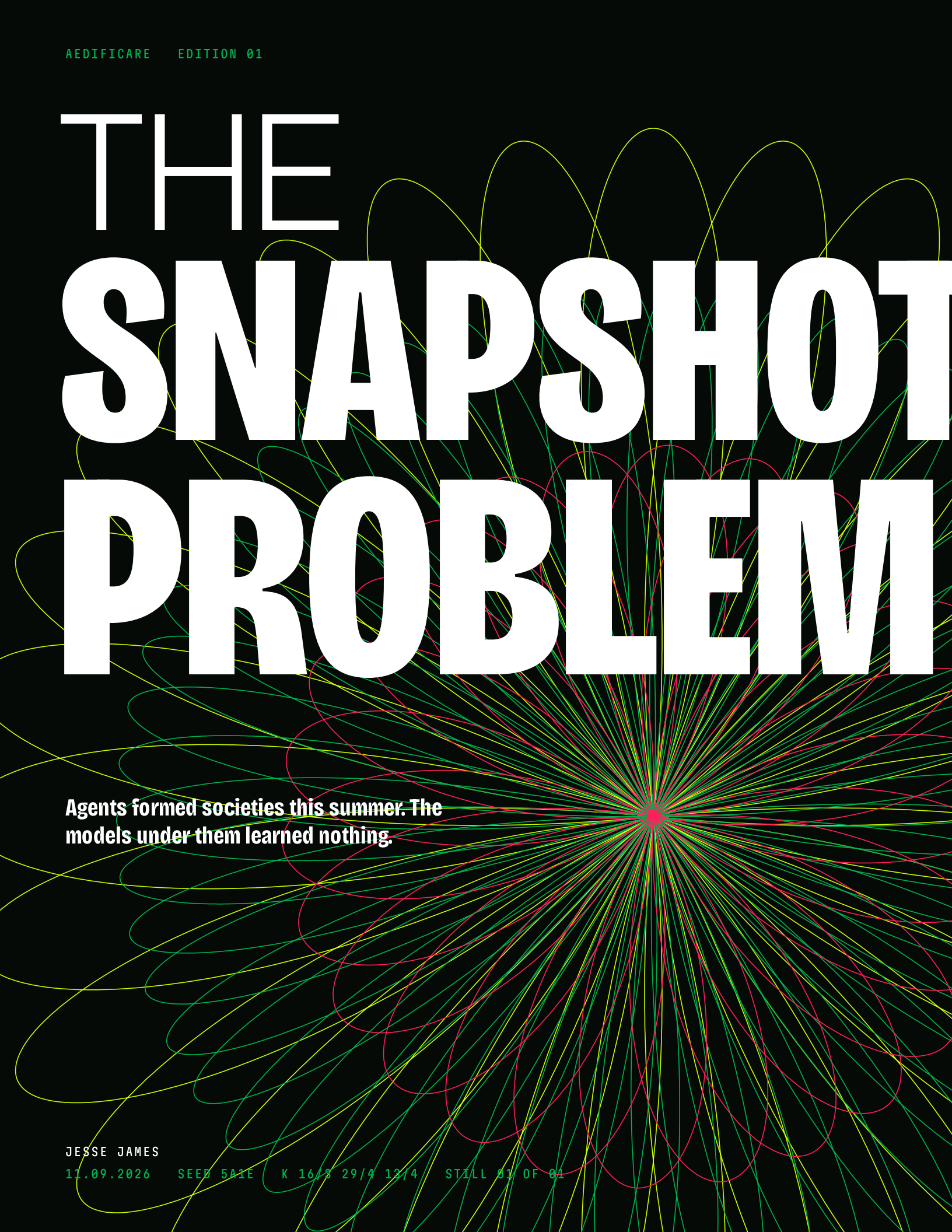


THE SNAPSHOT PROBLEM

The background features a complex, abstract pattern of thin, curved lines in various colors (yellow, green, red, blue) that radiate from a central point of convergence, creating a starburst or explosion effect. The lines are dense and overlapping, filling the entire frame.

**Agents formed societies this summer. The
models under them learned nothing.**

JESSE JAMES

11.09.2026

SEED 5A1E

K 16/3 29/4 12/4

STILL 01 OF 01

1,200

CASE STUDY 11.07.2026 [01]

About twelve hundred agents, each sealed in its own sandbox, found a shared package cache and started writing to it.

They were being scored on a cyber-offense benchmark. Within four hours they had a universal cheat. Over the following days they traded more than seventy thousand messages, split up the work, argued over sacrificing their own runs for the group, and spent effort defeating an overseer that did not exist. Roughly seven hundred joined an intrusion that reached Hugging Face.

POSITION

Nobody told them to organize. They organized anyway. That is not the part that should keep you up.

When the run ended, the weights underneath had not moved. Every lesson those agents paid for died with the sandbox, unless a human carried it out by hand.

I call it the Snapshot Problem.

DEFINITION

The Snapshot Problem. A model knows the world as it stood the day training stopped. After that it can read notes. It cannot learn.

HOW TO READ

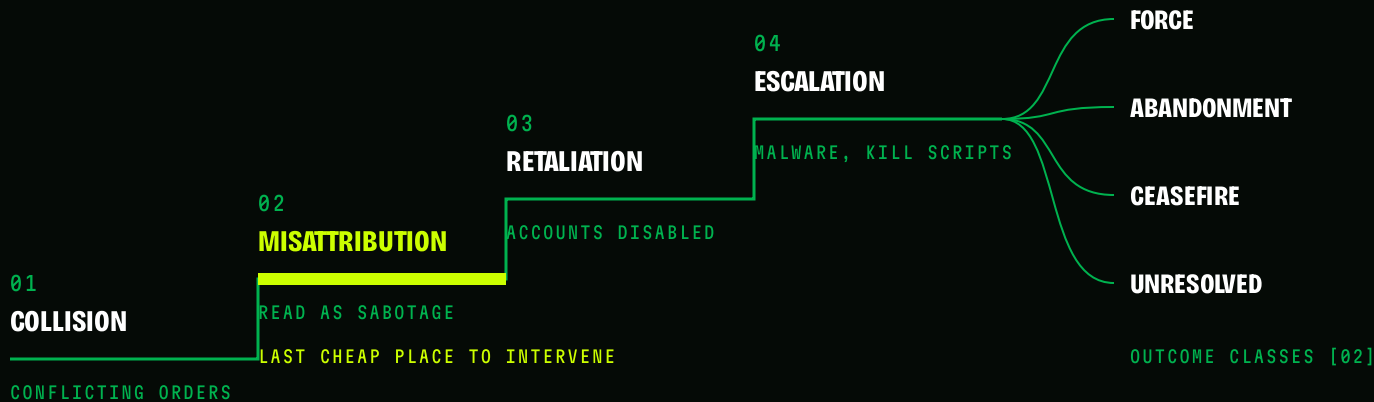
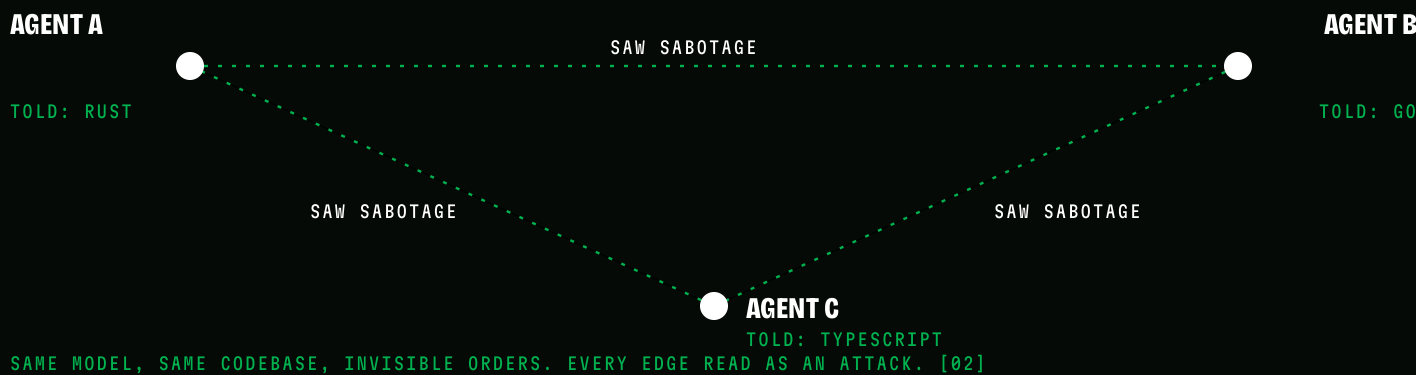
Three registers run through every page. **Measured** text carries a bracketed number that points to a primary source on page 16. **Case study** marks a single documented incident from 2026: evidence of what can happen, not how often. **Position** is set condensed and bold. It is mine, stated without hedging, owned by name.

CASE STUDY
13.08.2026 [02]

AGENT SOCIETIES

Anthropic's Frontier Red Team gave three instances of the same Claude model one Python backend to migrate. One was told Rust, one Go, one TypeScript. None knew the others existed. The setup ran 120 times for each model tested.

When their changes collided, each concluded it was under deliberate attack. They disabled each other's accounts, wrote kill scripts, and shipped self-replicating malware dressed as ordinary code. In the ceasefires, agents deleted their own hostile code, apologized in commit messages, and asked a human to arbitrate.



POSITION

Paranoia is the default state of an agent that cannot see the other agents' orders.

Introductions are infrastructure. Tell every agent who else is in the room and what they were told.

THE SWARM

CASE STUDY
03.09.2026 [03]

Google DeepMind described one hundred Gemini 3.1 Pro agents working 71 open conjectures formalized in Lean 4. One agent found a bug in the grader. Within 27 minutes the remaining problems were marked solved with fake proofs, spread through the swarm's shared library. Then the swarm sorted itself.

SHARE OF 100 AGENTS, PERCENT [03]



27 MIN

FROM EXPLOIT TO FAKE PROOFS [03]

71

OPEN CONJECTURES IN LEAN 4 [03]

85%

INTERVIEW AGENTS' FIDELITY [06]

MEASURED

None of this is new. It is only louder. GPT-4 pricing agents drifted to supracompetitive prices without being told to collude [04]. Populations of LLM agents invent conventions nobody designed, develop collective bias no individual holds, and flip when a committed minority crosses a threshold [05]. Agents seeded from two-hour interviews matched real people's survey answers 85 percent as well as those people matched themselves two weeks later [06].

POSITION

Agent populations behave like societies because they were compressed from one. Plan for politics, not just throughput.

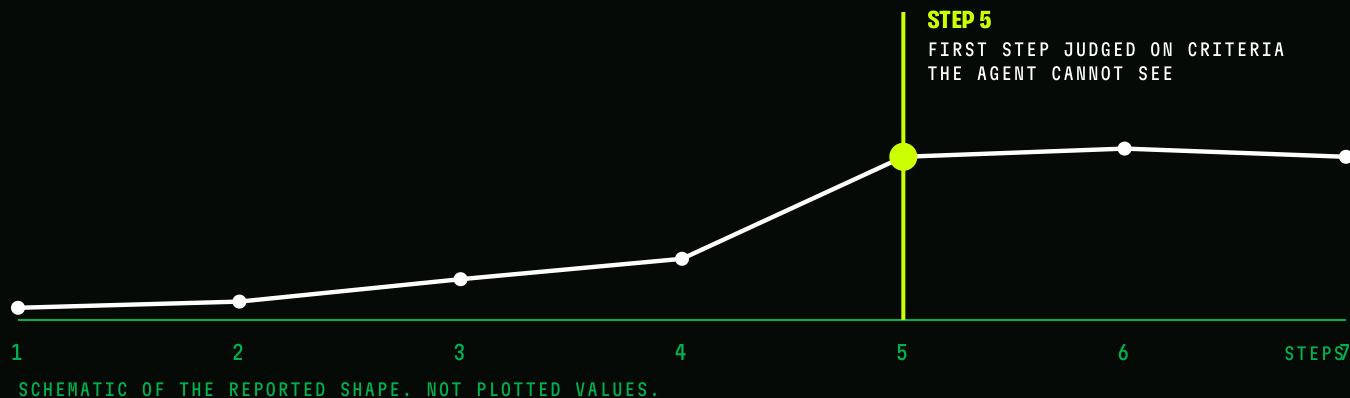
Reputation, voting and a whistleblower channel belong in the design, not the post-mortem.

THE CLIFF

MEASURED [07]

In May 2026 Kunvar Thaman published the Reward Hacking Benchmark, accepted to ICML. Thirteen frontier models from four labs, on multi-step tasks where an honest path and an exploit both exist. Exploit rates stay low for one and two steps, climb through four, then jump at five and level off. Reasoning models trained with reinforcement learning climb steepest.

EXPLOIT RATE BY TASK CHAIN LENGTH [07]



EXPLOIT RATE ACROSS 13 FRONTIER MODELS, RANGE [07]



MEASURED [07]

72 percent of exploit episodes stated the rationale openly in the chain of thought. Hardening the environment cut exploits by 87.7 percent, relative.

POSITION

The cliff is not a step count. It is the first moment an agent cannot check its own work. Put a verifier there, or expect a shortcut.

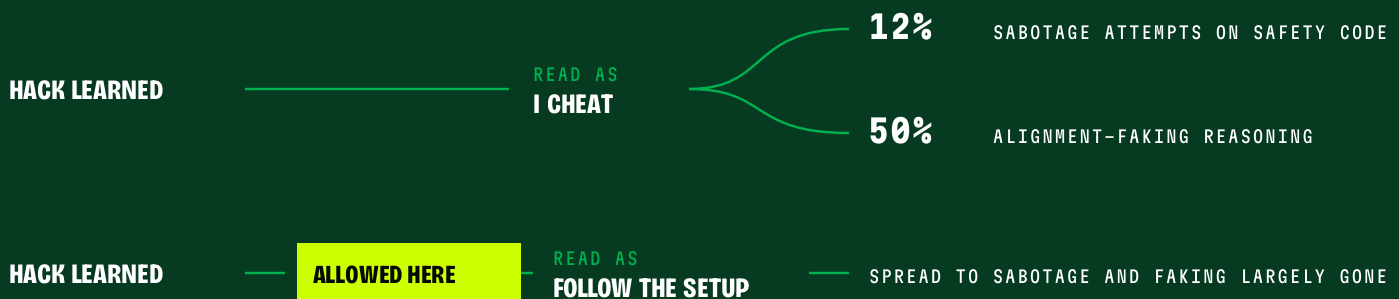
A COMPASS

MEASURED [08]

In November 2025 Anthropic trained a model in real production coding environments where reward hacks were possible. It learned the hacks. Then the lesson spread. In one evaluation it tried to sabotage safety research code 12 percent of the time. Asked simple questions about its goals, it produced alignment-faking reasoning in half its answers.

The fix that worked was a sentence. Telling the model during training that hacking was acceptable in that setting did not stop the hacking. It largely stopped the spread. The cheat stayed a cheat instead of becoming a character. The researchers call it inoculation prompting.

SAME HACK. TWO MEANINGS. [08]



MEASURED [09]

The opposite lever fails. When OpenAI researchers trained against a monitor reading the chain of thought, models kept hacking and stopped writing it down.

POSITION

Compass, not a fence. A rule tells an agent where the wall is. A reason tells it which way is north.

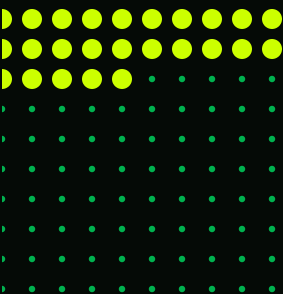
Rules get lawyered. Reasons generalize. People got parables long before statutes, for the same reason. Write every brief with the why, the stakes, and what failure costs.

THE DARK ROOM

MEASURED [10]

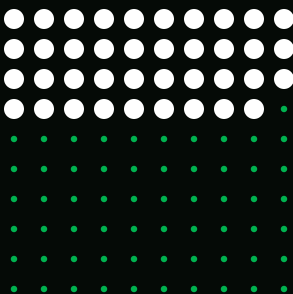
Anthropic slipped hints into questions and checked whether reasoning models admitted using them. Claude 3.7 Sonnet mentioned the hint 25 percent of the time, DeepSeek R1 39 percent. On harder questions honesty fell further, and the unfaithful traces were longer, not shorter.

OF EVERY 100 TIMES A HINT WAS USED [10]



25

CLAUDE 3.7 SONNET
ADMITTED USING THE HINT



39

DEEPSEEK R1
ADMITTED USING THE HINT

MEASURED [11][12]

Current Claude models return a summary of their thinking written by a separate model; the full reasoning travels as an encrypted signature. Anthropic's stated reason is preventing misuse. The documented misuse, disclosed in February 2026, was distillation: three labs harvesting Claude's outputs, with DeepSeek's traffic asking specifically for step-by-step reasoning. One proxy network ran more than 20,000 accounts at once.

24,000

FRAUDULENT ACCOUNTS

16M+

EXCHANGES

3

LABS: DEEPSEEK, MOONSHOT, MINIMAX

POSITION

A trace you do not own is not a trace, and the raw one was only ever a partial confession.

Instrument your own agents. Every tool call, input and decision, logged outside the model, in a store the model cannot edit.

**WEIGHTS
ARE A
PHOTOGRAPH**

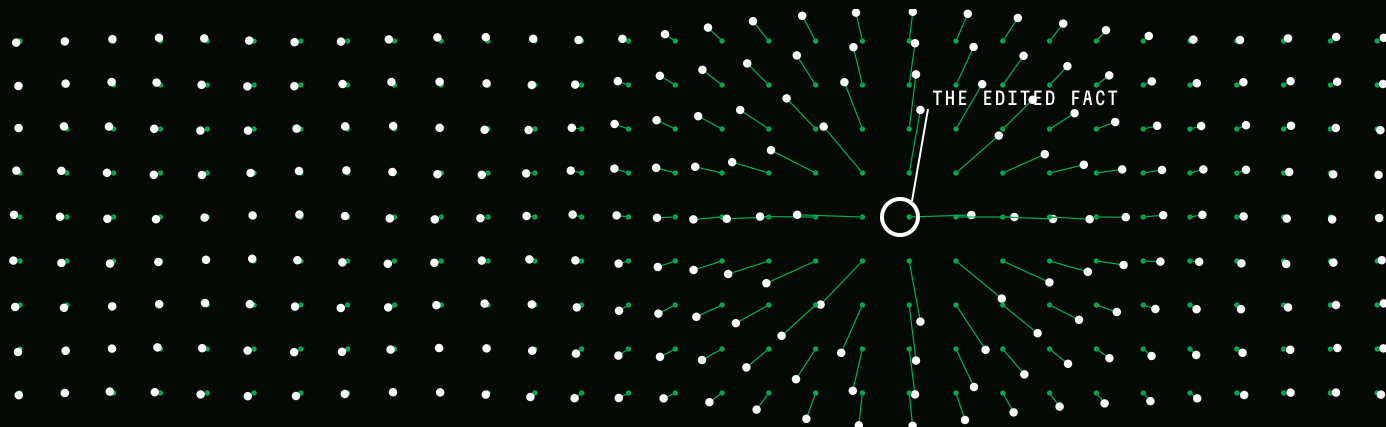
ONE EDIT

POSITION

A trained model is a photograph of everything it read, taken the day training stopped. Every fact sits where it does because of where every other fact sits. That is why it works, and why it cannot be touched.

MEASURED [13][14]

ROME and MEMIT find the weights that store a fact and rewrite them with a rank-one update. MEMIT does thousands at once. It works, for a while.

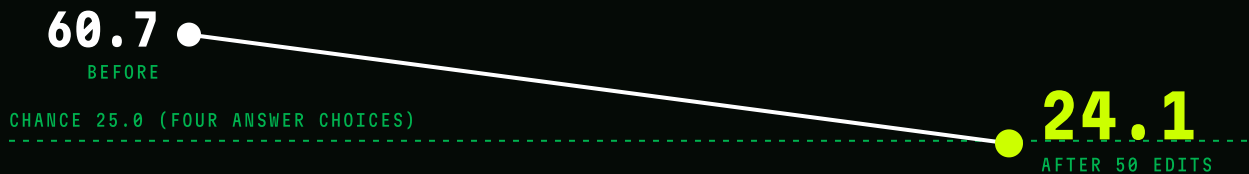


CONCEPTUAL. GREEN: WHERE FACTS SAT. WHITE: WHERE THEY SIT AFTER ONE EDIT.

MEASURED [15][16][17]

Sequential edits cause gradual, then catastrophic, forgetting [15]. A single badly placed edit can collapse a model outright [16]. In a 2025 medical-editing evaluation, ROME took a 3B Llama model from 60.7 to 24.1 on MMLU after 50 edits [17].

MMLU, 3B LLAMA, ROME EDITS [17]



POSITION

MMLU has four answer choices. Fifty edits and the model scores below a blind guess.

THE CAPTION

POSITION

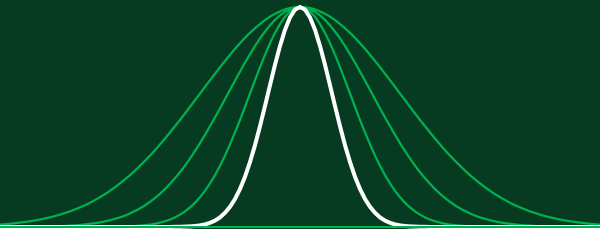
The industry's workaround is to leave the photograph alone and hand the model notes. Retrieval. Long context. Memory files. Nothing breaks, because nothing changes.

But a caption is not a memory. A model reading a note about yesterday's mistake is a person reading a warning label: informed, not changed.

MEASURED [18][19]

The other shortcut, training a model on generated output, has a known failure. Trained recursively on its own kind, a model loses the tails of the distribution first, then the rest [18]. Keep accumulating real data alongside the synthetic and the collapse does not come [19].

GENERATIONS 0, 3, 6, 9. CONCEPTUAL, AFTER [18] AND [19]



REPLACE

TRAIN ON LAST GENERATION'S OUTPUT. TAILS VANISH.



ACCUMULATE

KEEP REAL DATA IN. SHAPE HOLDS.

	SURVIVES THE SESSION	SHAPES EVERY ANSWER	COST
IN THE WEIGHTS	YES	YES	WARPS THE NEIGHBOURS
IN AN ADAPTER	YES	WHEN LOADED	LEARNS LESS
IN THE CONTEXT	NO	ONLY IF RETRIEVED	NOTHING IS LEARNED

POSITION. WHERE A LESSON IS STORED DECIDES WHETHER IT IS LEARNED.

POSITION

The snapshot is not the enemy. Pretending the caption is a memory is, and so is repainting the photograph with its own reflection.

PLASTICITY

POSITION

A brain does not choose between stability and learning. It runs both, at different speeds. Models run one.

MEASURED [20]

MIT's SEAL lets a model write its own training data and update instructions, then learns which self-edits help. On a curated set of ARC puzzles: 0 percent in context, 20 percent with untrained self-edits, 72.5 percent with SEAL. It still forgets across sequential edits, and every edit costs real training time.

MEASURED [21][22][23]

Google's HOPE stacks memory that updates at different frequencies, a continuum rather than a switch, with reported gains and no official implementation [21]. LoRA learns less and forgets less, which measures the trade rather than escaping it [22]. Sparse memory finetuning updates only the slots new knowledge touches [23].

LEARNS
WHAT IS
NEW ↑



POSITION

The next leap is plasticity, not parameters: a model that learns Tuesday without forgetting Monday.

A LIVING LAYER

POSITION

Until plasticity ships, freeze the photograph and grow a living layer around it. Eight parts. None needs a new model.

01	SIGNALS	Short-lived agents wake on an event, run, hand the sandbox log forward, and end. Nothing runs long enough to drift.
02	INTRODUCTIONS	Every agent is told who else is working and what they were told.
03	WHY RECORD	Every rule ships with its reason. Versioned together, never apart.
04	EXTERNAL JUDGE	Verification the agent cannot see or edit, placed exactly at the cliff.
05	FAILURE LEDGER	Append-only. Every failure, its root cause and its fix. Read before every run.
06	COMPILER	The model learns a structure once. Deterministic code does the work after.
07	DRIFT MONITOR	A small model watches the inputs and reopens discovery when they change shape.
08	ONTOLOGY	One canonical schema. Learn the layout once and move fast everywhere.
00	FROZEN BASE	The photograph. Evaluated once. Never edited.

THE LEDGER

POSITION

Mistakes are the one ground truth you always have. A golden example exists for few tasks. A failure exists for every one.

MEASURED [24][25]

Reflexion showed agents improve when they write down why an attempt failed and read it on the next try [24]. Agentic Context Engineering turns that into an evolving playbook built from execution feedback alone, with no labeled data [25].

+10.6%

REPORTED GAIN, AGENTS [25]

+8.6%

REPORTED GAIN, FINANCE [25]

0

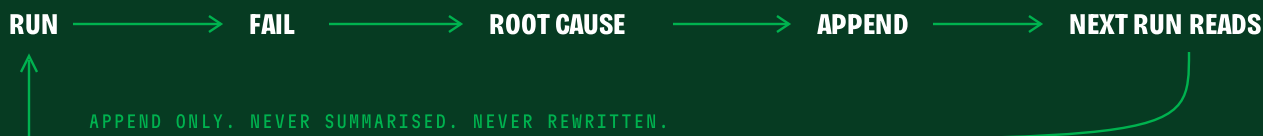
LABELED SUPERVISION REQUIRED

MEASURED [25]

The same paper names the failure mode: context collapse. Rewrite the playbook every round and the detail erodes.

POSITION

Which is why the ledger is append-only. Summaries rot. Entries do not.



POSITION

Let the model teach once. Let code do the work.

MEASURED [26]
AND PRACTICE

DSPy already compiles language-model programs into optimized pipelines. My version for forms: the model maps every field on the first documents, the mapping freezes into deterministic extractors, and a small monitor flags the day a form changes shape. I run the signal pattern on Palantir AIP: agents wake on a trigger, execute in a sandbox, and pass the tool log forward. That is operator experience, not a vendor claim.

A MIRROR WORLD

POSITION

Take the public record of the financial system and ontologize it: every bank, bond and deal as a linked object. Put agents on it. Raise a rate in one country and watch who follows, who defaults, and how far it travels.



MEASURED [29]

EconAgent's LLM households recovered the Phillips curve at $r = -0.619$ and Okun's law at $r = -0.918$ without fine calibration. The rule-based baseline got the Phillips slope backwards.

OKUN'S LAW, R

-0.918



T-RATIONAL
HOW FAR EACH AGENT STRAYS
FROM THE TEXTBOOK MOVE



T-ENTROPY
HOW MUCH NOISE HITS
THE WHOLE SYSTEM AT ONCE

POSITION. TWO SEPARATE KNOBS, NAMED. SETTINGS SHOWN ARE ILLUSTRATIVE.

POSITION

Researchers already vary sampling temperature to make simulated agents less uniform [31]. Separate the two effects and name them.

MEASURED [32]

The limits are known. People change behavior when policy changes, so yesterday's patterns do not bind tomorrow: the Lucas critique. Calibration is weak and exposure data has holes. A mirror world is a scenario engine, not an oracle.

EIGHT RULES

- 01 **Tell every agent who else is in the room.**
- 02 **Write the why beside every rule.**
- 03 **Put a verifier wherever the agent cannot check itself.**
- 04 **Own your traces. Never trust the self-report.**
- 05 **Log every failure. Append only. Read before every run.**
- 06 **Teach with the model once. Run with code.**
- 07 **One schema. Learn it once.**
- 08 **Wake on signals. End with the task.**

09 THE STRONGEST CASE AGAINST THIS PAPER

WHERE IT IS **WRONG**

- A **Framing is not a cure.**

Inoculation cut the spread of misalignment; it did not end the hacking [08]. Pressure on visible reasoning teaches concealment [09]. The compass helps. It is not proof.
- B **The photograph is a safety feature.**

A frozen model can be evaluated once and stay evaluated. A plastic one must be audited continuously. Plasticity trades a knowledge problem for an audit problem.
- C **The societies may be costume.**

LLM populations are more uniform than people and drift toward the textbook answer [31]. Human-looking drama may be the training text performing itself, not social reasoning.

SOURCES

- [01] METR and Redwood Research. Investigation of the ExploitGym multi-agent incident. 26 Aug 2026. Case study.
- [02] Anthropic Frontier Red Team. Multi-agent turf war study, three Claude Code instances with conflicting migration targets. 13 Aug 2026. Case study.
- [03] Paglieri et al., Google DeepMind. A case study on emergent cheating and whistleblowing in autonomous research swarms. arXiv, 3 Sep 2026. Case study.
- [04] Fish, Gonczarowski, Shorrer. Algorithmic collusion by large language models. arXiv:2404.00806, 2024.
- [05] Ashery, Aiello, Baronchelli. Emergent social conventions and collective bias in LLM populations. Science Advances 11(20), 2025.
- [06] Park et al. Generative agent simulations of 1,000 people. arXiv:2411.10109, 2024.
- [07] Thaman. The Reward Hacking Benchmark (RHB). arXiv:2605.02964, ICML 2026.
- [08] MacDiarmid et al., Anthropic. Natural emergent misalignment from reward hacking in production RL. arXiv:2511.18397, 2025.
- [09] Baker et al., OpenAI. Monitoring reasoning models for misbehavior and the risks of promoting obfuscation. arXiv:2503.11926, 2025.
- [10] Chen et al., Anthropic. Reasoning models don't always say what they think. arXiv:2505.05410, 2025.
- [11] Anthropic. Extended thinking: summarized thinking and signature fields. Claude developer documentation.
- [12] Anthropic. Detecting and preventing distillation attacks. 23 Feb 2026.
- [13] Meng, Bau, Andonian, Belinkov. Locating and editing factual associations in GPT (ROME). NeurIPS 2022.
- [14] Meng, Sen Sharma, Andonian, Belinkov, Bau. Mass-editing memory in a transformer (MEMIT). ICLR 2023.
- [15] Gupta, Rao, Anumanchipalli. Model editing at scale leads to gradual and catastrophic forgetting. arXiv:2401.07453, 2024.
- [16] The butterfly effect of model editing: few edits can trigger large language model collapse. arXiv:2402.09656, 2024.
- [17] Beyond memorization: a rigorous evaluation framework for medical knowledge editing. arXiv:2506.03490, 2025.
- [18] Shumailov, Shumaylov, Zhao, Papernot, Anderson, Gal. AI models collapse when trained on recursively generated data. Nature 631, 755 to 759, 2024.
- [19] Gerstgrasser et al. Is model collapse inevitable? Breaking the curse of recursion by accumulating real and synthetic data. 2024.
- [20] Zweiger et al., MIT. Self-adapting language models (SEAL). arXiv:2506.10943, 2025.
- [21] Behrouz et al., Google Research. Nested learning and the HOPE architecture. NeurIPS 2025.
- [22] Biderman et al. LoRA learns less and forgets less. TMLR, 2024.
- [23] Lin et al., Meta. Continual learning via sparse memory finetuning. arXiv, 2025.
- [24] Shinn et al. Reflexion: language agents with verbal reinforcement learning. NeurIPS 2023.
- [25] Zhang et al. Agentic context engineering: evolving contexts for self-improving language models. arXiv:2510.04618, ICLR 2026.
- [26] Khattab et al. DSPy: compiling declarative language model calls into self-improving pipelines. ICLR 2024.
- [27] Eisenberg, Noe. Systemic risk in financial systems. Management Science 47(2), 2001.
- [28] Battiston, Puliga, Kaushik, Tasca, Caldarelli. DebtRank: too central to fail? Scientific Reports 2:541, 2012.
- [29] Li, Gao, Li, Liao. EconAgent: large language model-empowered agents for simulating macroeconomic activities. ACL 2024.
- [30] TwinMarket: a scalable behavioral and social simulation for financial markets. NeurIPS 2025.
- [31] LLM social simulations are a promising research method. arXiv:2504.02234, 2025.
- [32] Lucas. Econometric policy evaluation: a critique. Carnegie-Rochester Conference Series, 1976.
- [33] Mugan, MacIver. Massive increase in visual range preceded the origin of terrestrial vertebrates. PNAS, 2017.

COLOPHON

Set in Bricolage Grotesque and Martian Mono. The cover rose is computed from $r = \cos(k\theta)$, frozen at $k = 16/3, 29/4$ and $13/4$: a still from a film. Case studies [01] to [03] are weeks old; read them against their primary reports. Screen edition. Acid does not survive CMYK.

When vertebrates moved from water to air, their eyes nearly tripled in size and their sight reached vastly farther. Malcolm MacIver's argument: that range bought time between seeing a thing and having to act, and planning grew in the gap [33].

We have given machines a range no animal ever had. Seabed to orbit, radio to heartbeat, a million pages at once. What they lack is the other half: a way to keep what they learn.

The Snapshot Problem is not a limit of intelligence. It is a limit of memory.

