in 11C) (1 mark).

Q12. One mark per rung-shaped question. E.g. “**Do bamboo buildings hold up across many earthquakes in many towns over many years, or was this one quake lucky?**” (1 mark) and “**Do they hold up against the other dangers too — fire, rot, insects and storms — not just shaking?**” (1 mark). Questions must match a rung: the evidence the rung names is the answer the question seeks.

Q13. Any two lenses, each named and used, earn a mark; the third mark is for a verdict. Example. **Lens 1 (assumptions): the two pots are assumed identical in everything but music — soil, light, water, seeds — and none of it is checked or even stated (1 mark).** **Lens 2 (sources of error): one plant per condition, measured once — no repeats, no mean, so any difference could be luck (random error at full strength), and if one pot was measured with a stretchy ruler the error would ride along unseen (1 mark).** Verdict: **the conclusion “music makes plants grow taller” is bigger than its evidence — the result is a first observation, not a tested conclusion (1 mark).** (Lens 4 used well — e.g. *which music, which plants, how much taller, for how long?* — also earns the marks.)

Q14. **Random error scatters readings high and low around the true value, so averaging three trials cancels much of the scatter and the mean lands closer to the truth (1 mark).** **Systematic error shifts every reading the same way — three shifted readings average to a shifted mean — so the only fix is the setup itself (zero the balance, mend the ruler) (1 mark).**

Sources

- Victorian Curriculum and Assessment Authority, *Victorian Curriculum F–10 Version 2.0, Science, Levels 7 and 8*, content description VC2S8I06 and its elaborations, 2026. vcaa.vic.edu.au — retrieved 2026-08-21.

All numerical series in the figures and questions (the parachute group means, the five stopwatch timings, the ten-trial running means, the never-zeroed balance readings, the 1024 g balance, the 205 g and 200 g masses, the trundle-wheel 24.5 m) are **simulated data for practice**, labelled as such in each figure that plots them; they are teaching data, not measurements, and carry no source claim. The parachute, chalk-powder, sugar-cube, paper-towel and bamboo cases are original teaching scenarios; the bamboo claim is evaluated as a method-of-evaluation exercise only, with the earthquake science itself reserved for Ch12 (D9).

Elaboration contexts used: VC2S8I06.E01 (the random/systematic error pairing and improvement-matching), .E02 (the parachute method and its assumption chart), .E03 (the widely accepted fizz-means-chemical-change assumption under the chalk case), .E04 (the chalk-powder claim weighed against conflicting evidence), .E05 (the bamboo claim and its evidence ladder), .E06 (the two-group comparison and secondary-source habit), .E07 (the conclusion-to-next- question loop).

Teacher-facing note (not student text). 5A owns VC2S8I06 outright (single delivery, D6) and is taught in the inquiry front block (D8), so every science-context reference is, by design, an everyday observation or a forward reference: forces and why wind matters are Ch15's, the gas behind the fizz and its test are 11C's, earthquakes are Ch12's, fair-test planning is 3A's and graphing conventions are 4A/4C's — none are taught here. The percentage-error work is the D11 interlock with VC2M8N05 (with VC2M8ST02/VC2M8ST03 on spread and trial number and the VC2M7N04 rounding cap); students meeting it before the maths chapter should treat the rule as given. The random/systematic distinction is taught and named here per §5.2 — the content descriptions never use those words, so 5A names and defines them in the Word watch. No Aboriginal and Torres Strait Islander content is drawn on here: VC2S8I06 carries no such elaborations, and D10's sourced treatments sit in the subtopics the TOC names.

Building an evidence-based argument: source credibility, evaluating claims, and the ethics and cultural protocols of using secondary data

What we teach here — VC2S8I07 (Victorian Curriculum F–10 Version 2.0, Science, Levels 7 and 8):

“evidence-based arguments can be constructed to support conclusions or evaluate claims, including consideration of ethical issues and protocols associated with using or citing secondary data or information” **This subtopic owns the whole CD:** there is no partner subtopic (single delivery, D6) — 5B

carries every clause of the CD. It is the construction-and-evaluation subtopic of the book: where 5A taught you to take apart a method and a conclusion after the measuring, this one teaches you to build an argument that survives checking — and to check the arguments other people hand you. **In our words:** every day someone asks you to believe something — an ad, a video, a headline, a friend. Science’s answer is always the same: show your evidence, say where you got it, and show the link between the two. Here you get the tools for both directions: a scorecard for deciding whether a source is worth believing, a ladder for grading how strong evidence is, a test for how far a claim reaches past what was actually checked, and the rules for using someone else’s data honestly. **Achievement-standard extract served here (D13):** *“They provide science-based explanations for findings, and use evidence to support conclusions and evaluate claims.”* 5B owns VC2S8I07 outright, so the extract appears here in full. **Elaborations drawn on:** VC2S8I07.E01 (determining whether a source is credible), .E02 (evaluating the quality of evidence of primary and secondary sources used when constructing an argument), and the first limb of .E03 (acknowledging sources when using or citing secondary data — the ethics of citation). **ICIP (D10):** no Aboriginal or Torres Strait Islander content is delivered here. The second limb of VC2S8I07.E03 (respecting cultural protocols about the sharing of particular information), VC2S8I07.E04 and VC2S8I07.E05 are reserved under the ICIP decision and are not taught; see the teacher-facing note at the foot of this subtopic. The title keeps the words *cultural protocols* because TOC §3 titles are fixed verbatim. **Scope notes:** the ceiling is held (D9). Australia’s warming number from *State of the Climate 2024* appears here as a real claim to evaluate — the *causes* of a changing climate, and any mechanism behind it, are Levels 9–10 content and are not taught in this book; the number is a setting for argument skills, never something we explain. The maths we lean on is the D11 interlock: ethical and fair methods for drawing conclusions about a whole group (VC2M8ST04), which appears here as the fairness rules for data about people. Planning the investigation that generates data is 3A’s job; reading trends from processed data is 4C’s; here we argue from data and judge the sources it came from.

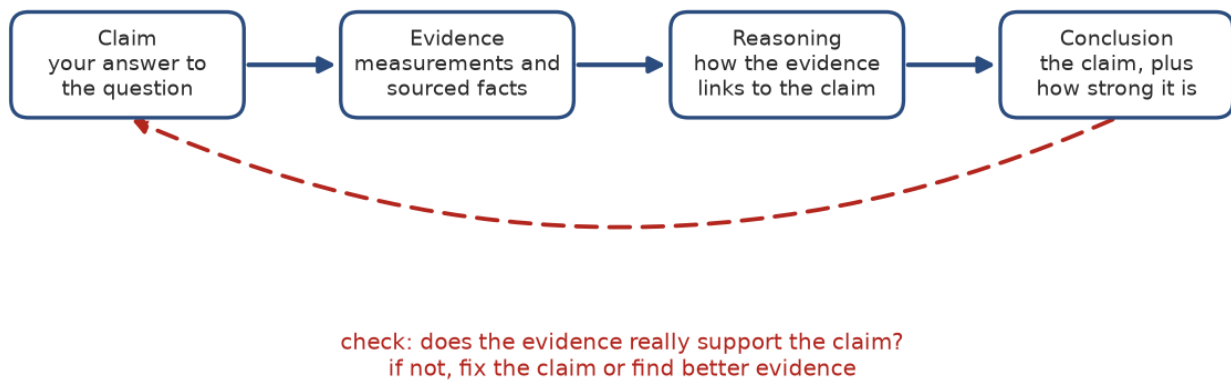
Word watch — *argument, claim, evidence, reasoning, credibility, citation, protocol*. In science an **argument** is not a fight; it is a claim with its evidence and reasoning attached. A **claim** is a statement someone asks you to believe (the same word 5A used); **evidence** is the observations and measurements standing behind it; **reasoning** is the link — the sentence that says *why* that evidence supports that claim. **Credibility** needs one special note: the Victorian Curriculum uses the term *credibility* (VC2S8I07.E01) but does not define it. This subtopic therefore treats it as *how much a source deserves to be believed when it makes a claim*, decided by the five scorecard questions below, and every question here is answerable on that definition alone. A **citation** is a written note that tells the reader exactly where a piece of information came from; to **cite** is to write one. A **protocol** is an agreed set of rules about how something is to be done correctly.

Theory

The parts of an evidence-based argument

Here is an argument with nothing behind it: *chips should stay on the canteen menu, because they are the best lunch*. Somebody’s claim, standing alone, asking to be believed. Science does not accept arguments shaped like that — and neither should you. An evidence-based argument has four parts, and it is only finished when all four are in place.

The parts of an evidence-based argument



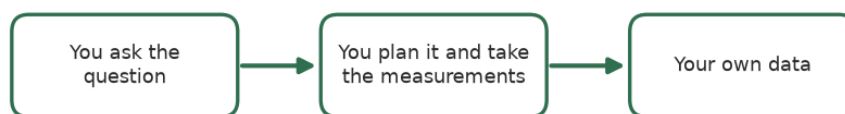
Flow chart titled “The parts of an evidence-based argument”: four blue boxes in a row — Claim (your answer to the question), Evidence (measurements and sourced facts), Reasoning (how the evidence links to the claim), Conclusion (the claim, plus how strong it is) — with a dashed red arrow looping back from the conclusion to the claim, labelled: check — does the evidence really support the claim? If not, fix the claim or find better evidence

The **claim** is your answer to the question. The **evidence** is the measurements and sourced facts you are standing on. The **reasoning** is the part people forget: the sentence that joins the evidence to the claim, so a reader can see *why* the numbers point where you say they point. And the **conclusion** is the claim with a strength attached — not just *this is what I found*, but *this is how strongly the evidence holds it up*. The dashed arrow is the habit that makes the whole thing honest: before you hand the argument to anyone, walk back around and ask whether the evidence really supports the claim. If it does not, you fix the claim or you find better evidence. You never leave the gap there.

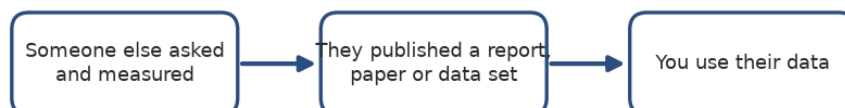
Two routes to evidence: primary and secondary data

Every number in an argument reached somebody by one of two routes.

Primary data — you generate it



Secondary data — someone else generated it



e.g. you read the Bureau of Meteorology's temperature record for Australia

Two-row chart titled “Two routes to the data in an argument”: the top green row, primary data — you generate it — flows through you ask the question, you plan it and take the measurements, your own data, with the example “your class times how long ice takes to melt in three cups”; the bottom blue row, secondary data — someone else generated it — flows through someone else asked and measured, they published a report, paper or data set, you use their data, with the example “you read the Bureau of Meteorology’s temperature record for Australia”

Primary data is data you generated yourself — the measurements from an investigation you planned and ran (that is 3A’s territory). **Secondary data** is data somebody else generated and published: government reports, journal papers, data sets, museum records, news articles. You were not standing there when it was measured, so every secondary number arrives with a question glued to it: *who measured this, and why should I believe them?*

That question is exactly what this subtopic is about. The CD’s phrase “using or citing secondary data or information” points at the kind of data you will use most for the rest of your life — nobody measures the national temperature record themselves; we use the people whose job it is. Secondary data is not automatically weaker than primary data: a careful national record beats a class experiment on any question about the whole country. What decides is not who measured it but how well the source stands up when you ask it questions. (Reading processed and secondary data for patterns and trends is 4C’s skill; here we grade the source itself.)

Is the source credible? The five questions (VC2S8I07.E01)

A **source** is wherever information comes from — a report, a website, a video, a person. To decide whether a source deserves to be believed, ask five questions.

A source-credibility scorecard		
The question to ask	A strong answer	A weak answer
Who is behind it?	Bureau of Meteorology — the national climate agency	a video from an account with no name or qualifications
Are they expert on this topic?	climate scientists reporting measured temperature records	someone selling a product that the claim supports
Are they independent?	no gain from one result or another	paid to get the result the ad wants
Is it up to date — and dated?	the 2024 report, with its data and methods listed	no date, no data, no way to check
Can you check it somewhere else?	the same finding appears in other agencies’ records	no other source has it

A source that answers weakly on several rows is not a reliable source — keep looking.

Table-style chart titled “A source-credibility scorecard”: five rows, each with a blue question box and a green strong answer and a red weak answer. Who is behind it? — strong: Bureau of Meteorology, the national climate agency; weak: a video from an account with no name or qualifications. Are they expert on this topic? — strong: climate scientists reporting measured temperature records; weak: someone selling a product that the claim supports. Are they independent? — strong: no gain from one result or another; weak: paid to get the result the ad wants. Is it up to date — and dated? — strong: the 2024 report, with its data and methods listed; weak: no date, no data, no way to check. Can you check it somewhere else? — strong: the same finding appears in other agencies’ records; weak: no other source has it. Caption: a source that answers weakly on several rows is not a reliable source — keep looking

Notice what the five questions are really doing: they are asking *who*, *why*, *when* and *can anyone confirm* — the same instincts you use when a friend tells you something wild, written down so they can be applied on purpose. One caution: credibility is always about *this* claim. A famous science agency can still be wrong outside its own field, and a

shop can tell you the truth about its own prices. The scorecard is never “is this a famous name”; it is always “does this source deserve to be believed *about this question?*”

Try it 1 — work the scorecard (VC2S8I07.E01)

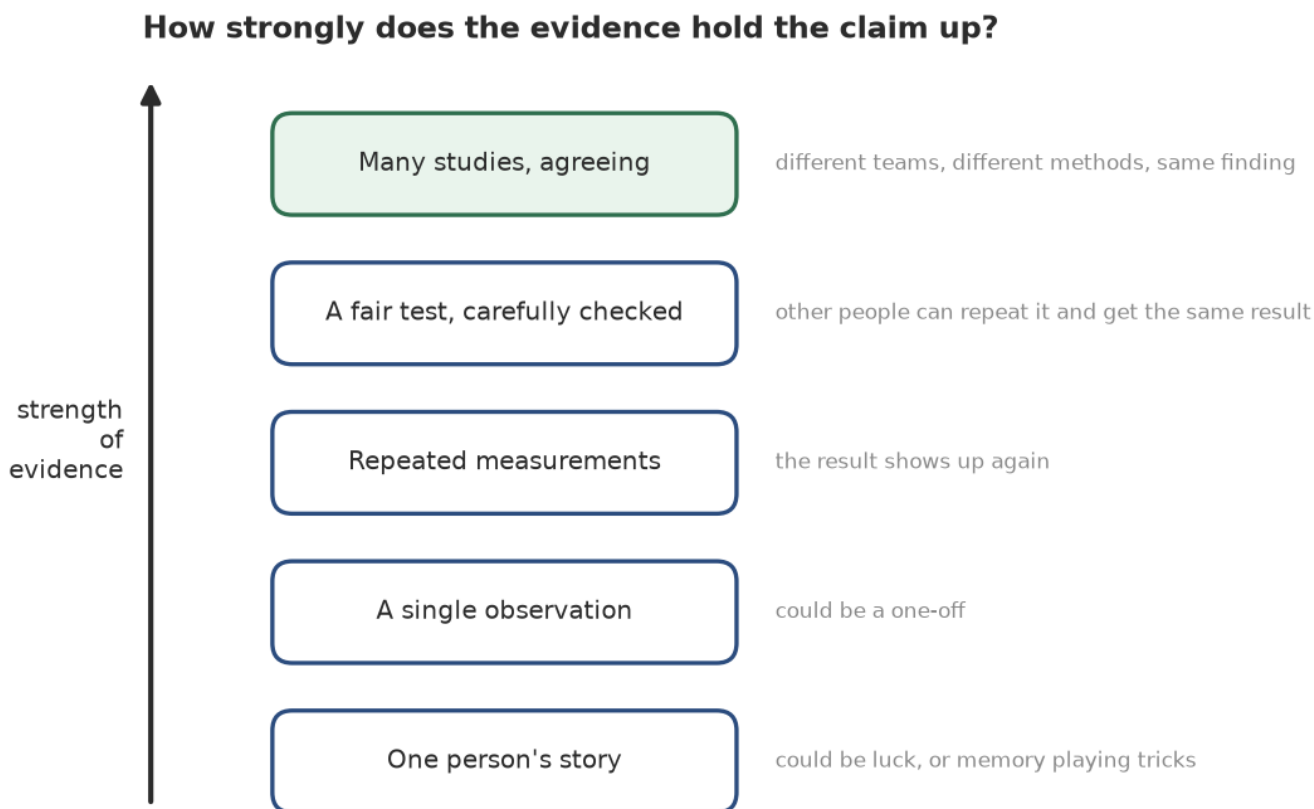
A claim is going around: “*There are fewer bees in the school garden than there were five years ago.*” Two sources take a position (both sources are simulated, for practice).

- **Source A:** a report by insect scientists at a state museum, dated 2025, listing the counting method, the sites and the years covered.
- **Source B:** a post on social media from an account with no name: “bees disappearing!!! i saw like 2 bees today.”

Pick two rows of the scorecard and work them for both sources, one line each. Then give a verdict: which source is more credible about the bee claim, and why?

How strong is the evidence? The ladder (VC2S8I07.E02)

A credible source is only half the job; the evidence itself has a strength. The ladder grades it.



A claim resting on the bottom rung may still be true — it is just not well supported yet.

Stacked-box diagram titled “How strongly does the evidence hold the claim up?”: five equal boxes rising from the bottom, with an arrow labelled strength of evidence beside them — one person’s story — could be luck, or memory playing tricks; a single observation (could be a one-off); repeated measurements (the result shows up again); a fair test, carefully checked (other people can repeat it and get the same result); many studies, agreeing (different teams,

different methods, same finding). Caption: a claim resting on the bottom rung may still be true — it is just not well supported yet

The bottom rung is one person's story. A single story is not a lie — it is a real experience, and it is often where a question starts. But one story can be luck, or memory playing tricks, which is why it can start a question and never finish one. Each rung up is the same idea — *more of the evidence, checked more carefully* — until the top rung, where different teams using different methods all land on the same finding. That is the strongest shape evidence comes in.

Two habits from 5A carry straight over. Better-tested evidence outweighs weaker evidence — and the ladder tells you what a claim still needs, not whether the claim is false. A bottom-rung claim might be completely true; it is just not well supported yet.

A real claim to practise on: the trend underneath the wiggles (VC2S8I07.E02)

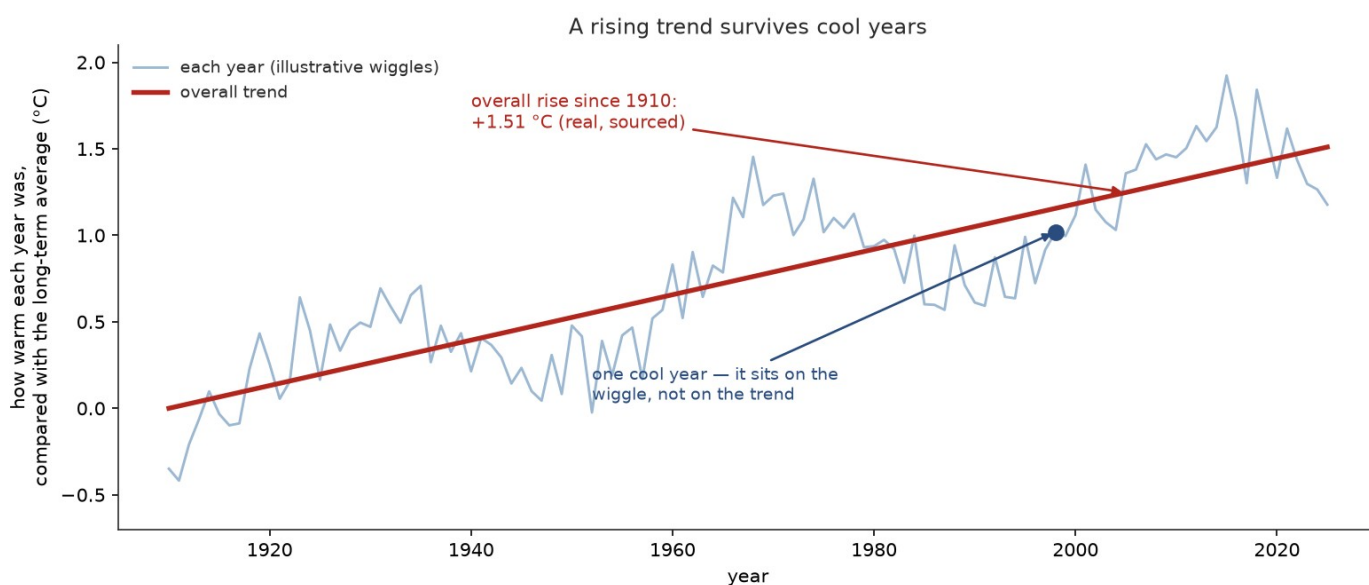
Here is a real claim, sourced, that you will meet again in this subtopic. Australia's official climate report says this:

Australia's climate has warmed by an average of 1.51 ± 0.23 °C since national records began in 1910.

Bureau of Meteorology and CSIRO, "State of the Climate 2024 — Report at a glance", 2024.

<https://www.bom.gov.au/weather-and-climate/past-weather-and-climate/state-of-the-climate-2024/report-at-a-glance> — retrieved 2026-08-21.

Run the scorecard on it: the national weather agency and the national science agency, reporting measured records, with no gain from either answer, dated, and checkable against other agencies' records. Five strong answers. Now look at what the number actually says — because the shape of the record is where arguments get won and lost.



The year-to-year wiggles in this drawing are illustrative only; the size of the rise is the Bureau of Meteorology and CSIRO figure for Australia's near-surface air temperature (State of the Climate 2024).

Line graph titled "A rising trend survives cool years": a light blue line of year-by-year values wiggles up and down between 1910 and 2025 while a bold red trend line rises steadily through it; one cool year, 1998, is marked with a blue dot and labelled "one cool year — it sits on the wiggle, not on the trend"; the red line is labelled "overall rise since 1910: +1.51 °C (real, sourced)". A note under the axis states that the year-to-year wiggles in the drawing are illustrative only; the size of the rise is the Bureau of Meteorology and CSIRO figure for Australia's near-surface air temperature (State of the Climate 2024)

Every year sits somewhere on the wiggles; the trend is what survives when you average over all of them. And this is where a very common move appears. Somebody picks one cool year — or three — and announces that the warming "stopped". Picking only the bits of the record that suit your claim is called **cherry-picking**, and it is the single most common way evidence gets misused.

The strongest form of the cool-year objection deserves a real answer, not a shrug: *the records can't be right, because I remember some bitterly cold winters, and 2012 was a cool summer where I live*. That objection is honest, and it has an honest answer. Your memory of one place in one season is exactly what the record calls a wiggle — real, but local. The claim is about the average of a whole continent across every year since 1910, and one cool year, or ten, does not remove a trend built from over a hundred of them. The test is the same one 5A used: before you trust the numbers, check what the method covered. A claim built on one picked year has covered one year.

The tool cuts both ways — that is the point of it. Someone quoting a single hot summer to claim *the hottest year ever* — *it's all over!* is standing on the same bottom rung. Trend claims are argued from the whole record, from both sides.

Try it 2 — sort the evidence onto the ladder (VC2S8I07.E02)

Four students offer evidence for the claim “*the creek behind the school is getting dirtier*”. Put each one on a rung of the ladder (bottom to top), one line of reason each.

1. “My aunty says the creek used to be clear when she was at school.”
 2. One class measures the creek’s flow once, on a Tuesday.
 3. The class measures the creek’s flow every week for a term.
 4. Three nearby schools test the same creek the same year and all find the same drop in flow.
-

Try it 3 — spot the cherry-pick (VC2S8I07.E02)

Someone arguing that *Australia is not warming* lists three cool years from the record — 1998, 2008 and 2012 — and says: “there you go, it keeps cooling.” Using the trend figure, answer: (a) what did the argument leave out? (b) What would a fair use of the record look like? (c) Would your answer change if the same trick were played with three hot years to argue the warming is *faster* than 1.51 °C? One line.

How far does the claim reach? (VC2S8I07.E01, .E02)

Advertisements are the best practice ground in the world for these tools, because their whole job is to make claims sound strong. Work on this one (a fictional ad, drawn for practice — no real product is being talked about).

Claim reach: how far does the claim go past the evidence?

what was actually checked:
20 people tasted the cereal and
said if they felt fresher (simulated)

the claim reaches far
beyond what was checked

the ad's claim: “BrainBoost helps you concentrate better in class — energy all morning!”

The bigger the gap, the weaker the support — no matter how confident the claim sounds.

BrainBoost is a fictional product — simulated ad for practice.

Diagram titled “Claim reach: how far does the claim go past the evidence?”: a short green box shows what was actually checked — 20 people tasted the cereal and said if they felt fresher (simulated) — while a long red arrow labelled with the ad’s claim, “BrainBoost helps you concentrate better in class — energy all morning!”, stretches far beyond it; a bracket marks the gap, “the claim reaches far beyond what was checked”. Caption: the bigger the gap,

the weaker the support — no matter how confident the claim sounds. The ad is for a fictional product, BrainBoost, simulated for practice

Claim reach is the distance between what a claim says and what was actually checked. The ad checked whether twenty people *felt* fresher; the claim is about *concentrating better in class, all morning*. Those are different claims — the second one is bigger, and nothing was measured to hold it up. The bigger the reach, the weaker the support, and it does not matter how confident the claim sounds.

Now set the ad’s claims next to what was checked.

BrainBoost: claims against evidence	
The claim (the advertisement)	The evidence (what was checked)
"Helps you concentrate better in class"	the only check: 20 people paid to taste it — no concentration test was done (simulated)
"Energy all morning — no crash"	one serve has 9 g of sugar — more than many 'regular' cereals (simulated)
"Recommended by experts"	no experts are named, so the claim cannot be checked

A confident claim is not the same as a well-supported one. Check the reach against the evidence.

BrainBoost is a fictional product — simulated ad for practice.

Two-column chart titled “BrainBoost: claims against evidence”: left red column, the ad’s claims — “helps you concentrate better in class”, “energy all morning — no crash”, “recommended by experts”; right blue column, what was checked — the only check was 20 people paid to taste it, with no concentration test done (simulated); one serve has 9 g of sugar, more than many regular cereals (simulated); no experts are named, so the claim cannot be checked. Caption: a confident claim is not the same as a well-supported one — check the reach against the evidence. The ad is for a fictional product, BrainBoost, simulated for practice

Notice the third claim: “recommended by experts” — with no experts named. Run the scorecard on a claim with no named source: *who is behind it?* has no answer, so the claim cannot even be checked. And the ladder puts *twenty people paid by the seller* exactly where it belongs: near the bottom. The reach test is 5A’s big-words habit in new clothes: read a claim’s biggest words — *every, always, best, all morning* — and ask what evidence would take to hold them up.

Try it 4 — measure the reach (VC2S8I07.E01, .E02)

A poster for a fictional drink, ZoomFizz, says: “ZoomFizz — helps you remember facts faster! Tested by scientists!” In small print: the only test was ten people paid to drink it, then asked if they felt sharper. (Fictional product, for practice.)

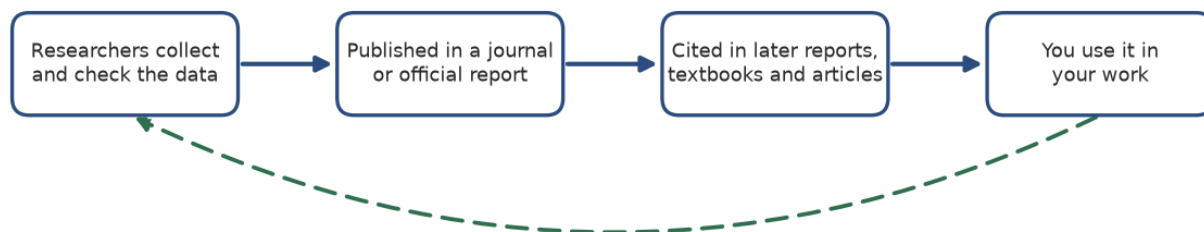
- (a) State the claim the poster is making.
- (b) State what was actually checked.
- (c) In one sentence, measure the reach: does the evidence hold the claim up, and what is missing?

Using secondary data honestly (VC2S8I07.E03)

Now the other direction: not judging other people's claims, but using other people's data in your own argument. This is where ethics comes in — and the ethics are concrete, not vibes.

Start with where secondary data travels. Somebody measured it, somebody published it, it got cited in later works, and eventually it reached you.

The citation chain: where secondary data travels



cite every step you use, so any reader can walk the chain back to the original

Flow chart titled “The citation chain: where secondary data travels”: four blue boxes in a row — researchers collect and check the data; published in a journal or official report; cited in later reports, textbooks and articles; you use it in your work — with a dashed green arrow looping back from your work to the researchers, labelled: cite every step you use, so any reader can walk the chain back to the original

A **citation** is how you keep the chain unbroken. When you use a number, a sentence or an idea that came from a source, you write a line that lets any reader walk back to where it came from. Three reasons, all of them real:

1. **Credit.** The people who measured and checked the data did the work; the citation says the work is theirs.
2. **Checking.** Your reader is allowed to ask the same questions you asked — *is this source credible?* — and the citation is the address they need to go and check.
3. **Honesty.** Using someone's words or numbers without saying where they came from is passing their work off as yours. In every school and every science, that is the one thing you do not do.

Every source line in this book has the same five parts.

The parts of a source line



Every source line in this book answers all five questions. Copy the shape when you write your own.

A source line with its five parts labelled, titled “The parts of a source line”: the line “Bureau of Meteorology and CSIRO, ‘State of the Climate 2024 — Report at a glance’, 2024. bom.gov.au/.../report-at-a-glance — retrieved 2026-08-21” with five colour-coded labels pointing down to boxes — who (the people or group behind it), what (the title, in

quote marks), when (the year it came out), where (the full web address), retrieved when (the day you saw it). A note says the web address is shortened on the figure; in your own source line, write it out in full. Caption: every source line in this book answers all five questions — copy the shape when you write your own

That figure labels the five parts of the very source line for the warming number you met above: *who* made it, *what* it is called, *when* it came out, *where* to find it, and *when you saw it* — because a web page can change, and the retrieval date tells the reader which version you were reading.

Two more rules before you use anything. First, **permission**: not everything published may be copied. Some sources publish their work with a clear *you may share this* attached; the rules an owner sets for how their work may be used are called a **licence**, and keeping to them is part of honest use. If there is no permission and no licence, you ask before you use. Second, **faithfulness**: quoting a source means keeping its meaning. Snipping half a sentence so a source seems to say the opposite of what it said is not citation — that is the cherry-pick, wearing a source line.

One kind of data carries extra rules: data **about people**. If your evidence is a survey — *what do Year 8 students like best?* — the method has to be fair as well as careful, and fairness here has three parts.

Data about people: fair and unfair methods		
The rule	A fair method	An unfair method
Ask permission first	everyone is told what the data is for, and agrees to take part	names taken from a class list without asking anyone
Ask a fair spread of people	students from every year level, so all groups have a say	only the people likely to give the answer you want
Keep identities safe	results reported for the group as a whole — no names	named students' answers posted on the board

When the data is about people, the method has to be fair too: permission, a fair spread, and no names.

Three-row chart titled “Data about people: fair and unfair methods”: blue rule boxes beside green fair and red unfair columns. Ask permission first — fair: everyone is told what the data is for, and agrees to take part; unfair: names taken from a class list without asking anyone. Ask a fair spread of people — fair: students from every year level, so all groups have a say; unfair: only the people likely to give the answer you want. Keep identities safe — fair: results reported for the group as a whole, no names; unfair: named students’ answers posted on the board. Caption: when the data is about people, the method has to be fair too — permission, a fair spread, and no names

Those three rules are the D11 maths interlock (VC2M8ST04): ethical and fair methods for drawing conclusions about a whole group from the people you actually asked. Ask only your friends and you have not measured the year level — you have measured your friends. And notice that the third rule, *only the people likely to give the answer you want*, is the cherry-pick again, this time applied to people instead of years. Same move, same weakness.

Everything in this section fits on one checklist.

Before you use secondary data: four checks

 Trust it does the source pass the credibility scorecard?	 You are allowed to use it is there permission to use it, or a clear “you may share this”?
 It stays faithful does it keep the source's meaning, without stretching it to fit your claim?	 It is credited is there a source line a reader can follow?

If any check fails, don't use it yet — fix that first.

Two-by-two checklist titled “Before you use secondary data: four checks”: trust it (does the source pass the credibility scorecard?); you are allowed to use it (is there permission, or a clear “you may share this”); it stays faithful (does it keep the source’s meaning, without stretching it to fit your claim?); it is credited (is there a source line a reader can follow?). Caption: if any check fails, don’t use it yet — fix that first

Try it 5 — fix the citation (VC2S8I07.E03)

A student writes this in a report:

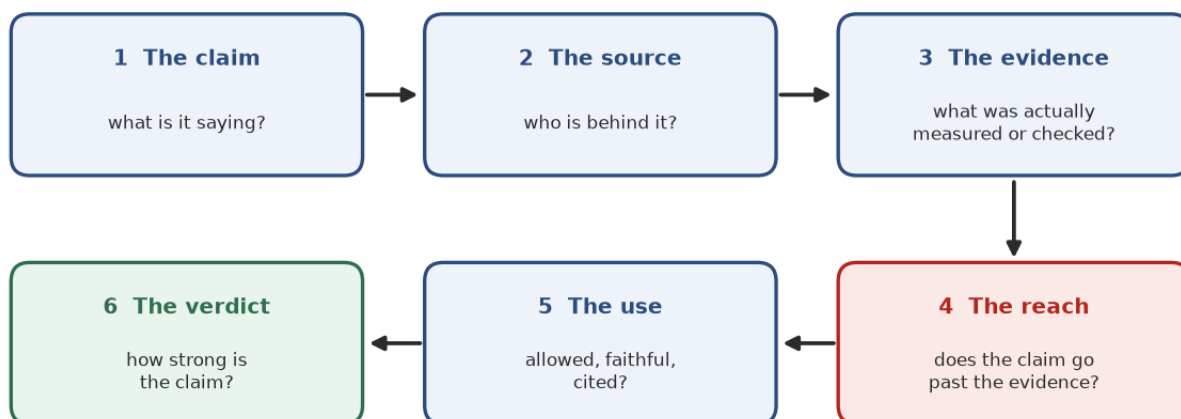
Australia’s climate has warmed by an average of 1.51 °C since national records began in 1910. This shows the warming is real.

The first sentence is word-for-word from the *State of the Climate* page, and the student has not said so. (a) Name the two problems with that paragraph. (b) Rewrite the two sentences so the source is credited. (c) Write the source line for it, using the five-part shape.

Evaluating a claim, start to finish

The whole subtopic fits in one walk-through. When anyone hands you a claim — an ad, an article, a friend, a report — this is the order of operations.

Evaluating a claim, start to finish



Each step uses one of this subtopic's tools. The verdict says how strong the claim stands — “well supported”, “weakly supported”, or “cannot be checked”.

Six-step flow chart titled “Evaluating a claim, start to finish”: top row left to right — 1 the claim (what is it saying?), 2 the source (who is behind it?), 3 the evidence (what was actually measured or checked?); then down to a bottom row running right to left — 4 the reach (does the claim go past the evidence?), 5 the use (allowed, faithful, credited?), 6 the verdict (how strong is the claim?). Caption: each step uses one of this subtopic's tools; the verdict says how strong the claim stands — “well supported”, “weakly supported”, or “cannot be checked”

Keep the verdict words exact. **Well supported** means the source is credible, the evidence is high on the ladder, and the claim does not reach past what was checked. **Weakly supported** means the claim outruns its evidence — it might still be true, but it is not there yet. **Cannot be checked** means no source is named or the source cannot be found — and a claim that cannot be checked gets no belief at all until it can be.

The whole story in one pass

An evidence-based argument is a claim, evidence, reasoning and a conclusion with its strength stated — checked on the way back before it is handed to anyone. The evidence arrives by one of two routes: primary (you measured it) or secondary (someone else did, and published it). Before you use secondary data you run the source through five questions — who, expert, independent, dated, checkable — and the evidence itself through the ladder, from one person's story up to many studies agreeing. You watch two favourite tricks: cherry-picking the convenient years (or the convenient people), and letting a claim reach farther than anything measured. And when you use someone else's data in your own argument, you keep it honest: trust it, be allowed, keep it faithful, and credit it with a source line any reader can follow. None of this is about catching people out. It is about building arguments that deserve to survive — and knowing exactly how strong yours is when you hand it over.

Worked examples

Example 1 — scaffolded

A claim is made: “*The creek through town is cleaner than it was ten years ago.*” Two sources take a position (both sources simulated, for practice).

- **Source A:** a water-quality report from the state government's environment department, 2025, written by water scientists, listing the method and every measurement across ten years.
 - **Source B:** a web page selling land next to the creek: “*Our creek is pristine — come for a swim!*” No data, no date, no author named.
- (a) Work the *who is behind it?* row of the scorecard for both sources.
 - (b) Work the *are they independent?* row for both sources.
 - (c) Which source is more credible about the creek claim, and why?

Working	Comment
(a) Source A: water scientists at the government environment department — named experts whose job is measuring water quality. Source B: a land seller; no author named, no expertise shown.	The first question asks who, and whether they are expert on <i>this</i> claim.
(b) Source A: no gain from either answer — the department’s job is to report the measurements. Source B: a direct gain — clean creek, higher land prices.	Independence means no benefit from one result or the other. A seller is not automatically wrong, but the claim must then be checked somewhere independent.
(c) Source A is more credible: experts, independent, dated, method and data listed, checkable. Source B answers weakly on every row of the scorecard.	A source that fails several rows is not a reliable source — keep looking.

Example 2 — scaffolded

A student’s report contains this sentence, copied without credit from the *State of the Climate* page: “Australia’s climate has warmed by an average of 1.51 °C since national records began in 1910.”

- Name the two problems with using the sentence this way.
- Rewrite the sentence so the source is credited in the text.
- Write the source line for it, using the five-part shape.

Working	Comment
(a) No credit — the measured and checked work is presented as the student’s own; and no way for the reader to check it — the citation is the address.	Credit and checking are two of the three reasons citations exist; honesty covers both.
(b) Australia’s climate has warmed by an average of 1.51 °C since national records began in 1910 (Bureau of Meteorology and CSIRO, <i>State of the Climate 2024</i>).	Crediting in the text can be as simple as naming the makers and the report the number came from.
(c) Bureau of Meteorology and CSIRO, “State of the Climate 2024 — Report at a glance”, 2024. https://www.bom.gov.au/weather-and-climate/past-weather-and-climate/state-of-the-climate-2024/report-at-a-glance — retrieved 2026-08-21.	Who, what, when, where, retrieved when. The retrieval date matters because pages change.

Example 3 — unscaffolded

Evaluate the BrainBoost ad (fictional, for practice) using the three tools of this subtopic — scorecard, ladder, reach — and give a verdict.

Scorecard. The source of every claim in the ad is the company that sells the cereal. *Who is behind it?* — the seller. *Independent?* — no: the company gains directly from every claim being believed. *Checkable?* — “recommended by experts” names no experts, so there is nothing to check against. Two failed rows and one empty one. **Ladder.** The only thing checked was twenty people paid to taste the cereal and say if they felt fresher — near the bottom rung: no concentration test, no comparison, no repeat by anyone else. **Reach.** The claims are about concentrating better in class and having energy all morning; what was checked is whether twenty people *felt* fresher. The claim reaches far past the evidence, and the 9 g of sugar per serve sits awkwardly next to “energy all morning — no crash”. **Verdict:** weakly supported. The ad’s claims outrun everything that was measured; none of them has earned the strength they are spoken with.

Example 4 — unscaffolded

Construct an evidence-based argument from the warming number: write the claim, the evidence, the reasoning and the conclusion — and then state two limits of your argument.

Claim: Australia’s climate has warmed, on average, since national records began in 1910. **Evidence:** the Bureau of Meteorology and CSIRO report *State of the Climate 2024*: an average warming of 1.51 °C, with a range of 0.23 °C either side of that average, across the national record since 1910. **Reasoning:** the number is an average over the whole continent and every year of the record, so it is not the story of any single year; a trend that survives over a hundred and

sixteen years of wiggles, measured by the national weather agency and checkable against other agencies' records, is the shape of well-supported evidence. **Conclusion:** the claim is well supported — credible source, top-of-ladder evidence, no reach beyond what was measured. **Limits.** (1) The argument says *that* the warming happened, and nothing about *why* — the causes of a changing climate are not part of this book, so the argument stops at the number and does not step past it. (2) The number is an average with a range; individual years sit on the wiggles, so the argument never supports a claim about any one year.

Practice questions

Scaffolded

- Q1.** About the parts of an argument. (a) Name the part that is “your answer to the question”. (b) Name the part that joins the evidence to the claim. (c) What does the check-back loop ask before you hand the argument to anyone?
- Q2.** Primary or secondary? Say which for each, and one line of reason. (a) Your class grows beans in three soils and measures the heights. (b) You read the Bureau of Meteorology’s temperature record on its website. (c) You download a museum’s published records of bird sightings.
- Q3.** About the scorecard. (a) Name the five questions. (b) Which question asks whether the source gains from one result or another? (c) A source answers weakly on three rows. What should you do?
- Q4.** About the evidence ladder. (a) Name the bottom rung and the top rung. (b) Why can the bottom rung never finish an argument? (c) What does the top rung have that no single study has?
- Q5.** About the trend figure. (a) What does the red line show? (b) One cool year is marked on the graph. In one sentence, why does it not remove the trend? (c) State the real, sourced size of Australia’s rise since 1910.
- Q6.** About source lines. (a) Name the five parts of a source line. (b) Which part answers “which version of the page did you read”? (c) Why does that part matter?

Un scaffolded

- Q7.** A poster for a fictional sunscreen, SuperShield, says: “*Blocks 100% of the sun’s rays — doctors agree!*” In small print: the company asked five of its own staff if they liked the smell. (Fictional product, for practice.) Evaluate the claim: state what was checked, say where it sits on the ladder, and give a verdict using the verdict words of this subtopic. (3 marks)
- Q8.** Someone says: “*It was cooler in 1998 than in 1997 — the warming stopped.*” Using the trend figure, explain in two or three sentences what is wrong with the argument — name the move being made. (2 marks)
- Q9.** A student wants to answer “*what subject do Year 8 students like best?*” She asks her five friends, writes their names next to their answers, and posts the results on the noticeboard. Name the three fairness rules for data about people, and say which one each part of her method breaks (VC2M8ST04). (3 marks)
- Q10.** A student’s source line reads: “*that weather website, the page about rain, looked at it a while ago.*” List the parts of the source line that are missing or unusable, using the five-part shape. (2 marks)
- Q11.** Claim: “*Plants grew taller by the window.*” Evidence: “*Five plants by the window averaged 12 cm; five plants on the shelf averaged 7 cm.*” Write the missing reasoning sentence that joins the evidence to the claim. (2 marks)
- Q12.** Two sources argue about a town’s water (both simulated, for practice). Source A: the council’s web page, dated 2 July 2026, written by the council’s named water officers, listing the restriction level and the dam measurements behind it. Source B: a video from an account with no name, claiming the council is hiding a water crisis, with no data shown. Use two rows of the scorecard to decide which source is more credible, one line per row. (3 marks)
- Q13.** A school garden trial (simulated data for practice): five plants got fertiliser and grew 13, 15, 14, 16 and 13 cm; five plants got none and grew 9, 10, 11, 9 and 10 cm. (a) Calculate the mean height for each group. (b) Write the claim, the evidence and the reasoning as one evidence-based argument. (c) State one limit of your argument. (4 marks)

Q14. In two or three sentences, give the two reasons citations exist that matter most when a reader wants to check your argument. (2 marks)

Answers

Q1. (a) The claim. (b) The reasoning. (c) Does the evidence really support the claim — and if not, fix the claim or find better evidence. **Q2.** (a) Primary — your class planned it and took the measurements. (b) Secondary — the Bureau measured and published it; you are reading their data. (c) Secondary — the museum generated and published the records. **Q3.** (a) Who is behind it? Are they expert on this topic? Are they independent? Is it up to date — and dated? Can you check it somewhere else? (b) Are they independent? (c) Keep looking — a source weak on several rows is not a reliable source. **Q4.** (a) Bottom: one person's story. Top: many studies, agreeing. (b) It is one person's story — it could be luck or memory playing tricks, so it can start a question but not finish one. (c) Different teams using different methods all finding the same result. **Q5.** (a) The overall trend through the year-by-year wiggles. (b) One cool year sits on the wiggle, not on the trend, which averages every year since 1910. (c) An average of 1.51 °C, with a range of 0.23 °C either side (Bureau of Meteorology and CSIRO, 2024). **Q6.** (a) Who made it; what it is called; when it came out; where to find it; when you retrieved it. (b) The retrieved-when part. (c) Pages change; the date tells the reader which version you were reading. **Q7.** What was checked: five staff asked about the smell — nothing about sun protection. Ladder: below the bottom rung's usefulness — not even an observation about blocking rays. Verdict: weakly supported (or: cannot be checked, since no test of the claim exists). **Q8.** One year against the year before is a wiggle, not a trend; the move is cherry-picking — choosing the convenient year and ignoring the whole record. The trend is what survives when every year is averaged, and one cool pair of years does not remove a rise built from over a hundred of them. **Q9.** Permission (she never asked the people her data is about — and friends are not the whole year level anyway); a fair spread (five friends cannot speak for every group in Year 8 — only people likely to agree with her were asked); identities safe (names posted where anyone can read them). **Q10.** Who (no organisation named), what (no title), when (no year), where (no address) and retrieved when (“a while ago” is not a date) — all five parts are missing or unusable. **Q11.** Any sentence that makes the link explicit, e.g. the plants by the window averaged 5 cm more than the plants on the shelf, and the two groups differed only in position, so the difference in height goes with the position. **Q12.** Example: *who is behind it?* — A names its officers and its data; B has no name at all. *Are they independent?* / *can you check it?* — A lists measurements any reader can check; B shows no data, so the claim cannot be checked. Source A is more credible. **Q13.** (a) Fertiliser 14.2 cm; none 9.8 cm. (b) Claim: for these plants, fertiliser goes with taller growth. Evidence: means of 14.2 cm against 9.8 cm for five plants each. Reasoning: the only planned difference was the fertiliser, and the fertilised group averaged 4.4 cm more. (c) Any one: five plants per group is small; one trial; these plants in this garden only. **Q14.** Credit — the citation says whose work it is; checking — the citation is the address a reader follows to judge the source and the number for themselves.

Fully-worked solutions

Q1. (a) **The claim** — your answer to the question. (b) **The reasoning** — the sentence that says why the evidence supports the claim. (c) **Does the evidence really support the claim?** If it does not, you fix the claim or find better evidence — the gap is never left in place.

Q2. (a) **Primary** — your class asked the question, planned the trial and took the measurements; the data is yours. (b) **Secondary** — the Bureau of Meteorology measured and published the record; you are using their data. (c) **Secondary** — the museum generated and published the sightings; you are downloading their data. The route is decided by who generated the data, not by who is holding it.

Q3. (a) **Who is behind it? Are they expert on this topic? Are they independent? Is it up to date — and dated? Can you check it somewhere else?** (b) **“Are they independent?”** — independence means no gain from one result or another. (c) **Keep looking** — a source that answers weakly on several rows is not a reliable source, and a weak row is a reason to check the claim somewhere else, not a reason to believe the opposite.

Q4. (a) **Bottom: one person's story. Top: many studies, agreeing.** (b) **Because it is one person's story — it could be luck, or memory playing tricks; it is real, but it can start a question and never finish one.** (c) **Agreement across different teams using different methods** — the same finding surviving many routes to it, which is the strongest shape evidence comes in.

Q5. (a) **The overall trend — what the record does once the year-by-year wiggles are averaged out.** (b) **One cool year sits on the wiggle, not on the trend; the trend averages every year since 1910, so a single cool year cannot remove it.** (c) **An average warming of 1.51 °C, with a range of 0.23 °C either side of that average** (Bureau of

Meteorology and CSIRO, *State of the Climate 2024*). The wiggles in the figure are illustrative; the size of the rise is the real, sourced number.

Q6. (a) Who made it; what it is called (the title); when it came out; where to find it (the full web address); when you retrieved it. (b) The retrieved-when part. (c) Because pages change; the retrieval date tells the reader which version of the page you were reading.

Q7. One mark per tool used correctly. What was checked: five of the company’s own staff were asked if they liked the smell — nothing about blocking the sun’s rays was measured (1 mark). Ladder: the claim sits on no rung at all — the only “test” was about the smell, so nothing measured touches the sun-blocking claim (1 mark).

Verdict: weakly supported — and because no test of the claim exists, “cannot be checked” is equally acceptable; “100%” and “doctors agree” with no doctors named cannot survive the scorecard (1 mark).

Q8. One year against the year before is a wiggle, not the trend; the trend is what survives when every year of the record is averaged together (1 mark). The move has a name — cherry-picking: choosing the year that suits the claim and leaving out the other hundred and fifteen; the same trick with hot years would fail the same test (1 mark).

Q9. One mark per rule correctly paired with the breach. Permission — her data is about people who were never told what it was for and never agreed to take part (1 mark). A fair spread — five friends cannot speak for Year 8; only people likely to share her views were asked, which is the cherry-pick applied to people (1 mark). Identities safe — she posted named answers, so everyone in the data can be identified (1 mark).

Q10. One mark for listing at least three missing parts, two for all five in the shape. Who — no organisation named; what — no title; when — no year; where — no address; retrieved when — “a while ago” is not a date. All five parts are missing or unusable, so the line is not a citation at all (2 marks).

Q11. Two marks for a sentence that both links and covers the comparison. Example: “The plants by the window averaged 5 cm more than the plants on the shelf, and the two groups were treated the same apart from position — so the difference in height goes with the position by the window” (2 marks). One mark for stating the 5 cm difference without the link; the reasoning must say why the evidence points at the claim, not just repeat it.

Q12. One mark per row worked correctly, one for the decision. Who is behind it? — Source A names its water officers and lists its data; Source B has no name behind it at all (1 mark). Can you check it somewhere else? / is it dated? — Source A is dated 2 July 2026 and lists measurements a reader can follow up; Source B shows no data, so its claim cannot be checked (1 mark). Source A is the more credible source about the town’s water (1 mark).

Q13. (a) Fertiliser: $(13 + 15 + 14 + 16 + 13) \div 5 = 71 \div 5 = 14.2$ cm; none: $(9 + 10 + 11 + 9 + 10) \div 5 = 49 \div 5 = 9.8$ cm (1 mark). (b) Claim: for these plants, fertiliser goes with taller growth. Evidence: five fertilised plants averaged 14.2 cm; five unfertilised plants averaged 9.8 cm. Reasoning: the only planned difference between the groups was the fertiliser, and the fertilised group averaged 4.4 cm more, so the difference in height goes with the fertiliser (2 marks). (c) Any one honest limit: five plants per group is small and the means could move with different plants; one trial only; the claim covers these plants in this garden, not plants everywhere (1 mark).

Q14. One mark per reason. Credit — the citation says the work belongs to the people who measured and checked it, so their work is not passed off as yours (1 mark). Checking — the citation is the address a reader follows to judge the source for themselves: is it credible, is the number really there, does the source say what it is quoted as saying (1 mark).

Sources

- Victorian Curriculum and Assessment Authority, *Victorian Curriculum F–10 Version 2.0, Science, Levels 7 and 8*, content description VC2S8I07 and its elaborations, 2026. vcaa.vic.edu.au — retrieved 2026-08-21.
- Bureau of Meteorology and CSIRO, *State of the Climate 2024 — Report at a glance*, 2024. <https://www.bom.gov.au/weather-and-climate/past-weather-and-climate/state-of-the-climate-2024/report-at-a-glance> — retrieved 2026-08-21.

The only real datum in this subtopic is the warming figure of 1.51 ± 0.23 °C since 1910 (Bureau of Meteorology and CSIRO, 2024, source line above); it is used solely as a claim to evaluate, and the causes of a changing climate are not taught at this band (D9). Everything else quantitative is **simulated data for practice**, labelled as such wherever it appears: the garden-trial heights in Q13 and the means built from them, the creek and bee scenarios of Example 1 and Try it 1, and the fictional products BrainBoost, ZoomFizz and SuperShield, which stand in for advertisements so that no real product is named or rated. The source-line shape taught here is the house form used in every ClassBasics book.

Elaboration contexts used: VC2S8I07.E01 (the source-credibility scorecard and every scorecard worked in the Try its, examples and questions), .E02 (the evidence ladder, the trend-and-cherry-pick treatment of the warming figure, and the claim-reach work on advertisements), .E03 first limb (the citation chain, the five-part source line and the fix-the-citation exercises).

Teacher-facing note — reserved elaborations and scope (not student text). 5B owns VC2S8I07 outright (single delivery, D6; TOC §4 row 28) and teaches the whole CD: constructing evidence-based arguments and evaluating claims, with the ethical limb delivered as citation practice. The ICIP-reserved content is as follows. The second limb of VC2S8I07.E03 (respecting cultural protocols about the sharing of particular information), VC2S8I07.E04 (the megafauna debate and early Aboriginal involvement) and VC2S8I07.E05 (commercial products founded on traditional knowledges, biopiracy and intellectual property) are reserved under the ICIP decision (D10; CONTENT_POLICY rule 1): no Aboriginal or Torres Strait Islander content is delivered in this book. The subtopic title retains the words “cultural protocols” because TOC §3 titles are fixed verbatim; the reserved content is not taught, and no figure covers it. TOC §2 D10’s proposal to deliver that content (recorded 2026-08-07, flagged “reviewer decision required”) was superseded by CONTENT_POLICY rule 1 (2026-08-15) and the DJB First Nations decisions of 2026-08-19; the newest delivered section in this tree, 10D, follows the same reservation pattern. Scope notes: the *State of the Climate 2024* datum is evaluation context only — no greenhouse mechanism and no cause of climate change is taught (D9 ceiling; Levels 9–10 content). The D11 interlock is VC2M8ST04 (ethical and fair methods for drawing conclusions about a whole group), delivered as the three fairness rules for data about people; students meeting it before the maths chapter should treat the rules as given. CONTESTED_TERMS applications: rule 1 — “credibility” is undefined in the curriculum and receives an operational definition at first use in the Word watch, on which every question here is answerable; rule 5 — the cool-year objection to the warming figure is stated in its strongest form before it is answered; rule 4 — the same evaluation tools are applied to official sources and to advertisements alike.