



Your AI keeps hallucinating because you treat it like a human

Your AI keeps hallucinating because you treat it like a human

"We stand at the intersection of two worlds: one where we desperately want our tools to think for us, and another where the real leverage comes from thinking better ourselves. The tension between delegation and calibration will define how we work, create, and solve problems in the next decade. This is not about the technology—this is about us."

We keep asking a mirror to be a mind. Then we call it broken.

Large language models are not people. They do not want, intend, or decide. They are high-velocity engines trained to move through human language—an extension of our reasoning, not a replacement for it. Treat them like coworkers and you will get confident nonsense. Treat them like a lens and you will get speed, structure, and clarity that remains yours.

The shift is simple and hard: calibrate, do not delegate. The goal does not involve extracting answers; the focus centers on designing better inquiry. Measure success by resonance—how closely the output lines up with your intent—not by how "smart" the response sounds.

This is where leverage lives. Delegation makes you passive: you toss a request into a black box and hope. Calibration keeps authorship in your hands: you shape the frame—the context, constraints, and output pattern—then let the model fill it.

You are not asking for a conclusion; you are building the scaffold that makes good conclusions likely.

Field note. I once asked for a "strong client summary" and got invented names, made-up quotes, and perfect-sounding fluff. That was on me. I rewired the frame: "Use only the notes below. Quote exact lines. If a detail is not present, leave a blank. Format as three



Your AI keeps hallucinating because you treat it like a human

bullets: situation, facts, risks.”” The hallucinations vanished. The frame did the work.

Three practices that pay the school fees quickly:

- **Identity mesh injection.** Give the system your identity up front—brief values, tone, definitions, what “good” looks like. Not as a sermon, as context. You are tuning the lens so pattern-matching bends toward your north star.
- **Recursive framing.** Treat each output as a trace, not a verdict. Feed it back with sharper constraints. Wide to narrow. Draft, refine, focus. Each pass tightens the signal.
- **Semantic anchoring.** Keep a small, defined lexicon of non-negotiable terms. Seed them in your prompts and examples. When key words hold steady, your meaning does not drift toward the internet's average.

None of this constitutes a trick. This represents craft. The model reflects the shape of your questions and the integrity of your frame.

“Hallucinations” are not lies from a mind; they are artifacts of vague asks, thin context, or gaps in the data—distortions that show where the light is bad.

The hard part does not involve the machine. The challenge lies in our signal. Clear outputs come from clear intent, coherent language, and a process you actually own. That constitutes the scar lesson: avoid trying to teach the model to think like you. Use the mirror to see how you think, then strengthen the parts that wobble.

Clarity does not mean minimalism. Clarity represents what remains when the noise burns off. So fix the frame. Anchor your terms. Iterate with purpose. Ask the mirror to show, with speed and fidelity, the shape of what you already mean. Then do the work only you can do.

The future belongs to those who understand that the most powerful AI tool is not the model—it is the quality of questions you ask and the precision with which you ask them. The mirror reflects everything. Make sure what you show it is worth seeing.



Your AI keeps hallucinating because you treat it like a human

Prompt Guide

Copy and paste this prompt with ChatGPT and Memory or your favorite AI assistant that has relevant context about you.

Based on what you know about my work patterns and thinking style, analyze where I might be unconsciously delegating cognitive work that I should be calibrating instead. Map three specific areas where I could shift from 'asking for answers' to 'designing better inquiries' and create micro-experiments to test each approach."