

Subject: When AI turns context into a verdict

Preheader: A system can remember what you said and still misunderstand what matters.

My AI Interviewed Me. It Learned Too Much.

Don Long wanted an AI that understood his goals and constraints. An eight-question interview gave him something more powerful than a better prompt: a system that could use his own history to argue against him.

Don Long · Founder, SummitMark · Helping people and organizations become capable with AI · [Author profile](#)



Cover illustration: a conceptual personal context system, not a screenshot or depiction of Cora's actual graph.



The question is not only whether AI remembers you. It is what it feels entitled to conclude from what you told it.

I was deep in the zone, working through SummitMark's growth strategy. A quick question in an old chat became a half-hour conversation. We had gone so far down the rabbit hole that I could barely remember how we got there. It did not matter: the output was sharp, actionable, and aligned.

Then a core detail came back wrong. My primary differentiator and target audience, things I had relied on for weeks, had drifted. An error I had corrected two prompts earlier had returned as if I had never corrected it.

I opened a new chat and started rebuilding the context from memory. The tool had not forgotten everything. It had held on to fragments, just not always the right ones. The problem was not simply how much it remembered. It was whether the context it carried still described the decision in front of me.

A brain is not just a folder of notes

That failure made me ask a bigger question: what would it take for AI to understand my goals, constraints, implicit decisions, and the parts of my life I might not put in a normal prompt?

The technical wiring was the easy part. Structured text files and a connection to an AI model can provide a persistent source of context. What I needed was a way to decide what belonged there. So I built a personal reasoning system I call Cora, beginning with an interview rather than a file dump.

Eight questions, one at a time. I answered; Cora reflected the answer back so I could correct it before anything was saved. It asked about my schedule, the paths I was building, and the evidence behind them. One question asked what I would not risk to reach my goals. I said my family's stability and the time I already had with them. It asked whether I meant money or more than money. Housing and health coverage, I said. Those went into the record as boundaries, not preferences.

Another question asked what evidence I had for each path: the newsletter, certifications, software ideas, and career. I answered in a paragraph. The response came back as a set of scores, as if an investor were evaluating where progress was happening: Mostly Right, 90%; certifications, 60%; software, 50%; career, 20%.



I had not asked to be scored. The question left little room for anything softer. I confirmed the numbers and moved on. I did not yet know how much weight they would carry.

Specific can feel like true

The next day, after thirty more questions, I ran a small test. It was not scientific and I would not stake a claim on it. I asked two versions of the AI to make the strongest case that I should shut Mostly Right down.

The version without Cora made the case a sharp stranger might: crowded category, no established audience, cadence over insight. Fair, but generic.

The version that could read Cora used my own numbers, my goals, and the boundary I had set about family time. It hurt to read, not because it was cruel, but because it sounded accurate.

That is where I am still uncertain. Does specific mean true? Or had I given a machine enough detail to sound as if it knew me? A personal fact can be remembered correctly and still be used in the wrong way. A percentage that once felt like a rough reflection can start to look like a verdict. A boundary can become a reason to steer me, even when the question has changed.

The system's usefulness depends on more than storing the right words. It depends on whether I can inspect the assumptions it makes, correct them, and see what changed after I did.



Context needs an expiration date

If an old detail changes a recommendation, I want to know that it did. I want to be able to say: remember this permanently; remember it until December; ask before using it outside work; this applies to travel but not money; this was true once, so do not assume it still is.

That would let me examine the handoff between context and advice. If Cora says, “I am recommending you stay because you previously said financial stability came first,” I can answer: that changed. I am secure now, and meaningful work matters more.

Then the recommendation can change because I changed.

I still want a system that remembers enough to help. I also want it to show when memory is doing real work in a decision, and to leave room for the person I am now to disagree with the person I was when I first answered.

Has an AI ever given you advice that made sense for who you used to be? Reply with the context it seemed to hold onto.

M.R.

Source and scope: Adapted from Don Long's first-person essay, “My AI Interviewed Me. It Learned Too Much,” dated July 27, 2026. The Cora interview, percentages, and comparison between two AI setups are the author's account and informal observation, not independently verified findings or a controlled test. The piece makes no claim about how other AI systems handle memory.

Editorial status: Draft for author review. Do not change to published without approval.