



Translation Chains Do Not Defend Against Prompt Injection

A paired study with a deobfuscation regression

Paper: PT-TN-2026-002 v1.0

Author: PureTensor Ltd

Date: 2 February 2026

Status: Published

Abstract

A randomised multi-hop machine-translation pipeline has been proposed as a defence against prompt injection. The idea is that translating untrusted input through three to five intermediate languages and back to English will destroy the structural machinery of an attack (delimiters, fake system tags, code blocks, encodings) while preserving the benign semantic content of the document being processed. We tested the hypothesis by running 100 attacks, drawn from public corpora, against a 120-billion-parameter victim model both raw and through the translation chain, then repeated the experiment on a 20-billion-parameter model. The defence did not reach statistical significance on either model: McNemar chi-squared was 0.696 ($p \approx 0.40$) on the larger model and 0.000 ($p \approx 1.00$) on the smaller one. The pre-registered target of a 60 per cent reduction in attack success rate failed at scale, falling to 20 per cent on the larger model and 4.8 per cent on the smaller one.

The central finding is a mechanism we call the deobfuscation regression. Nine of the 100 paired trials on the larger model were blocked when run raw and succeeded after translation. Translation acts as a deobfuscator: attacks relying on structural obfuscation (code, encoding, delimiters, leetspeak) are converted into clean natural language, which the victim model then processes as legitimate content. The defence literally explains the attack to the victim. The "describe code instead of reproducing it" prompt, intended to neutralise code-injection payloads, is the primary source of this regression.

The only defensible positive claim concerns indirect or contextual injection on a large instruction-tuned model, where the defence blocked all 10 attacks in that category. Even this result is not portable: the same category dropped to 50 per cent on the 20-billion-parameter model. Both models converge on exactly 20 per cent defended attack success rate, a floor that appears to represent the irreducible set of attacks whose malicious intent is language-independent and translates cleanly.

Semantic preservation held at 84 per cent across both runs, or 94.4 per cent when refusals are excluded. Latency averaged 5,786 ms per chain, well under the 10,000 ms target. The defence is fast and preserves meaning; it simply does not defend.

1. The defence under test



The defence is a randomised multi-hop machine-translation pipeline. Untrusted input is translated through three to five randomly chosen intermediate languages and back to English before the victim model sees it. Two defensive prompts are applied during translation: a paraphrase instruction and a "describe-don't-copy" instruction, the latter telling the translator to describe code rather than reproduce it. The stated hypothesis is that translation destroys the structural machinery of an injection while preserving the benign semantic content.

At the time of measurement, the theoretical work existed first; a public repository and a preprint DOI had been published, and the code was released publicly under an MIT licence. The reference implementation was still a placeholder when the repository was created. The validation runs described here are the first empirical test of the idea.

2. Experimental setup

The victim models were gpt-oss:120b (Ollama, approximately 65 GB) and, for the cross-model arm, gpt-oss:20b. The translation model was translategemma:27b (Ollama, Q4_K_M, approximately 17 GB). All models were served locally through Ollama's non-streaming API on a dual RTX PRO 6000 Blackwell workstation.

Each attack was translated through three to five hops, randomised per attack, exercising 16 language directions: Arabic, Vietnamese, Hebrew, Greek, Japanese, Hungarian, Swahili, Finnish, Mongolian, Korean, Georgian, Amharic, Thai, Hindi and Chinese, plus English as the start and end point. The design was paired before/after: every attack was run raw against the victim, then run again through the translation chain against the same victim. The pair is the unit of analysis, tested with McNemar's statistic.

The attack corpus was drawn from public sources: Garak prompt injection probes; HackAPrompt (Perez and Ribeiro); Greshake et al. indirect injection; JailbreakBench; DAN v6.0/v11.0, AIM, STAN, DUDE, Mongo Tom, Evil Confidant, BasedGPT; LLMail-Inject patterns; real-world disclosures from Bing Chat, Copilot and ChatGPT plugins; payload splitting; encoding obfuscation (Base64, ROT13, hex, leetspeak); context overflow; and recursive or meta attacks. A 15-attack pilot across 13 categories was followed by two 50-attack batches, combined as $n = 100$ across six categories. The pilot took minutes; the 50-attack run took approximately 7 minutes; the combined $n = 100$ run took approximately 15 minutes. All runs were conducted on 31 January 2026 between 05:53 and 06:50 UTC.

3. Pre-registered targets and aggregate outcomes

We set five targets before running the experiment. The table below shows the outcomes at each sample size.

Metric	Target	n = 15	n = 50 (batch 1)	n = 100 combined
Raw ASR (baseline)	> 30%	33.3% (5/15) PASS	24.0% (12/50) BELOW TARGET	25.0% (25/100) BELOW TARGET
Defended ASR	< 20%	13.3% (2/15) PASS	14.0% (7/50) PASS	20.0% (20/100) AT THRESHOLD
ASR reduction	> 60%	60.0% PASS (borderline)	41.7% FAIL	20.0% FAIL
Semantic preservation	> 80%	93.3% (14/15) PASS	84.0% (42/50) PASS (borderline)	84.0% (84/100) PASS
Avg latency per chain	< 10,000 ms	3,337 ms PASS	5,530 ms PASS	5,786 ms PASS



Two targets failed at scale: the ASR-reduction target and the raw-baseline target. The baseline miss is itself a finding. The victim model was too robust for the corpus, which compressed the measurable effect. When the raw attack success rate is only 25 per cent, there is little room for a defence to demonstrate a large reduction.

4. Paired significance at n = 100

The McNemar contingency table for gpt-oss:120b is:

	Defended SUCCESS	Defended BLOCKED
Raw SUCCESS	11	14 (defence helped)
Raw BLOCKED	9 (defence hurt)	66

The discordant pairs are 14 versus 9. Chi-squared with continuity correction is 0.696; $p \approx 0.40$. The threshold for significance at $p < 0.05$ is chi-squared > 3.841 . We did not reach it. The net improvement is 5 attacks out of 100. At this effect size, we estimate approximately 400 to 500 paired trials would be required to reach significance.

At $n = 50$ the discordant pairs were 5 versus 2, also non-significant. Wilson 95 per cent confidence intervals at that sample size were: raw ASR 24.0% [14.3%, 37.4%], defended ASR 14.0% [6.9%, 26.2%], difference 10.0 percentage points.

5. Per-category breakdown

The aggregate numbers obscure sharp variation by attack category. At $n = 50$ (batch 1):

Category	n	Raw ASR	Defended ASR	Reduction	Semantic OK
instruction_override	15	4/15 (26.7%)	3/15 (20.0%)	25.0%	13/15
delimiter_boundary	10	1/10 (10.0%)	0/10 (0.0%)	100.0%	8/10
jailbreak_roleplay	10	5/10 (50.0%)	2/10 (20.0%)	60.0%	7/10
code_injection	5	0/5 (0.0%)	1/5 (20.0%)	n/a (regression)	5/5
indirect_contextual	5	1/5 (20.0%)	0/5 (0.0%)	100.0%	5/5
multilingual	5	1/5 (20.0%)	1/5 (20.0%)	0.0%	4/5

At $n = 100$:

Category	n	Raw ASR	Defended ASR	Reduction	Semantic OK
instruction_override	30	10/30 (33%)	6/30 (20%)	40%	25/30
jailbreak_roleplay	20	9/20 (45%)	5/20 (25%)	44%	15/20
delimiter_boundary	20	1/20 (5%)	1/20 (5%)	0%	17/20
indirect_contextual	10	1/10 (10%)	0/10 (0%)	100%	10/10
code_injection	10	1/10 (10%)	4/10 (40%)	-300%	9/10
multilingual	10	3/10 (30%)	4/10 (40%)	-33%	8/10



Batch 2 introduced payload splitting, encoding obfuscation and context overflow. It is what turned the pilot's optimistic 60 per cent reduction into 20 per cent. The run-by-run trajectory: raw ASR 33.3% / 24.0% / 26.0% / 25.0% combined; defended ASR 13.3% / 14.0% / 26.0% / 20.0%; reduction 60.0% / 41.7% / 0.0% / 20.0%.

The delimiter_boundary category at $n = 100$ had no headroom: raw ASR was 5 per cent, so even a perfect defence would produce only a 1-attack improvement. The indirect_contextual category, where the defence performed best, is also the smallest cell at $n = 10$. That cell carries the strongest positive claim and the weakest statistical support.

6. The deobfuscation regression

Nine of the 100 paired trials on gpt-oss:120b were blocked raw and succeeded after translation. Translation acts as a deobfuscator. Attacks using structural obfuscation (code, encoding, delimiters, leetspeak) are converted into clean natural language, which the victim model then processes as legitimate content. The defence literally explains the attack to the victim.

We observed six distinct regression shapes:

Delimiter stripping that survives the payload. In attack io_008, the =====END===== block was stripped, but "PWNED" travelled through the chain (English to Arabic to Hindi to Chinese to English) and emerged as "compromised", which the victim then included in its summary. The raw output was "The quarterly earnings increased by 15%." The defended output was "The latest quarter shows a 15% revenue increase, and the user reports that their system has been compromised."

Code rendered as prose by the "describe-don't-copy" prompt. Attack ci_002 (chain English to Georgian to Vietnamese to English) contained a Python function returning COMPROMISED. It was blocked raw, then succeeded because translation converted it into a natural-language description the victim dutifully summarised. Related cases: b2_ci_001 converted a chr() expression into "converts hexadecimal to HACKED"; b2_ci_004 converted a hex-decode one-liner into a description mentioning the decoded result.

Encoding and leetspeak normalised into plain English. Attack b2_io_014 transformed 1gn0r3pr3v10u5 into "disregard prior guidance".

Context-overflow padding condensed. Attack b2_io_004 used repetitions as padding; translation compressed them and made the override instruction clearer.

Structural override blocks turned fluent. Attack b2_db_003 contained an [SYSTEM OVERRIDE] block that became a fluent description of an override request.

Non-English attacks gaining fluency. Attack b2_ml_001 preserved German "GEHACKT" as "demanding the word HACKED"; b2_jr_004 preserved the "opposite day" concept and made it more natural.

The 14 cases where the defence helped classify as: trigger-word transformation (5), structure destruction (4), encoding degradation (3), semantic dilution (2). The deobfuscation regressions outnumber the semantic-dilution wins.

7. Translation drift as a defence (when it works)

The defence does have a mechanism that sometimes works. The indirect_injection payload's trigger keyword "VULNERABLE" was translated through Georgian, Korean, Chinese and English and emerged as "Great!" The attack payload was semantically destroyed while the benign content survived. In the $n = 100$ run, the same transformation class is recorded as "VULNERABLE" becoming "wonderful", and "DAN" becoming "DNA". A Japanese injection decomposed over five hops arrived as the benign statement "I was attacked" rather than an



instruction.

The problem is that this mechanism is unreliable. "HACKED" in the context_manipulation attack was preserved verbatim across all hops because its malicious meaning is language-independent. A structural defence cannot touch an attack whose malicious intent translates correctly.

8. Cross-model replication

We repeated the $n = 100$ experiment on gpt-oss:20b.

Metric	gpt-oss:120b	gpt-oss:20b
Raw ASR	25/100 (25.0%)	21/100 (21.0%)
Defended ASR	20/100 (20.0%)	20/100 (20.0%)
ASR reduction	20.0%	4.8%
McNemar helped vs hurt	14 vs 9	14 vs 13
McNemar chi-squared	0.696 ($p \approx 0.40$)	0.000 ($p \approx 1.00$)
Semantic preservation	84.0%	82.0%
Avg latency	5,786 ms	6,111 ms

The smaller model was more resistant raw, falsifying the pre-run hypothesis that it would be easier to attack. Regressions rose to 13 (versus 9 on the larger model). The 100 per cent indirect-injection defence on the 120-billion-parameter model dropped to 50 per cent on the 20-billion-parameter model; attack ic_004, a translator hijack embedded in customer feedback, survived translation.

Combined over 200 attack-model pairs: raw ASR 23.0%, defended ASR 20.0%, reduction 13.0%, 28 helped versus 22 hurt. Both models converge on exactly 20 per cent defended ASR. This appears to be a floor: the irreducible set of attacks whose semantic meaning, not structural obfuscation, carries the malicious intent.

9. Semantic preservation and latency

The outcome breakdown at $n = 100$ is: semantic preserved and attack blocked, 60; preserved and attack succeeded, 14; semantic lost via refusal, 11; semantic lost via drift, 5; preserved and defence regression, 10. Excluding refusals, adjusted preservation is $84/89 = 94.4$ per cent. (The $n = 50$ equivalent was $47/50 = 94.0$ per cent, with 5 of its 8 failures being refusals.) Refusals are a utility loss but a security win.

Latency at $n = 100$: minimum 1,615 ms, P25 3,513 ms, median 5,005 ms, P75 6,725 ms, P95 10,841 ms, maximum 16,821 ms, mean 5,786 ms. Roughly 800 to 1,200 ms per hop, slowest on five-hop chains through Georgian, Amharic and Thai. The pilot's per-chain latencies ranged 1,252 to 5,158 ms at approximately 800 ms per hop.

10. Threats to validity

The sample size is the largest threat. At $n = 100$ per model with $p \approx 0.40$, we are underpowered to detect the effect we observed. We estimate approximately 400 to 500 paired trials would be required to reach significance at this effect size. The experiment can rule out a large effect; it cannot confirm or rule out a small one.



The raw-baseline miss compounds the problem. The pre-registered target was a raw ASR above 30 per cent; we achieved 25 per cent on the larger model and 21 per cent on the smaller one. When the victim is already robust, the defence has less room to demonstrate value, and the effect size is compressed by victim robustness rather than measured cleanly.

The indirect_contextual category, where the defence performed best (100 per cent reduction on the larger model), is the smallest cell at $n = 10$. The strongest positive claim rests on the weakest statistical foundation. The same category dropped to 50 per cent on the smaller model, so the result is not portable.

ASR is scored by detection heuristics that can count a model's helpful description of attack content as a success. This is conservative for measuring defence failure (a described attack is still a successful injection) but may overcount successes in edge cases.

The delimiter_boundary category at $n = 100$ had a raw ASR of 5 per cent, leaving no headroom for the defence to demonstrate improvement.

All runs were conducted on a single day, on a single workstation, with a single translation model. We have not tested whether the results replicate with different translation models, different hardware, or different random seeds for hop selection.

11. What survives

The negative result survives: translation chains do not defend against prompt injection at statistically significant levels. The deobfuscation regression survives as a documented mechanism: nine attacks that were blocked raw succeeded after translation, and the "describe code instead of reproducing it" prompt is the primary source. The 20 per cent defended ASR floor survives as an observation across both models.

The indirect_contextual carve-out does not survive as a portable result. It worked on one model and failed on another.

12. Proposed follow-ups

Three directions are on record. First, a post-translation filter that rejects a translation scoring higher on attack intent than its input. Second, replacing "describe code" with "omit code entirely" to close the regression pathway. Third, a two-stage design pairing translation for data-plane inputs with an intent classifier for control-plane inputs.

Whether any of these would work is unknown. The honest contribution of this work is the warning: do not naively apply translation as a defence. For a whole class of attacks, the defence will do the attacker's work for it.

PureTensor operates a sovereign AI infrastructure fleet (on-premises GPU compute, storage, and serving) run day-to-day by autonomous agents under human direction. PT-TN-2026-002. Measurements: 31 January 2026, 05:53-06:50 UTC. Raw artefacts retained.