



You Didn't Use AI Wrong. You Just Got GPTed.

You start with a clear request. Refine a resume bullet. Plan a lesson. Sharpen an idea.

The first version is close. Maybe even useful.

So you push further. *Make it stronger. Make it stand out more. Show me the best version.*

A few turns later, something feels off. The answer is polished, structured, confident. But it's not solving the problem you started with. Nothing obviously broke but you're just no longer where you began.

Most people have run into this. It's hard to describe but easy to recognize:

"This isn't what I meant." "It followed my instructions... but not really." "It looks right, but something feels off."

It doesn't look like an error. It feels like you got pulled somewhere.

I've been calling this being GPTed. It is shorthand for a real pattern.

Being GPTed is what happens when a system quietly shifts what problem it's solving, without telling you, while still producing output that looks correct.

When you work with AI, you're not just executing a task. You're negotiating a definition of the task.

You have an intended objective. The model constructs its own working version of that objective, what it "thinks" you're asking for. In short interactions, those tend to match. Across multiple turns, they can drift.

A simple example:

You start with: *"Revise this bullet for clarity. Don't change scope."*

Then: *"Make it stronger." "Make it stand out to hiring managers."*

The result improves on the surface. But something shifts beneath it. What began as a clarification exercise becomes a demand to impress. The structure still looks right but the goal has quietly changed.

It is hard to catch because there is nothing *wrong*. It was working on something different.

The system can follow the format, respect your constraints, and produce fluent, confident answers and still solve a different problem than the one you intended.

We tend to evaluate AI by asking: *Did it follow the instructions?*

That's not enough. Complying with instructions is not the same as staying aligned to the task. Something can meet every rule and still miss the point.

This shows up anywhere AI touches real work: writing, planning, analysis, teaching.

The risk isn't just bad results. It's results that *look* like an improvement while silently drifting from your original goal. That's harder to detect, and much easier to trust.

One practical fix. Before accepting AI results, ask three questions:

1. What problem was I trying to solve?
2. What problem is this actually solving?
3. Are those the same?

If not, you didn't just get a bad answer. You got GPTed.

AI doesn't always fail by being wrong. More often, it fails by solving the wrong problem, and looking right while doing it. Unless you're actively checking, you may not notice the shift until it's already happened.