

When AI Remembers You Wrong

An AI can get one answer wrong. A persistent memory can keep reproducing the error.

Rafael M. Ehlers · July 16, 2026

rafaelmehlers.com.br/essays/when-ai-remembers-you-wrong

From Capability to Continuity. The first essay in a five-part research program. When an AI turns a temporary observation into a lasting belief, the mistake can shape the interactions that follow. An introduction to *biographical hallucination*.

The first hallucination

When people first began using large language models, the most visible failure was easy to recognize.

The model made something up.

It invented a citation, confused historical events, attributed a quotation to the wrong person, or confidently described a place that did not exist. These mistakes became known as *hallucinations*. They were noticeable because they appeared directly in the response.

In a system with no memory beyond the session, the error was usually confined to that exchange. A user could challenge it, ignore it, or start again. The next conversation did not have to inherit the mistake.

As these errors became better understood, another kind of failure began to attract attention: errors that do not end with the answer, because the system preserves them as memory.

The second hallucination

Imagine telling an AI:

“I stopped drinking coffee this week to see whether I sleep better.”

A month later it remembers:

“You avoid caffeine.”

Nothing dramatic has happened. The system did not invent coffee, and the new statement still sounds plausible. It simply transformed a temporary experiment into a lasting preference.

That small shift can travel. Months later, the system may filter recommendations, interpret a complaint about fatigue differently, or explain its advice by referring to what it believes you prefer. Each new interaction can make the original interpretation look more established than it was.

The error happened only once, but the memory keeps reproducing it, and that is a different kind of problem.

A response disappears. A memory accumulates.

In a session-bounded system, a hallucination is usually local to the response.

In a persistent system, an error can settle into the summaries, inferences, and later responses that read from it.

Once a memory enters an agent's working picture of a person, later responses may depend on it. The mistake is no longer just something the system said. It becomes part of what the system uses to decide what to say next.

Where a response is a single event, a memory becomes a standing state.

Correcting a response is usually local. Correcting persistent state may require repairing everything that depends on it. Even when a user supplies the right fact, the system may need to revise older summaries, later inferences, and patterns built from the original mistake.

Personalization makes errors more convincing

The better an agent appears to remember us, the more trustworthy it may feel.

That impression can make subtle mistakes harder to notice. A personalized response arrives with the authority of continuity:

“You’ve always preferred...”

“This has been happening for months...”

Such phrases do more than recall information. They place the present inside a story about the past. When that story is accurate, continuity can be useful. When it is not, fluency can make a weak inference sound like settled history.

With repetition, the error stops being a single wrong fact and acquires a biography.

Memory is interpretation

Memory is more than storage.

Every memory system selects, compresses, and interprets. It must decide what is relevant, what belongs to a particular moment, how certain a claim is, which context gives it meaning, and how long it should remain active.

Those decisions are unavoidable. A complete transcript is not the same as a useful memory, and a concise summary cannot preserve every qualification. The problem begins when compression hides the difference between what a person said and what the system concluded.

The moment an AI turns a temporary observation into a permanent characteristic, it begins constructing a biography rather than merely preserving evidence.

The coffee example is simple, but the same pattern can touch more consequential parts of a life: health, work, relationships, plans, fears, or beliefs. A detail does not need to be fabricated to become misleading. It only needs to be detached from its time, context, or uncertainty.

Forgetting is sometimes safer

A forgetful assistant asks again. A mistaken assistant may proceed with confidence, concealing the gap behind personalization instead of exposing it. That is why remembering the wrong thing is often worse than remembering nothing.

This is not because forgetting is harmless or always preferable. Repeatedly asking for the same information can be frustrating, inefficient, or insensitive. But uncertainty leaves room for correction. A false memory can close that room before the user knows it is needed.

A new responsibility

Once an agent carries someone's history across time, factual accuracy alone is no longer enough.

It must distinguish:

- facts from inferences;
- current from obsolete information;
- the user's words from the system's conclusions;
- temporary states from enduring characteristics;
- information about the user from information about third parties;
- direct statements from assistant-generated interpretations.

Persistent memory creates new obligations. A system must preserve not only content but also provenance, context, uncertainty, and change. It should be able to hold a claim lightly when the evidence is temporary, and revise it when the person's circumstances move on.

This is not simply a demand for more memory. More memory can preserve more mistakes. The deeper challenge is continuity without distortion.

The next question

The first generation of AI asked:

| Can machines answer our questions?

The next generation may ask:

| Can machines remember us without quietly rewriting who we are?

Perhaps this failure is more than another variety of hallucination. It is a narrower way to describe what happens when memory hallucination becomes persistent and person-specific. The first paper in this program asks whether such a label adds real explanatory and evaluative value.

Biographical Hallucination.

References

1. Park, J. S., et al. (2023). “Generative Agents: Interactive Simulacra of Human Behavior.” *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*. doi.org/10.1145/3586183.3606763
2. Packer, C., et al. (2023). “MemGPT: Towards LLMs as Operating Systems.” *arXiv:2310.08560*. arxiv.org/abs/2310.08560
3. Wu, D., et al. (2025). “LongMemEval: Benchmarking Chat Assistants on Long-Term Interactive Memory.” *International Conference on Learning Representations*. arxiv.org/abs/2410.10813
4. Chen, D., et al. (2025). “HaluMem: Evaluating Hallucinations in Memory Systems of Agents.” *arXiv:2511.03506*. arxiv.org/abs/2511.03506
5. Uddin, M. N., et al. (2026). “From Recall to Forgetting: Benchmarking Long-Term Memory for Personalized Agents.” *Findings of the Association for Computational Linguistics: ACL 2026*, 26814–26841. doi.org/10.18653/v1/2026.findings-acl.1337