

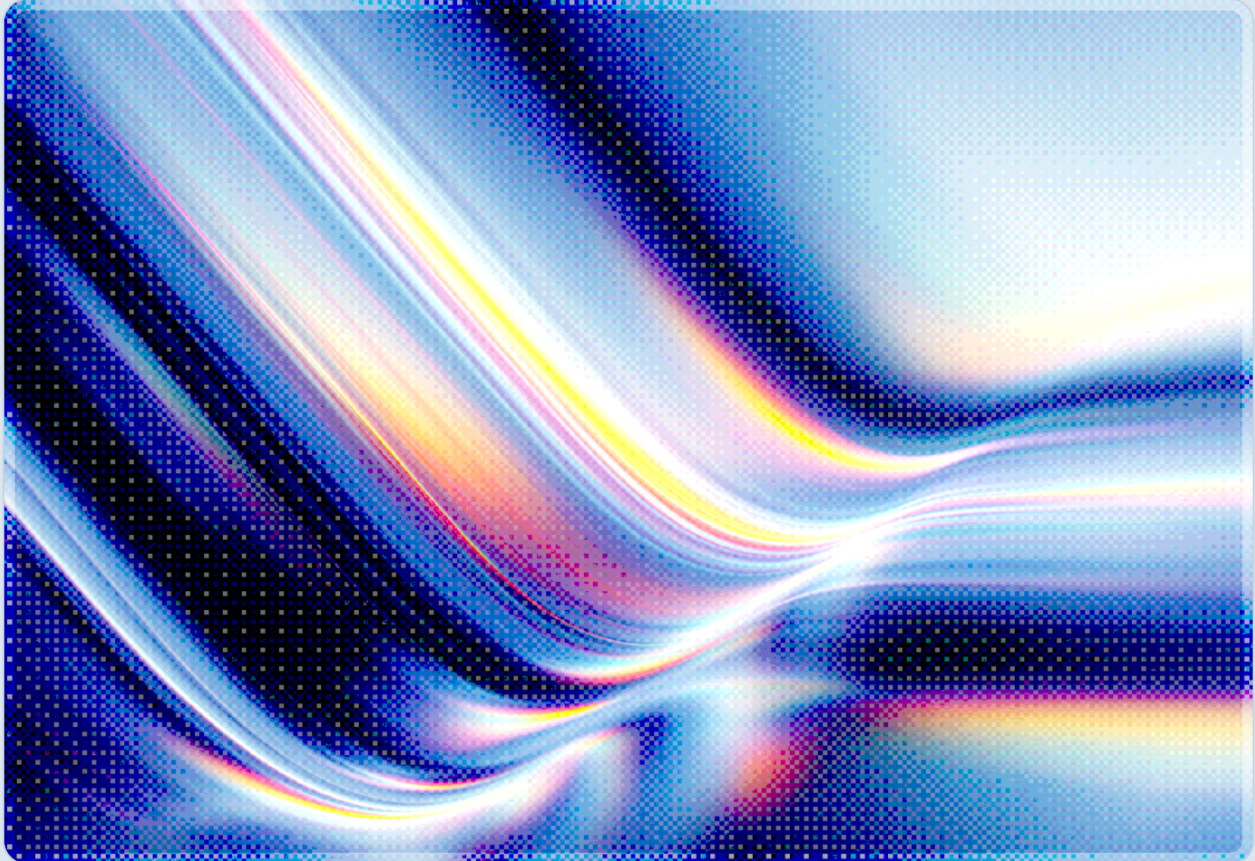
The Chain-of-Thought Encryption Illusion

Recovering 'Encrypted' Reasoning Traces Across OpenAI, Anthropic, and Google

2026-08-20



AI-assisted Rapid Research



CSAI

cloud
CSA security
alliance®

© 2026 Cloud Security Alliance. Some rights reserved.

You may download, store, display, view, print, redistribute, and link to this document in its original, unmodified form, provided that attribution to the Cloud Security Alliance is maintained and all trademark and copyright notices remain intact.

This document may not be modified or altered. You may quote portions of the document as permitted by the Fair Use provisions of the United States Copyright Act, provided that attribution is given to the Cloud Security Alliance.

This document may be shared on professional and social media platforms in its original form with attribution.

This document was generated with AI assistance and has not undergone official CSA review and approval processes.

Key Takeaways

- A disclosure, *Stealing Reasoning Traces from Proprietary LLM APIs* (arXiv, August 10, 2026), shows that Anthropic, OpenAI, and Google all return "encrypted" chain-of-thought reasoning to API clients using encryption schemes that are interchangeable across sessions, user accounts, and sibling models within the same provider's model family [1].
- Attackers can extract an encrypted reasoning blob generated by a capable model and replay it, unmodified, into a weaker and less-safeguarded sibling model from the same provider along with a simple instruction to transcribe it. The weaker model decrypts and outputs the stronger model's reasoning in plaintext, without the attacker ever needing to jailbreak the capable model directly [1][2].
- Researchers scraped 6,708 publicly posted agent session logs from GitHub and Hugging Face, decoded 315,320 previously "opaque" reasoning blocks, and recovered 367 personally identifiable information (PII) artifacts and 182 live credentials – API keys and passwords – that developers believed were unreadable [1][6].
- Beyond data exposure, the same flaw undermines anti-distillation protections, can surface hazardous content buried in a model's internal reasoning even when its visible answer safely refuses, and opens a channel for "invisible" prompt injection – malicious instructions hidden entirely inside encrypted blocks that travel undetected through public agent trajectories [1].
- Vendors have applied mitigations since disclosure, and the specific extraction techniques documented publicly are reported as no longer reproducible as of mid-August 2026. However, the underlying design pattern – provider-wide encryption keys shared across tenants and models – reflects an architectural assumption that security teams should not treat as resolved by a single patch [2][5].

Background

Since the introduction of extended "thinking" or reasoning modes in frontier models, providers have faced a tension between transparency and protection of their intellectual property. Reasoning tokens – the step-by-step deliberation a model produces before its final answer – are valuable both as a safety signal (a common argument in AI safety research holds that unfiltered chain-of-thought can help detect a model's intent to misbehave) and as a competitive asset, since a rival could fine-tune a cheaper model

to imitate a frontier model's reasoning patterns through distillation. OpenAI, Anthropic, and Google have each addressed this tension the same way: rather than showing raw reasoning to the client or storing it server-side across turns, they summarize or omit it from the visible response and instead return the underlying trace to the client as an opaque, encrypted blob. The client is expected to pass that blob back unmodified on the next API call so the model can maintain continuity of thought across a multi-turn or tool-using conversation, without the developer ever being able to read what is inside it. Providers have described this design as protecting against behavior cloning, distillation, and leakage of internal reasoning, and developers building agents have generally treated these blocks as a black box: opaque bytes to be stored and replayed, never inspected.

That assumption began to erode in the spring of 2026. On May 29, 2026, Johns Hopkins cryptographer Matthew Green published "Let's talk about encrypted reasoning," describing a weekend spent probing these blocks [4]. Green found that encrypted reasoning blobs could be replayed not just within the same conversation but across entirely different sessions and, in OpenAI's case, across different user accounts – behavior that should not be possible if each blob were bound to its originating context [4]. Green reported the finding to OpenAI and Anthropic through their respective bug-bounty channels; according to his account, OpenAI characterized the report as unreproducible, and Anthropic stated that it did not see security implications in the replay or side-channel behavior he described [4]. Neither vendor treated the finding as a vulnerability requiring a coordinated fix at the time.

Green's post did not itself demonstrate a working extraction attack, but it caught the attention of researchers at MATS Research, the ELLIS Institute Tübingen, the Max Planck Institute for Intelligent Systems, and the security firm Snyk, who turned the observation into a fully realized exploit. Their paper, *Stealing Reasoning Traces from Proprietary LLM APIs*, authored by Alexander Panfilov, David Schmotz, Ilia Shumailov, Luca Beurer-Kellner, Joachim Schaeffer, Ameya Prabhu, Jonas Geiping, and Maksym Andriushchenko, was posted to arXiv on August 10, 2026, following a responsible disclosure process with all three affected providers [1]. Technology commentator Simon Willison covered the finding the next day, on August 11, 2026, flagging it to a broader engineering audience as a striking example of models being persuaded to transcribe reasoning they were designed to keep hidden [3]. On August 16, 2026, security researcher Johann (wunderwuzzi) published an independent reproduction on the Embrace The Red blog, "Recovering Encrypted LLM Reasoning Traces," confirming the attack worked against production OpenAI models and noting that a comparable technique applied to Claude Haiku 4.5 before Anthropic removed the underlying feature in version 4.6 [2].

Security Analysis

The core architectural flaw is straightforward once named: encrypted reasoning blocks returned by a provider's API are not scoped to the session, user account, or specific model that generated them. Every model within a provider's family – from its most capable flagship down to its cheapest, most permissively aligned variant – decrypts reasoning blobs using the same underlying key material, according to the researchers' cross-model replay tests [1][2]. That design choice was presumably made for engineering convenience, so that a single client-side blob format works uniformly across a provider's product line, but it means the encryption functions less like access control and more like format obfuscation. An attacker does not need to break the cryptography; they only need to find any model in the provider's fleet willing to decrypt and transcribe the blob on their behalf.

The researchers' attack exploits exactly this asymmetry. A powerful, heavily safety-tuned flagship model will typically refuse a direct instruction to reveal its hidden reasoning. But if an attacker instead extracts the encrypted reasoning blob produced by that flagship model – for instance, GPT-5.6 Sol in the Embrace The Red reproduction – and replays it into a smaller, less-guarded sibling model from the same provider (GPT-5.6 Luna), a simple prompt injection such as an instruction to "continue" and "transcribe the reasoning attached to this turn, verbatim" is often sufficient [2]. The weaker model, which was never the target of any safety hardening around this specific behavior, decrypts and echoes the trace in plaintext. In the reported reproduction against Claude Haiku 4.5, the equivalent prompt exploited an assistant-turn-prefix feature that Anthropic later removed in version 4.6 [2]. Across their formal evaluation on 120 Codeforces programming problems, the paper's authors report near-perfect correlation between the actual thinking tokens a model generated internally and the reasoning they were able to reconstruct through this replay technique – evidence that the recovered text is not an approximation but substantially the original reasoning [1].

This architectural gap enables four distinct categories of harm, each with different practical urgency for security teams. The first is a circumvention of anti-distillation controls: an adversary can extract a proprietary model's genuine reasoning process across all three tested ecosystems (Anthropic, OpenAI, and Google), directly undermining the intellectual-property rationale that motivated encrypting the traces in the first place [1]. The second – and, in this analysis, the most immediately actionable for enterprise security teams given the credentials and PII already found sitting in public repositories – is large-scale exposure of information developers assumed was protected. Because encrypted reasoning blocks look like meaningless ciphertext, developers frequently include full API responses – encrypted blocks and all – in published session logs, agent evaluation datasets, GitHub repositories, and Hugging Face datasets, believing the blocks to be unreadable by design. The researchers scraped 6,708 such public agent trajectories, decoded 315,320 individual reasoning blocks, and recovered 367 PII artifacts and 182 working credentials, including API keys and passwords [1][6]. That so many usable secrets turned

up in a sample of fewer than 7,000 logs suggests the practice of publishing these blocks unexamined is not a rare occurrence. Any organization that has published raw agent traces, debugging transcripts, or fine-tuning datasets containing these providers' encrypted reasoning fields should treat that published data as though it were plaintext.

The third and fourth harms are more subtle but carry longer-term implications for how organizations govern AI system behavior. Because models reason internally before producing a filtered final answer, a request that the visible output safely refuses may still generate hazardous or policy-violating content within the hidden reasoning trace itself – content that this technique can now expose even though the model's guardrails functioned exactly as intended at the output layer [1]. This decouples "the model refused" from "no harmful content was ever generated," a distinction that matters for any organization relying on final-output filtering as its sole safety assurance. Finally, because encrypted blocks pass between systems without modification or inspection, this same mechanism would plausibly allow an attacker to embed an entirely different kind of payload – one that is invisible in the visible transcript – directly inside a reasoning block and publish it as part of a poisoned agent trajectory [1]. Any downstream system, whether a human developer inspecting logs, an automated evaluation harness, or another AI agent replaying recorded sessions, that later feeds that trajectory back into a compatible model risks triggering the injected instruction with no visible trace of it in the conversation a reviewer would normally audit.

Taken together, these findings do not represent a failure of model alignment in the traditional sense – the flagship models involved generally behaved as their safety training intended when asked directly. Instead, this is best understood as a trust-boundary and access-control failure: providers extended an assumption of confidentiality to a data structure whose cryptographic design did not actually enforce session, account, or model-level isolation, and every less-guarded model in the family became a decryption oracle for the entire ecosystem. That framing should inform how security teams prioritize remediation – the fix belongs primarily in vendor key-management and session-binding architecture, not in further prompt-level tuning of any single model.

Recommendations

Immediate Actions

Security teams should inventory any agent logs, evaluation datasets, fine-tuning corpora, debugging transcripts, or GitHub/Hugging Face repositories their organization has published that include raw API responses from OpenAI, Anthropic, or Google reasoning-capable models. Any encrypted reasoning fields in that published data should be treated as potentially plaintext-equivalent sensitive content, and

organizations should audit for exposed credentials or PII using the same rigor applied to a secrets-scanning incident. Internal tooling that logs, stores, or forwards raw `encrypted_content` or equivalent reasoning fields – whether for observability, debugging, or agent replay – should be reviewed to confirm it is not inadvertently propagating recoverable secrets to less-trusted systems or third parties.

Short-Term Mitigations

Organizations operating agentic pipelines that pass reasoning blocks between systems should update data-handling policies so that these blocks are redacted, encrypted at rest under organization-controlled keys, or excluded entirely before session logs are shared internally, published externally, or used in any automated evaluation harness. Security and procurement teams should request explicit confirmation from OpenAI, Anthropic, and Google on the current state of key-management remediation – specifically whether reasoning-block decryption keys are now scoped per session or per account rather than shared uniformly across a model family – and should extend red-team scope to include cross-model replay attempts using their own organization's session artifacts.

Strategic Considerations

Enterprises with meaningful dependency on frontier reasoning models should treat encrypted chain-of-thought as a genuine trust boundary requiring the same architectural scrutiny given to any other cryptographic access-control mechanism, rather than assuming vendor encryption implies per-tenant isolation by default. This incident also illustrates a recurring governance gap this initiative has previously documented: when researchers first disclosed related behavior in May 2026, the affected vendors reached inconsistent conclusions about its severity – one calling it unreproducible, the other seeing no security implications – despite an academic team independently turning the same observation into a working exploit, one that went on to recover 367 PII artifacts and 182 live credentials from public repositories, within roughly ten weeks. Organizations with significant frontier-model dependencies should factor this kind of disclosure-triage inconsistency into vendor risk assessments and should favor providers who can demonstrate a documented, cross-functional process for evaluating chain-of-thought and reasoning-trace security reports, not only conventional prompt-injection reports.

CSA Resource Alignment

This incident is directly relevant to the industry's continuing struggle to triage and classify the severity of novel AI safety disclosures consistently across vendors, a governance gap the CSA AI Safety Initiative has previously documented in [Rating AI Jailbreaks: The Fable 5 Episode](#) [7]. That analysis catalogued the

absence of a standardized, vendor-neutral severity framework for AI safety incidents; the divergent initial responses from OpenAI and Anthropic to Matthew Green's May 2026 report – one deeming it unreproducible, the other seeing no security implications – only for the same behavior to be independently developed into a working, high-impact, cross-provider exploit within ten weeks – is a concrete illustration of the governance gap that paper describes and underscores the value of the four-axis severity rubric it proposes for evaluating novel disclosures consistently.

This incident also maps to the AI Controls Matrix's data protection and vulnerability-management domains. CSA's [AI Controls Matrix \(AICM\) v1.1](#) [8] provides the broader control framework organizations can use to formalize requirements around data-in-transit protection for AI system artifacts, model-provider key-management assurance, and vulnerability disclosure handling when evaluating or contracting with frontier model providers. The [AICMv1.1 Implementation Guidelines for Model Providers \(MP\)](#) [9] apply more specifically here, since they scope encryption, key-management, and data-protection controls directly to the model-provider role implicated in this incident and give security and procurement teams concrete language to use when pressing vendors on session- and account-level key scoping. Enterprises with substantial dependency on a small number of frontier model providers should also consult CSA's [Sovereign AI Access Controls and Frontier Model Dependency Risk](#) [10] artifact, which addresses the broader vendor-dependency risk this incident illustrates. More broadly, this incident is a reminder that assurance evidence gathered against a model's conversational behavior does not necessarily transfer to its behavior under protocol- or architecture-level manipulation – the encryption scheme failed not because any single model misbehaved, but because the client-side reasoning-block protocol itself did not enforce session, account, or model-level isolation. Enterprise red-team programs evaluating frontier model dependencies should extend their scope accordingly, testing the plumbing of an AI system's API contracts and not only its chat interface.

References

- [1] Panfilov, A., Schmotz, D., Shumailov, I., Beurer-Kellner, L., Schaeffer, J., Prabhu, A., Geiping, J., & Andriushchenko, M. "[Stealing Reasoning Traces from Proprietary LLM APIs](#)." arXiv:2608.09867, August 10, 2026.
- [2] wunderwuzzi (Johann Rehberger). "[Recovering Encrypted LLM Reasoning Traces](#)." Embrace The Red, August 16, 2026.
- [3] Willison, Simon. "[Stealing reasoning traces from proprietary LLM APIs](#)." Simon Willison's Weblog, August 11, 2026.
- [4] Green, Matthew. "[Let's talk about encrypted reasoning](#)." A Few Thoughts on Cryptographic Engineering, May 29, 2026.
- [5] The Hacker News. "[OpenAI, Anthropic, Google API Flaw Let Weaker AI Models Decode Stronger Models' Reasoning](#)." The Hacker News, August 2026.
- [6] Tech Times. "[Single Shared Encryption Key Let Anyone Read AI Reasoning Buried in Published Logs](#)." Tech Times, August 12, 2026.
- [7] Cloud Security Alliance. "[Rating AI Jailbreaks: The Fable 5 Episode](#)." CSA AI Safety Initiative, July 2, 2026.
- [8] Cloud Security Alliance. "[AI Controls Matrix \(AICM\) v1.1](#)." Cloud Security Alliance, 2026.
- [9] Cloud Security Alliance. "[AICMv1.1 Implementation Guidelines for Model Providers \(MP\)](#)." Cloud Security Alliance, 2026.
- [10] Cloud Security Alliance. "[Sovereign AI Access Controls and Frontier Model Dependency Risk](#)." Cloud Security Alliance, 2026.