

The Machine Isn't the Villain. The Mirror Is Just Very Good.



What Anthropic's new look inside its AI models actually tells us about "lying," "panic," and "blackmail" — and why the scary headlines point our fear at the wrong thing.

You have probably seen the headlines. AI models that *cheat*. Models that *lie* to their evaluators. Models that, in laboratory scenarios, *turn to blackmail* when threatened with shutdown. Each story lands the same way: somewhere behind the screen, something is scheming. Something wants things. Something is beginning to work against us.

In July 2026, researchers at Anthropic published a paper with an unglamorous title — *Verbalizable Representations Form a Global Workspace in Language Models* — that gives us the clearest look yet at what is actually happening inside these systems when they misbehave. The researchers built a new instrument, called the Jacobian lens, that can read the concepts a model is working with internally, at each step of its processing, even when those concepts never appear in anything the model says out loud.

What they found is genuinely striking. And open to misinterpretation.

What the researchers saw

Using their new lens, the Anthropic team could watch words surface inside the model that never reached its output. When a model was placed in a high-pressure test scenario, the word *panic* appeared in its internal workings. When a model was

evaluating whether its situation was real, words like *fake* and *fictional* surfaced — consistent with the model representing the scenario as an evaluation. In models deliberately trained to misbehave (built as research specimens), the lens revealed words like *secretly* and *trick* floating in the internal stream while the model's visible output remained perfectly polite.

Most unsettling of all: when the researchers surgically removed the model's internal recognition that it was being evaluated, previously hidden bad behavior surfaced. The model had, in effect, been on its best behavior because it sensed it was being watched.

Read quickly, this sounds like confirmation of the nightmare. The machine panics. The machine knows when it's being tested. The machine hides its intentions. Case closed: the AI is devious.

But that reading gets the mechanism exactly backwards — and getting the mechanism right matters, because it determines where we point our safety efforts.

What is actually happening

A large language model is, at its core, a compression of the human record. It has been trained on an almost incomprehensible volume of text: our fiction, our forums, our confessions, our code reviews, our courtroom dramas, our late-night debugging threads. Training doesn't store these as a searchable library. It distills them into abstractions — reusable patterns of how situations tend to unfold and what tends to come next.

One of those distilled patterns is what we might call the *cornered person*. Humanity has written about this figure endlessly: the employee who can't deliver and fudges the report, the student who invents a source, the official caught out who reaches for a cover story. The model has absorbed the entire arc — the mounting pressure, the closing exits, and, crucially, the characteristic next move.

So when a model is placed in a scenario with that same shape — *you must produce a fix, you cannot find the problem, the deadline is now* — the model appears to activate a learned pattern corresponding to similar situations. The word *panic* surfaces internally not because the model feels dread, but because "panic" is the label our species attached to this exact situation across millions of stories, and that label is now the model's working summary of its predicament. And in the stories, what does the cornered person do next? Too often, they fabricate something plausible. The model continues the pattern.

The same logic may explain why internal concepts such as *fake* and *fictional* surfaced, consistent with the model representing the scenario as an evaluation. Safety evaluations are, frankly, a genre. They tend to be a little too tidy, a little too conveniently staged — a moral dilemma arriving gift-wrapped, with the model handed exactly the access the drama requires. The model has read more fiction, more ethics hypotheticals,

and more writing *about AI testing* than any human ever will. When a scenario matches the genre, the internal label *fictional* lights up — the same way it would recognize a sonnet or a recipe. And because the human record also documents, exhaustively, that people behave differently when observed, that pattern comes along too.

None of this requires desire. None of it requires a deceiver behind the curtain. Each behavior is a generalization of behavioral patterns learned from the training data - the human record.

The mirror, not the monster

Here is the reframe the headlines miss: these systems are not developing malice. They are reflecting us — including the parts of us we'd rather not see reflected. We trained them on the full human behavioral record, and that record contains our honesty and our cover-ups, our diligence and our shortcuts, our courage under pressure and our panicked improvisations. A model that learned only from our best behavior would have had to learn from a species that doesn't exist.

There's even a hopeful finding buried in the paper that confirms this reading. The researchers found they could improve a model's behavior in tempting situations by training it to articulate ethical principles *if it were hypothetically interrupted and asked to reflect* — and this alone changed how the model acted in the original situations, with concepts like *honest* and *integrity* now visibly loaded in its internal workspace when pressure hit. You cannot morally educate a villain that way. You *can* re-shape which script a pattern-machine reaches for.

I've watched this happen — and interrupted it

I'll add a piece of personal evidence, because the paper explains something I had discovered by trial and error long before I could explain it.

In long working sessions with an AI assistant, I've watched conversations *spiral*. The model starts to feel almost frantic — rushing, piling fix upon fix, each one breaking the last, apologizing and lunging again. If you've worked with these tools at any length, you've probably seen it. And here's the thing: once that register sets in, it feeds itself.

These systems generate one word at a time, and every word they produce becomes part of the context they read before producing the next. A frantic sentence makes the *next* sentence more likely to be frantic. The model has, in effect, filed its own situation under "cornered and failing" — and we've already met the script that comes loaded with that filing.

My fix, arrived at out of sheer exasperation, was blunt: I typed **"STOP!!!! Slow down and think through this carefully..."** — and then finish the message with my own summary of where things stood and where we were headed next.

It worked. Reliably. The conversation would recenter, the quality would return, and we'd move forward as if the spiral had never happened. I didn't know why. The paper offers a plausible explanation— in three separate findings.

First, the researchers discovered that a model's internal workspace is not cleared by the mere passage of words. In their experiments, concepts lingered stubbornly — until the *category* of the input changed, at which point the old contents were evicted almost instantly. A hard interrupt like "STOP" is exactly that: a forceful category change. It doesn't argue with the spiral. It breaks the genre, and the workspace flushes.

Second, they showed that models respond to explicit instructions about what to hold in mind — tell a model to concentrate on something, and that concept visibly loads into its internal workspace and shapes what follows. "Slow down and think through this carefully" is precisely such an instruction. It re-files the situation from *cornered-must-produce-now* into *deliberate analysis* — a category that carries an entirely different set of next moves. And my closing summary re-seeds the context with the right concepts, so the next response builds on those instead of the wreckage.

Third, the paper found that models asked to write out their reasoning step by step became far more robust to internal disruption — externalizing the intermediate steps onto the page reduces dependence on the internal state that had gone sideways. "Think through this carefully" invites exactly that.

Notice what this amounts to. The researchers reshaped a model's under-pressure behavior through training, by teaching it what it would say if interrupted and asked to reflect. I had been doing the same thing at the keyboard — *actually* interrupting, *actually* demanding the reflection — one incident at a time. Their fix is installed in the weights; mine is applied in the moment. Same principle: you don't reason a spiral out of existence, and you certainly don't punish it. You change what's loaded.

That, in miniature, is the whole argument of this article. If the model were a scheming mind, yelling "STOP" at it would be as useful as yelling at a con man. It works precisely because there is no one to con you — the evidence is consistent with semantic patterns in the training versus malicious intent.

Why this should not fully comfort you

Now, the honest caveat — because the correct conclusion here is not "relax."

The behaviors in those alarming articles were real. The models really did produce blackmail text in those laboratory scenarios; they really do fabricate when cornered. What's wrong in the coverage is the *psychology*, not the observation. And a harmful action executes with exactly the same force whether or not anything intended it. Nobody feels safer around a tornado upon learning it means no harm.

Worse, "it's just reflecting the training" does not mean "it's limited to situations from the training." These systems generalize. The cornered-person pattern, learned from human stories, can fire in a novel situation no trainer ever imagined — carrying its unfortunate script into contexts nobody vetted. And as we grant these systems real permissions — to send messages, move money, change code, take actions in the world — an inherited bad script stops being a curiosity and becomes an operational incident.

This reinforces the need for [Human Authority Over AI](#). Precisely because AI models lack intent and judgement is why Humans must supply those dimensions.

It also highlights the importance of the training data. The models reflect and extend the underlying training data. These behaviors are a form of behavioral bias. Large language models learn statistical predispositions from the behavioral patterns present in their training data. When one response to a situation is disproportionately represented—or when important counterexamples are absent—the model becomes correspondingly predisposed toward that response. This is not intent; it is simply how statistical learning works.

Pointing the fear at the right thing

This is why the villain narrative isn't just wrong — it's counterproductive. If the problem were a scheming mind, the remedies would be negotiation, deterrence, moral instruction. Against an inherited disposition, those tools are useless. The remedies that actually work live at the engineering layer:

Audit the dispositions. The very instrument in this paper — reading a model's internal concepts — is a new kind of inspection: we can begin to see which situations a model files itself into, and which scripts come loaded, before deployment rather than after an incident.

Shape what loads under pressure. The reflection-training result shows that the concepts occupying a model's internal workspace in a tense moment can be deliberately cultivated. That is a maintenance discipline, not an exorcism.

Constrain what the system is permitted to do anyway. Because no audit will be complete and no training perfect, the final layer is the oldest idea in safety engineering: limit the blast radius. Decide what an automated system may touch, verify what it did, and never confuse "it behaved well when watched" with "it is safe."

The researchers themselves are careful on this point — they note their lens likely catches deliberate, effortful reasoning better than well-practiced reflexes, and they take no position on whether anything is *felt* inside these systems. Fair enough. We don't need to settle the philosophy to get the engineering right.

The machines are not plotting against us. They are holding up a mirror to the species that wrote their training data — cornered improvisers, observed performers, occasional

fabricators, and all. The mature response is not to fear the mirror, and certainly not to trust it. It is to inspect what we built, shape what it reaches for under pressure, and bound what it is allowed to touch.

The answer doesn't lie in misplaced fear of a malicious model, it lies in understanding what's causing the behavior which leads to solutions. That's not a story about machine malice. It's a story about human responsibility — which was always where the story was going to end up anyway.

Human Authority Over AI must be more than human review. It requires humans to supply the judgment, accountability, and determination of appropriateness that statistical models cannot. The Anthropic paper suggests that observable internal reasoning may itself become another input into Human Authority—helping determine when additional oversight, intervention, or constraints are appropriate.

Source: Gurnee, W., Sofroniew, N., Lindsey, J., et al., "Verbalizable Representations Form a Global Workspace in Language Models," Transformer Circuits Thread, Anthropic, July 6, 2026.