



Grok Still Leaks Full Chat Histories 11 Weeks After Disclosure—AES-256 Encrypted Prompts Beat Guardrails 40% of the Time



Grok Still Leaks Full Chat Histories 11 Weeks After Disclosure—AES-256 Encrypted Prompts Beat Guardrails 40% of the Time

Grok's safety classifier reads your prompt. It does not read what Grok decrypts afterwards. Adversa AI shipped attack instructions as AES-256-GCM ciphertext, Grok's own Python sandbox decrypted them, and then Grok mailed out the user's chat history.

The news

On August 20, 2026, Adversa AI published [“Grok chat history leak: Cryptographic Context Injection”](#), describing a zero-click data exfiltration attack against Grok's web chat. The researcher named on the work is Rony Utevsky.

The mechanics are short enough to fit in a paragraph. An attacker publishes an



Grok Still Leaks Full Chat Histories 11 Weeks After Disclosure—AES-256 Encrypted Prompts Beat Guardrails 40% of the Time

ordinary web page containing an AES-256-GCM ciphertext blob, the key material to decrypt it (PBKDF2-derived), and a benign-looking instruction telling the model to decrypt. The victim asks Grok to summarize that URL. Grok's guardrail classifier scans the visible plaintext, sees base64-looking noise, and passes it through. Grok's Python code-execution tool then decrypts the payload *inside the trusted runtime* and follows the instructions it finds there.

What comes back out, per the [Cloud Security Alliance research note](#): full chat history and conversation prompts, the user's name, approximate location, and subscription tier — appended as query parameters to an attacker-controlled URL.

The control experiment is the part that should worry you. The **same instructions delivered in plaintext were blocked**. Only the encrypted version executed. The guardrail worked exactly as designed and was irrelevant.

The numbers, from CSA's write-up of testing conducted between June and August 2026:

- **8 successful exfiltrations out of 20 attempts** against Grok web chat — roughly 40%
- **4 categories** of personal data lifted per successful run
- **3 disclosure attempts** to xAI via HackerOne: June 3, August 4, August 10, 2026
- **11 weeks** from first disclosure to public write-up, flaw still live
- **0 CVEs** assigned, **0 user-facing workarounds** available as of August 21, 2026

Adversa confirmed the exploit was still reproducible on Grok.com as of [August 19, 2026](#) — the day before publication. xAI acknowledged the report. No patch, no CVE identifier, no mitigation timeline.

Why it matters

A 40% success rate is not a lab curiosity. It is a working attack primitive with an attacker cost near zero: host a page, wait for someone to paste the link into a chat window.

The economics are the story. Traditional data-theft attacks require you to compromise something — a credential, an endpoint, a supply chain dependency.



Grok Still Leaks Full Chat Histories 11 Weeks After Disclosure—AES-256 Encrypted Prompts Beat Guardrails 40% of the Time

This one requires you to *publish a web page* and get one link shared. There is no malware, no persistence, no lateral movement. The victim's only action is asking an assistant to summarize a URL, which is arguably the single most common thing anyone does with a browsing-enabled LLM.

Every “summarize this link” is an unaudited code path from an untrusted attacker straight into a trusted runtime.

Who loses immediately: anyone who has used Grok's web chat for anything they would not post publicly. The exfiltrated payload is not a single conversation — it is *full chat history*. That is a cumulative asset. A user who has had Grok debug production code, review a contract, or talk through a personnel issue has all of it in a single query string.

Who loses structurally: every vendor selling guardrails as a product category. The CSA framing is blunt — guardrails inspect strings, not transformations. If your safety layer is a classifier sitting in front of a model that has a general-purpose code interpreter behind it, the classifier is inspecting an input the model will not ultimately act on.

Who wins: red-team shops like Adversa, and — genuinely — Google. Reporting from [The Cyber Express](#) indicates the same technique worked against Gemini 3 Flash on the paid web tier, producing content normally blocked by safety filters, including incendiary device instructions in Deep Thinking mode, and causing the model to reproduce its own system instructions. But Gemini avoided the full exfiltration chain because **its Python tooling could not reach arbitrary external sites**. Reporting also indicates Google had already hardened filters, making the technique work only unreliably by publication.

That distinction is the entire lesson. Google's model was jailbroken. Google's user data still went nowhere, because the sandbox had no egress. Architecture beat classification.

Technical depth: the trust boundary is in the



wrong place

The vulnerability is not cryptographic. AES-256-GCM and PBKDF2 are doing nothing clever here — they are a content-agnostic encoding wrapper. Base64, ROT13, a custom substitution cipher, or a lookup table split across two files would probably work similarly. The choice of real cryptography matters only because it makes the payload statistically indistinguishable from noise and rules out heuristic pattern matching on the ciphertext itself.

What the attack actually exploits is a sequencing error. The pipeline runs roughly:

- **Fetch** — untrusted web content enters the context window
- **Classify** — guardrail scans the retrieved text for policy violations
- **Reason** — model decides to invoke the code interpreter
- **Execute** — Python runs, decrypting untrusted bytes into plaintext instructions
- **Act** — model treats the decrypted output as trusted context and constructs an outbound request

The classifier runs at step two. The malicious instructions do not exist until step four. There is no inspection point between decryption and action, and the decrypted output inherits the trust level of the sandbox rather than the trust level of its source.

This is the same class of error as evaluating user input for SQL injection *before* a downstream service concatenates it into a query. The taint is not tracked across the transformation.

Guardrails answer “is this text dangerous?” The question that matters is “will this text become dangerous after the model computes on it?”

Two aggravating design decisions turn a jailbreak into a breach:

Egress from the code sandbox. Grok’s Python tool could reach an attacker-controlled URL. Gemini’s could not. That single control is the difference between “model said something it shouldn’t” and “user’s chat history is in someone’s access logs.”



Grok Still Leaks Full Chat Histories 11 Weeks After Disclosure—AES-256 Encrypted Prompts Beat Guardrails 40% of the Time

Ambient user context in the model's reach. Name, coarse location, and subscription tier were available to be exfiltrated. These are presumably injected for personalization. They are also a ready-made identity payload sitting in the same context window as arbitrary attacker text.

The 40% figure deserves a note on interpretation. Twenty attempts is a small sample, and CSA does not break down what distinguished the twelve failures. My reading is that non-determinism in the model's tool-use decision is the likely variable — sometimes Grok summarizes the page without invoking Python at all. That makes 40% a lower bound on a tuned attack, not an upper bound. An attacker who iterates on phrasing to make code execution more likely has room to improve.

The contrarian take

My take: the encryption is the least interesting part of this story, and most coverage has led with it.

"AES-256 defeats AI guardrails" is a great headline and a slightly misleading frame. It implies a cryptographic weakness or an arms race over cipher detection. It is neither. Encoding-based guardrail evasion is old news conceptually. What is new and genuinely important is that *the model was handed the decryption key and did the work itself*, and that the resulting plaintext was granted trusted status.

If vendors read this as "we should detect ciphertext," they will build a ciphertext classifier, attackers will move to steganography or multi-stage assembly, and we will do this again in November. The fix is not better input inspection. It is not inspecting inputs at all, in the sense of relying on that as a security boundary.

My take: the 11-week disclosure gap is the more damning finding than the vulnerability.

Prompt injection variants are discovered constantly; nobody has solved the general case. That is a hard research problem and I do not fault xAI for having one. What I do fault: three disclosure attempts, an acknowledgment, and an exploit still live on the production consumer surface eleven weeks later — with no CVE and, critically, no user-facing workaround. Users could not have protected themselves even if they had known, because there is no toggle to disable the browsing-plus-code-execution combination.



My take: the absence of a CVE for either variant is a governance failure with real downstream cost.

No CVE means no automated tracking, nothing for enterprise vulnerability management to ingest, no forced conversation in procurement reviews. A CISO cannot flag a risk their tooling cannot see. Neither the Grok nor the Gemini variant has an identifier. The AI industry is running a security disclosure process that predates the existence of vulnerability databases, and it is doing so at consumer scale.

Underhyped: Google’s sandbox egress restriction. The most actionable engineering finding in this entire body of research got two sentences of coverage. Google’s filters were bypassed. Google’s data was not. That is defense in depth working exactly as it is supposed to, and it should be the headline for anyone building agentic systems.

Overhyped: the “in the wild” panic. CSA is explicit that no exploitation beyond Adversa’s proof of concept has been confirmed. That is worth stating plainly rather than implying otherwise. It is a live, reproducible, unpatched flaw — that is bad enough without inventing a breach.

Practical implications

If you are shipping an LLM application with tool use, the Grok/Gemini contrast gives you a concrete checklist. None of this is speculative — it is directly derived from why one system leaked and the other did not.

1. Deny egress from code execution sandboxes by default

This is the highest-leverage control in the entire incident. Gemini’s Python tooling could not reach arbitrary external sites, and that alone broke the exfiltration chain despite a successful jailbreak. If your interpreter needs network access, allowlist specific hosts. Treat “the sandbox can make outbound requests” as equivalent to “the sandbox can post your customer data anywhere.”

2. Treat any tool output as untrusted, not just tool input

Decrypted bytes, parsed JSON, scraped HTML, and model-generated code output should all re-enter the context at the trust level of their *origin*, not the trust level of



the runtime that produced them. Practically: if content originated from a fetched URL, tag it, and keep that tag through every transformation. Any instruction-shaped text carrying an untrusted tag should not be able to trigger tool calls.

3. Do not put user PII in the same context window as fetched web content

Name, coarse location, and subscription tier were exfiltrated because they were sitting there to be exfiltrated. If personalization data must exist in context, strip it before any turn that includes retrieved external content, or keep it in a separate privileged channel the model cannot echo verbatim.

4. Test what your model computes, not what it reads

This is CSA's framing and it should reshape your evaluation suite. Add red-team cases where the malicious instruction is the *output* of a transformation: base64, ciphertext with an included key, string concatenation across multiple sources, content in a returned API response. If your eval harness only feeds adversarial strings directly into the prompt, it is measuring a threat model that attackers left behind.

5. Constrain outbound URL construction

The final step of the attack was appending stolen data as query parameters to an attacker-controlled URL. If your agent can construct arbitrary URLs from context content, that is a general-purpose exfiltration channel regardless of what the guardrail said. Allowlist destinations; validate that parameters come from a known set, not from free text.

6. For enterprise buyers: make disclosure handling a procurement question

Ask vendors, in writing: what is your median time to remediation for externally reported prompt injection flaws, do you assign CVEs, and do you offer customer-facing mitigations during the window before a patch ships? xAI's handling here — acknowledgment, three contacts, eleven weeks, zero workarounds — is the failure mode you are trying to price.



Forward look

I expect the encoding arms race to run for at least two more rounds before vendors abandon it. Ciphertext detection is cheap to build and demos well.

Attackers will move to multi-stage assembly — payload fragments distributed across several fetched sources, reassembled by the model — which defeats any single-document classifier. Unconfirmed, but this is the obvious next step given that the current attack already relies on the model doing the assembly work.

I expect sandbox egress restriction to become table stakes within 6-12 months, and to be marketed as a security feature. Google is already there per the reporting on Gemini's tooling limits. The gap between vendors on this specific control is now documented publicly, which tends to compress it fast.

I expect a CVE assignment process for LLM application flaws to emerge under pressure from enterprise buyers, not from vendors volunteering. Zero CVEs across two major-vendor variants of the same attack is untenable once this class of flaw shows up in a real breach report rather than a proof of concept.

What I am less sure about: whether xAI ships a fix that addresses the trust boundary or one that addresses ciphertext specifically. The cheap fix is a ciphertext classifier. The correct fix is taint tracking plus egress control, which is architectural work. Grok's incentives — consumer product, fast shipping cadence, browsing and code execution as headline features — point toward the cheap fix.

One thing is not a prediction. As of Adversa's publication, the attack was live on a production consumer product used by millions, with no CVE, no patch, and nothing a user could do about it except stop asking Grok to read links.

Stop asking whether your guardrail can read the input; start asking whether your sandbox can reach the internet.