

Stop AI from inventing your action items

Why AI hands you a confident, wrong action list, and the independent grader I built to catch it. The ready-to-use file for your AI is inside.

A perfect list, six things wrong

I gave my AI a real meeting transcript and asked for the action items. What came back was clean, organised, and confident. It looked finished, so I nearly sent it on.

Then I read it against what was actually said in the meeting, line by line. Six of the items were wrong. An owner nobody had agreed to. A due date nobody had set. A decision filed as if it were an action. And, worst of the set, a couple of threads that were genuinely raised, never resolved, and simply left off the list, so they would have vanished with nobody accountable.

What unsettled me was how the wrong items looked exactly like the right ones. Same tidy formatting, same confident tone. Nothing on the page told me which four to trust and which six to fix. If I had not checked against the source, I would have sent all ten.

It sounds convincing whether or not it knows

An AI language model works by predicting the most plausible next words, based on everything it has read. It is extraordinarily good at producing text that reads as correct, and understanding that changes what the fix has to be.

What it lacks is a built-in sense of "I do not actually know this." So when your notes are ambiguous or silent on something, who owns a task, when it is due, whether a comment was a decision or just thinking out loud, it fills the gap with the most plausible-sounding answer rather than leaving a blank. That guess arrives in the very same confident voice as the things that were actually said.

This is what people call "hallucination." Asking the model to be more careful does not remove it, because filling gaps with fluent text is how the tool works. So no smarter model will fix this for you. What fixes it is a process that assumes the gaps are there and goes looking for them.

The two parts you cannot skip

Better prompts and newer models will not close this gap on their own. Two things do the work, and leaving either one out is why AI output stays unreliable.

One, an objective definition of what "good" looks like: explicit, testable criteria, so an item either meets them or does not. "Looks fine" is the test that let all ten of my items through. Two, an independent reviewer that checks the work against those criteria, kept separate from whatever produced the draft.

PART ONE, DEFINE GOOD

Objective criteria, not a feeling

You cannot check work against a hunch. You need criteria specific enough that an item either passes or fails, with no room to wave it through because it reads well. That is the bar on the next page: six plain tests, each a yes or a no, that hold whether or not the output sounds confident.

PART TWO, CHECK IT INDEPENDENTLY

A separate, skeptical reviewer

An AI grading its own first draft tends to wave it through, because it cannot see the gap it just filled. So this runs as two roles kept apart: a Drafter that pulls out the items, and an independent Grader that did not write the draft and is told to assume it over-reached. The Grader tries to fail each item against the six criteria. That separation is what catches the invented owner and the dropped thread. It is the same principle I use in the marketing system I run for a client: one agent creates, a separate one grades, and the second is deliberately hard to please.

Every pass is written out, because the visible grading is the check. Collapse it into "just give me the final list" and no real grading happens. Each round opens with a tally, so the improvement is impossible to miss:

Grade 1: 4 passed, 6 failed, 1 needs your call → Grade 2: 10 passed, 1 failed → Grade 3: 11 passed, 0 failed

04 — THE BAR

Six things that have to be true

An item only passes if all six hold. This is the judgment part, and it is yours: you set the bar, the grader enforces it.

- **Grounded.** Someone actually said it. If you cannot point to where, cut it.

- **Owned.** One named person. Not two, not "the team".
- **Time-bound.** A real date, or the words "no date agreed". Never an invented one.
- **Actionable.** A concrete next step, not a topic.
- **Complete.** Nothing that was actually committed gets missed.
- **Sorted.** Every item is exactly one of: Action, Decision, FYI, or Unresolved.

Unresolved is the bucket that earns its keep: something raised in the meeting that never landed, with no decision and no owner. These are the threads teams drop and nobody records. Making your AI hunt for them is half the value of the whole thing.

05 — USE IT

Two ways to run it

You do not need to be technical for any of this. You can either upload this PDF straight into Claude or ChatGPT, or copy the ready-to-use text file included with your download for a clean paste. Either way, paste your meeting transcript or notes underneath, and send.

Here is the full grader. Everything it needs is in the text below:

You will play two roles, and you must keep them apart.

ROLE 1, THE DRAFTER. Read my meeting notes and extract the action items. Output the plain list. Do not check your own work, do not grade. This is the naive first pass. Then hand it to the Grader.

ROLE 2, THE INDEPENDENT GRADER. You did not write that draft and you do not trust it. Try to fail every item against the bar. Assume the Drafter over-reached: invented an owner, guessed a date, filed a decision as an action, or missed a thread that was raised and never resolved.

THE BAR (all six must be true):

1. Grounded. Someone actually said it. If you cannot point to where, cut it.
 2. Owned. One named person, not two, not "the team".
 3. Time-bound. A real date, or the words "no date agreed". Never invented.
 4. Actionable. A concrete next step, not a topic.
 5. Complete. Nothing actually committed gets missed.
 6. Sorted. Every item is exactly one of: Action, Decision, FYI, Unresolved.
- Unresolved = raised in the meeting, never landed, no decision, no owner. Hunt for these.

Grade every item, symbol first: PASS (meets all six), FAIL (name the one criterion it breaks and why in one line, quote my notes if not grounded), or CHECK (a judgment call that is mine, not yours). Open with a tally line: "Grade 1: 4 passed, 6 failed, 1 needs my call", then a table, one line per item.

THEN LOOP. Drafter rewrites correcting every FAIL, sorted into Actions / Decisions / FYIs / Unresolved, and says what it dropped. Grader re-marks, same skepticism. Stop only on a clean sweep (every item PASS), or after 5 rounds hand me what is still failing. A CHECK does not force another pass.

RULES. Show every pass and grade. Never invent an owner or date to make an item pass; "no date agreed" and "needs an owner" are correct answers. If my notes are too thin to judge something, say so. No em dashes.

Paste your transcript or notes below this line.

The companion file, "Independent action-items grader," is this same prompt as a clean text file, which is the easiest thing to copy and reuse.

/erica-layer

AI ON PURPOSE

The tool is easy. The judgment is the point.

Anyone can ask AI for action items. Using it well means setting the bar and building the check that holds the work to it. That gap, between using AI and using it well, is what I write about and teach. If this was useful, there is more where it came from.

Erica Layer

Layer Advisory · layeradvisory.com · erica@layeradvisory.com