

Using AI Without Getting Burned

A plain-English guide to getting good answers from AI — and spotting the bad ones

Modern AI tools like ChatGPT, Claude, and Gemini are genuinely useful. They can draft, summarize, explain, translate, and brainstorm faster than anything we've had before. They are also confidently, fluently wrong on a regular basis — and they never sound wrong while doing it.

That combination is the whole problem. A tool that was *obviously* unreliable would be easy to handle; you'd double-check everything. A tool that is *usually* right and always *sounds* right is far more dangerous, because it trains you to stop checking at exactly the wrong moment.

This guide will not tell you AI is magic, and it will not tell you to avoid it. Both are lazy answers. Instead it explains, in plain terms, what's actually happening when AI "makes things up," where these tools are strong, where they quietly fail, and a handful of habits that let you get the benefit without getting burned.

No jargon. No hype. Just what you need to use these tools like someone who understands them.

1. What's really happening when AI "makes things up"

You've probably heard the word **hallucination** — when an AI states something false as if it were fact. Made-up statistics, invented quotes, court cases that never happened, a book

summary for a book that doesn't exist. It's worth understanding *why* this happens, because the reason tells you exactly when to be careful.

Here is the key idea, and it's the single most useful thing in this guide:

A language model does not look up facts. It predicts likely text.

When you ask a question, the AI isn't consulting a database of verified truth. It's generating the words that most plausibly *follow* your question, based on patterns in the enormous amount of text it was trained on. Most of the time, the most plausible-sounding answer also happens to be correct — because true statements are common in its training. But when it isn't sure, it doesn't stop and say so. It produces the most plausible-*sounding* answer anyway, in the same confident voice it uses for things it "knows."

A hallucination isn't a glitch or a malfunction. It's the same process that produces the correct answers, running over a gap in the model's knowledge. The machine has no internal signal that says "careful, I'm guessing now." To the AI, and to you, the guess reads exactly like the fact.

That leads to the one instinct worth burning into memory:

Fluency is not evidence. A polished, detailed, confident answer is not more likely to be true than a hesitant one. If anything, be *more* careful with the impressively complete answers — "impressive and complete" is precisely what a good-sounding guess looks like.

2. Three kinds of answer — and why they all sound the same

Every answer an AI gives you is really one of three things, even though they arrive in the same tone of voice:

Known. Well-established, widely documented information the model has effectively memorized. *“What is the boiling point of water at sea level?”* Reliable.

Worked out. The model reasons from things it knows to reach your answer. Useful, but the reasoning can be subtly wrong, or built on an outdated or mistaken starting point. *“Given these three figures, what’s the trend?”* Often right — verify the steps.

Guessed. Pattern-completion over a gap. It sounds like the kind of thing that would be true. *“What was the exact revenue of this small company in 2021?”* or *“Cite three studies proving X.”* This is where hallucinations live.

The trap is that the AI presents all three in the same steady, authoritative register. Your job as the user is to notice which kind of answer you’re probably looking at — because that tells you how hard to check it. The good news: you can usually guess correctly with common sense. Broad, well-known topics lean “known.” Specific numbers, names, dates, quotes, citations, and anything recent or obscure lean “guessed.”

3. Where AI is strong — and where it quietly fails

You don’t need to check everything with equal suspicion. You need to know the terrain.

AI is genuinely strong at:

- **Working with material you give it** — summarizing a document you paste in, reformatting your notes, answering questions about text that’s right in front of it.
- **Drafting and rewriting** — emails, first drafts, rephrasing something to be clearer or shorter, changing the tone.
- **Explaining well-known concepts** — how compound interest works, what a mortgage term means, the gist of a common process.
- **Brainstorming and structuring** — generating options, outlines, and starting points you’ll refine yourself.

- **Language transformation** — translation, simplification, converting a rough list into clean prose.

AI quietly fails at:

- **Specific facts it wasn't given** — exact figures, dates, prices, names, statistics. High risk of confident invention.
- **Citations and sources** — it will happily produce real-looking references that do not exist. Always verify every one.
- **Recent events** — models have a knowledge cut-off and don't know what happened after it unless they can search the web. Even then, check.
- **Arithmetic and precise logic** — it can slip on calculations that look simple. Use a calculator or a tool for anything that must be exact.
- **Anything with real consequences** — legal, medical, financial, tax, or safety decisions. Treat AI as a starting point, never the authority.
- **Niche or specialist topics** — the thinner the training material, the more it guesses.

The pattern underneath all of this: **AI is reliable when the truth is already in front of it or extremely well established, and unreliable the moment it has to reach for a specific fact from memory.** So whenever you can, put the material *in front of it* rather than asking it to recall.

4. Seven habits that keep you safe

You don't need to be technical to use AI well. You need a small set of habits.

1. Give it the material — don't make it remember. If you have the document, the figures, the email, paste them in and ask the AI to work *from that*. Grounding an answer in real text you provided is far safer than asking it to recall facts from memory.

2. Ask for its sources — then actually check them. “What are you basing that on?” is a powerful question. If it can’t point to something real and checkable, treat the answer as a guess. Never trust a citation you haven’t opened yourself.

3. Treat “confident and detailed” as a yellow light, not a green one. The more impressive an answer sounds on a specific, factual question, the more it’s worth a second look.

4. Verify anything you’ll act on. The rule of thumb: *scale your checking to the stakes*. Low stakes and easy to undo — go ahead. High stakes or hard to reverse — verify before you rely on it. Sending something on your behalf, spending money, making a commitment, deleting something: check first.

5. Let it draft; you decide. AI is a superb assistant and a poor authority. Use it to get to a good first draft or a shortlist of options fast — then apply your own judgment. The decision stays with you.

6. Ask twice, or ask differently. If something matters, ask the same question again in a fresh conversation, or phrase it another way. Answers that wobble between attempts are a sign the model is guessing.

7. Draw hard lines around high-stakes domains. For legal, medical, financial, or tax questions, use AI to understand the landscape and prepare better questions — never to replace a qualified human. It’s a briefing tool, not a professional.

5. The 30-second trust check

Before you rely on an AI answer, run it through this:

- 1. What kind of answer is this?** Well-known fact, worked-out reasoning, or a specific detail it pulled from memory?
- 2. Did I give it the material, or is it recalling?** Grounded answers are safer than remembered ones.
- 3. Are there specific facts, numbers, names, or sources?** If yes, verify them independently.
- 4. What happens if this is wrong?** The bigger the consequence, the harder you check.
- 5. Would I stake something real on this without checking?** If not, don't.

Most of the time this takes seconds, and most of the time the answer is fine. The point is to make the check a reflex — so that on the one occasion the AI is confidently wrong, you catch it.

6. How we work — the discipline behind trustworthy AI

Everything above is about using AI safely as an individual. Building AI *into a business* — where it's making or supporting real decisions — demands a level of rigor beyond good personal habits. This is the part most "AI experts" skip, and it's where we spend our time.

A few of the principles we build on:

Never let the same AI grade its own homework. In accounting you don't let the bookkeeper audit their own books — the conflict of interest is obvious. The same rule applies to AI. The system that produces a piece of work should not be the only thing that checks it. We use *independent* review — separate AI processes, and human oversight — so that mistakes get caught by something that had no stake in making them.

Deliberate, adversarial checking. We don't just ask "does this look right?" We actively try to break the output — challenge it, stress-test it, look for the failure — because problems you go looking for are problems you find before your customers do.

The right tool for the right job. Not every task needs the most powerful (and expensive) model, and not every task is safe on a cheap one. Hard reasoning gets the strongest model; routine work gets an efficient one. Matching the tool to the task is both more reliable and more economical.

Ground the system in real information. Wherever possible, we design AI systems to work from your actual data and documents rather than from the model's memory — the same "give it the material" principle from Section 4, engineered in properly.

Anti-fragile by design. We build systems so that if one part fails, the rest keeps working — rather than one bad output quietly taking everything down with it. Things will go wrong; the design decides whether that's a hiccup or a disaster.

Test before it ships, not after. The finished system is checked thoroughly before anyone relies on it — not debugged live on real customers.

None of this is exotic. It's ordinary engineering discipline, applied honestly to a technology that most people are still treating as magic. That gap — between the hype and the actual careful work — is exactly the problem we exist to solve.

7. The honest bottom line

AI is a powerful assistant and a poor oracle. Used well — with the material in front of it, its claims verified, and your own judgment firmly in charge — it will save you real time and genuinely raise the quality of your work. Used carelessly, it will hand you a confident, fluent, plausible answer that happens to be wrong, at the worst possible moment.

The whole skill is knowing the difference. If this guide gave you that, it did its job.

Who we are

We're an engineering consultancy that helps businesses use AI *well* — which sometimes means telling you where AI fits, and just as often means telling you where it doesn't. No hype, no magic, no selling you something you don't need.

The work is grounded in real engineering experience: decades in hardware and software — from chip design and network processors at Intel to years of building web applications — now focused on helping organizations adopt AI in a way that's honest, reliable, and built to last.

If you'd like a straight, jargon-free conversation about where AI could genuinely help your business — and where it couldn't — get in touch.

Rehm KI Consulting Christopher Rehm · Bavaria, Germany
christopher.rehm.63@protonmail.com

This guide is general educational information about using AI tools. It is not legal, medical, or financial advice. For decisions in those areas, consult a qualified professional.