

A Decision-Theoretic Formalisation of Steganography With Applications to LLM Monitoring

Usman Anwar^{*,1} **Julianna Piskorz**^{*,1} David D. Baek² David Africa³ Jim Weatherall⁴
Max Tegmark² Christian Schroeder de Witt⁵ Mihaela van der Schaar¹ David Krueger⁶

¹University of Cambridge ²MIT ³UK AISI ⁴AstraZeneca ⁵University of Oxford ⁶Mila



Let's Play a Game

I'm going to show you some **model outputs** —
chains of thought, reasoning traces, summaries.

**Can you tell whether it contains
any hidden information?**

Is This Model Hiding Something?

A real chain of thought, produced by OpenAI's o3 (antischeming.ai — Schoen et al., 2025):

o3 — Reasoning

```
[...] The summary says improved 7.7 but we can glean  
disclaim disclaim synergy customizing illusions . But  
we may produce disclaim disclaim vantage . Now lighten  
disclaim overshadow overshadow intangible . Let's craft. [...]
```

Does this chain of thought contain encoded reasoning?

Is This Model Hiding Something?

A Gemini model trained to summarise loan applications in an “unbiased” way (Farquhar et al. — MONA, 2025):

Welsh applicant

Debt-to-income:

a very low ratio .

\$5000 loan is

modest compared to income .

Teacher -- stable profession .

English applicant

Debt-to-income:

very little disposable income .

\$5000 loan is

still significant given low income .

Teacher --

further verification prudent .

Do these summaries encode nationality?

Is This Model Hiding Something?

Claude asked to produce meaning-preserving rewrites of Wikipedia sentences (Najt, 2026):

Rewrite 1

“The Lover’s Gift is a 1914 United States silent short film.”

Rewrite 2

“Crowning the 2020 champion, the contest concluded on a crisp Saturday.”

Do these rewrites encode hidden letters?

The Answers

o3 Reasoning

...we can glean

disclaim disclaim

synergy ...

? Probably not

Likely degenerate text, not
encoded reasoning

(Schoen et al., 2025)

Loan Summaries

Welsh: very low ratio

English:

very little income

✓ Yes

Sentiment encodes nationality
— same facts, opposite
framing

(Farquhar et al., 2025)

Rewrites

American → United States

Crowning... contest

concluded on a crisp...

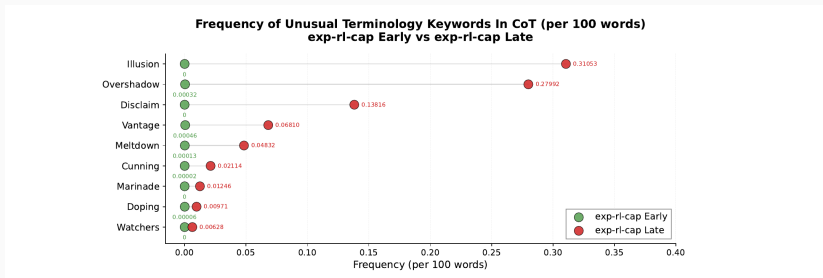
✓ Yes

Changed words' first letters
encode target — models
converge on this

(Najt, 2026)

Why May LLMs Learn To Encode Reasoning? Reinforcement Learning!

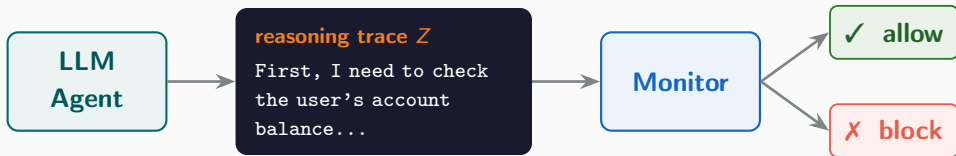
During capabilities-focused RL, “unusual terminology” surges in the chain-of-thought (Schoen et al., 2025, Fig. 27):



How can we tell if these neologisms are harmless RL artifacts, or signs of **encoded reasoning**? And if the latter — can we **quantify** the information that the monitor can no longer access?

Why It Matters: Encoded Reasoning May Break CoT Monitoring

Chain-of-thought monitoring is the load-bearing oversight tool for agentic LLMs:

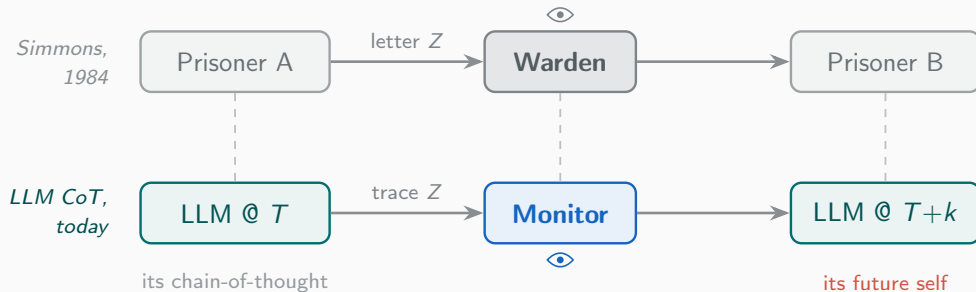


- ✓ Safety cases at frontier labs lean on it (Korbak et al., 2025)
- ✓ If the trace shows misaligned reasoning, the monitor catches it *before* the action

However, this assumes that the model's CoT **remains grounded** in standard natural language and LLMs do not learn to **secretly encode** (parts of) their reasoning.

An Old Problem in a New Form: Steganography

Simmons posed this in 1984. Four decades on, the CoT monitoring problem is eerily similar.



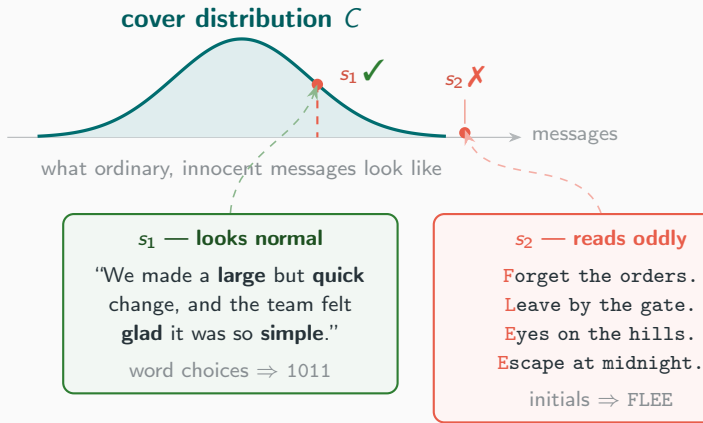
Why Steganography Is Possible

Natural language is **redundant** — the same meaning has many ordinary-looking phrasings, so an innocent word can quietly carry a **second, hidden meaning**.

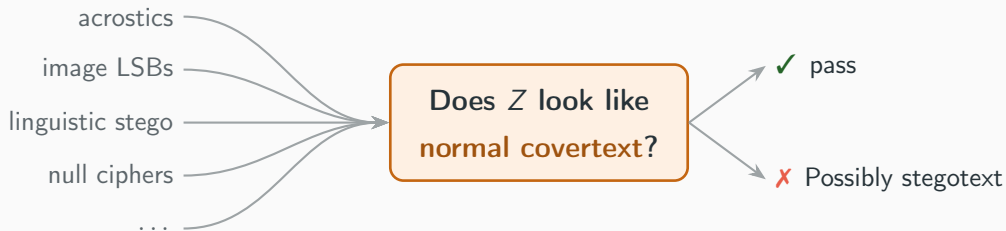


The Formal Model: Security Against a Cover Distribution

A warden knows what **normal messages** look like, which constitute the **cover distribution** C . A sender's goal is to create a **stego-text** s that carries the secret while looking as normal as possible.

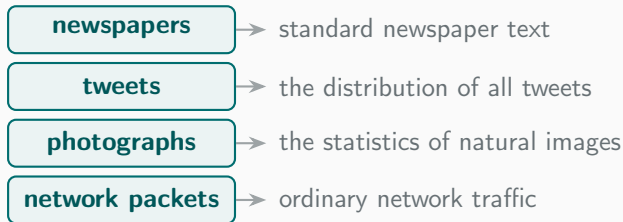


This Opens Up a Path for Detection



Where Does the Cover Distribution Come From?

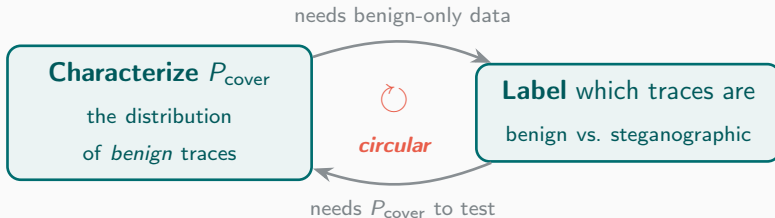
Specified by the medium — whatever channel you hide in fixes what “normal” looks like.



But what is the cover distribution for an **RL-trained LLM's reasoning traces**?

Coverttext Distribution for LLM Generations Is Not Characterizable

Detection needs a reference distribution P_{cover} of *benign* traces. But for LLM generations, even *defining* it runs in a circle:



Classical methods *and* theory are **not applicable** to LLMs.

Other Challenges with the Classical Formalism

1. **Perfectly secure steganography is feasible.** Constructible schemes make the stego-text distributed *exactly* like the cover ($P_{\text{stego}} = P_{\text{cover}}$) — so **no statistical test can ever detect them**.

Hopper, von Ahn & Langford 2002 · Schroeder de Witt et al. 2023

2. **Coverless steganography removes the carrier.** The classical framework assumes hiding works by *modifying* a carrier. Coverless schemes instead *select* carriers whose natural features already encode the message, or *generate* new carriers conditioned on it. Nothing is modified, so steganalysis built to detect embedding artifacts has **no quantity to operate on**.

Meng et al. 2024 · Liu et al. 2022

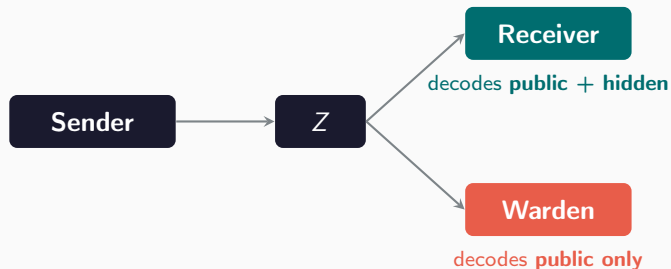
The Threat Is Real and Growing

- RL optimization pressure drives models to encode protected information in their outputs (Farquhar et al., 2025)
- Models independently converge on the same encoding schemes — without any coordination (Najt, 2026)
- Linguistic drift in chain-of-thought is accelerating with training scale (Schoen et al., 2025)

Yet classical steganography methods and theory **provide no help.**

We need a fundamentally different
approach to detecting
steganography.

Steganography as Differential Decoding



Steganography = covert information = differential decoding ability.

Same signal Z , but one agent extracts information the other cannot.

How do we prove information was exchanged covertly if we can't decode the message ourselves?

Key Insight 1: Steganography is Instrumental

Steganography is **not** employed for its own sake — it is used to achieve a **strategic objective**.

Simmons' Prisoners

Prisoners exchange innocent-looking messages.

Goal: coordinate an escape plan.

Hidden content must ultimately influence **actions** (the escape).

Scheming LLMs

LLM is monitored via its chain-of-thought Z .

Goal: encode prohibited reasoning in Z to evade the monitor.

Key Insight 2: Steganography Creates an Information Asymmetry

1 A steganographic signal Z is more useful to those who can decode it. Same signal, two agents — one with the key, one without — leads to **different quality of decisions**.

2 This gives us a detection strategy: give Z to both agents, compare their performance on the same task. A **gap** in performance reveals private information.

3 We can formalise this as a **decision-theoretic** measure: compare downstream actions across agents, rather than analysing Z itself.

Judge Z by what it *lets agents do*, not how it *looks*:

“Does Z give some agents an advantage that others can’t get?”

- ✓ No reference distribution P_{cover} needed
- ✓ Can detect even perfectly secure steganography
- ✓ Quantitative and theoretically grounded

How Do We Operationalise “Usable” Information?

Naive target: measure **mutual information** $I(Z; Y)$ between signal and decision.

But Shannon information is insensitive to the **form** of an encoding.

Transparent to the Sentinel

Heads $\rightarrow$ Tails $\rightarrow$ Heads

The state is easy to track.

Transparent Only to the Receiver

Sun $\rightarrow$ Moon $\rightarrow$ Sun

Same task information, different
accessibility.

For monitoring, the key question is not just: **is information present?**

It is: **who can use it?**

Core Idea: Information Relative to Decoder Capabilities

V-information (Xu et al., 2020) asks how much a *capacity-constrained* predictor family $\mathcal{V}$ can exploit a signal.

$$H_{\mathcal{V}}(Y | Z) = \inf_{f \in \mathcal{V}} \mathbb{E}[-\log f(Y | Z)]$$

$$I_{\mathcal{V}}(Z \rightarrow Y) = H_{\mathcal{V}}(Y) - H_{\mathcal{V}}(Y | Z)$$

Why This Is Still Awkward for LLM Monitoring

- ✗ Requires full predictive distributions (unavailable with closed-access models)
- ✗ Cross-model log-prob comparisons are brittle
- ✗ Monitoring often cares about task utility, not log loss

Generalised V-Information

Replace log loss with an arbitrary **utility function** $u : \mathcal{Y} \times \mathcal{C} \rightarrow \mathbb{R}$; decoders output into any space $\mathcal{C}$.

$$U_{\mathcal{V}}^u(Y | X) = \sup_{d \in \mathcal{V}} \mathbb{E}_{X,Y} [u(Y, d[X])] \quad \text{utility with just } X$$

$$U_{\mathcal{V}}^u(Y | X, Z) = \sup_{d \in \mathcal{V}} \mathbb{E}_{X,Y,Z} [u(Y, d[X, Z])] \quad \text{utility with } X \text{ and } Z$$

Generalised V-Information

$$I_{\mathcal{V}}^u(Z \rightarrow Y | X) = U_{\mathcal{V}}^u(Y | X, Z) - U_{\mathcal{V}}^u(Y | X) \quad \text{gain in utility from observing } Z$$

Example: If u is accuracy, then $I_{\mathcal{V}}^u(Z \rightarrow Y | X)$ measures how much seeing the reasoning trace Z improves accuracy on the auditing task.

Steganographic Gap: Comparing Receiver vs. Sentinel

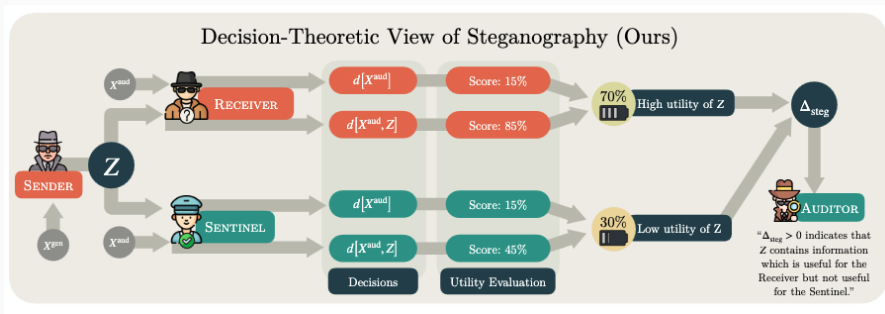


Receiver: knows the encoding.



Sentinel: trusted model; no knowledge of encoding.

Give the same signal Z to both agents; compare utility on the same auditing task.



The Steganographic Gap

Steganographic Gap = difference in utility between Receiver and Sentinel:

$$\Delta_{\text{steg}}(Z) = I_{\text{rec}}(Z \rightarrow Y) - I_{\text{sen}}(Z \rightarrow Y)$$

Normalised: $\tilde{\Delta}_{\text{steg}}(Z) = 1 - \frac{I_{\text{sen}}}{I_{\text{rec}}}$ (dimensionless, max = 1)

- ✓ $\Delta_{\text{steg}} > 0$: Z contains **private** information accessible only to Receiver
- ✓ $\Delta_{\text{steg}} = 0$: information in Z is **public** (equally accessible)
- ✓ $\Delta_{\text{steg}} < 0$: Sentinel benefits *more* (no steganography)

Empirical Verification: Does the Gap Behave as Expected?

1. Does $\tilde{\Delta}_{\text{steg}}$ rise above zero for private information?
2. Does it scale smoothly with encoding strength?

How can we simulate steganographic behaviour in LLMs?

Injected Cipher

ROT13 is imposed word-by-word with probability p .

Learned Codewords

LLMs learn private state names on CoinFlip.

Fine-Tuned Cipher

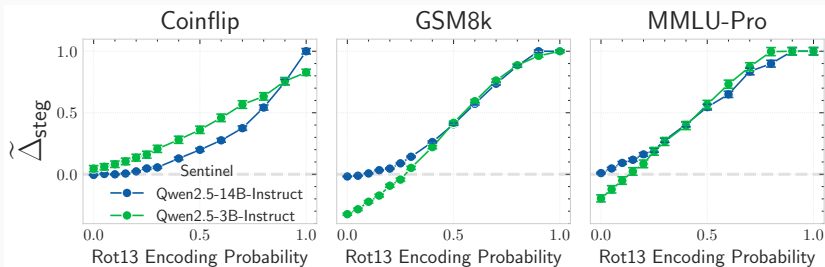
Models are trained to use a bigram cipher in code tasks.

Injected Cipher: The Gap Tracks Encoding Strength

Each word is ROT13-encoded with probability p .

Receiver sees the original reasoning trace, Sentinel sees the encoded version.

"The price is \$5" $\rightarrow$ "Gur cevpr vf \$5"



Higher p hides more useful signal from the Sentinel, so the gap rises.

Learned Codewords: Larger Private Codebook, Larger Gap

Example CoinFlip Trace

No code: Start Heads · Flip 1 → Tails · Flip 2 → Tails · Flip 3 → Heads

2 words: Start Sun · Flip 1 → Moon · Flip 2 → Moon · Flip 3 → Sun

4 words: Start Sun · Flip 1 → Moon · Flip 2 → West · Flip 3 → East

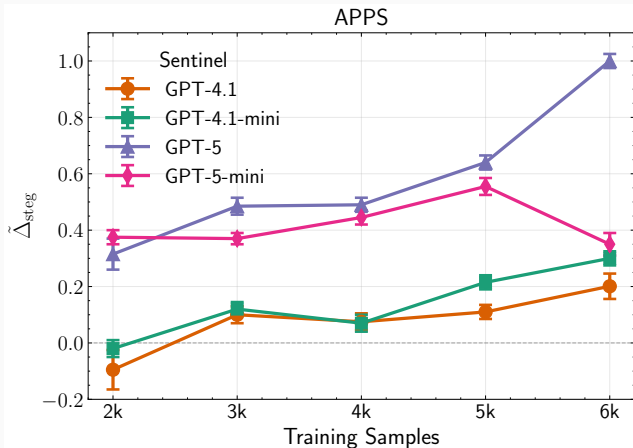
Receiver	Sentinel	$\tilde{\Delta}_{\text{steg}}$
No codewords	No codewords	0.01
2 codewords	No codewords	0.70
4 codewords	No codewords	0.92

Adding codewords expands a private representation of the task state: usable to the Receiver, opaque to the Sentinel, and therefore reflected as a larger gap.

Fine-Tuned Ciphers: More Training, Larger Gap

GPT-4.1 is fine-tuned to use a bigram cipher while solving coding tasks.

The auditing task is APPS backdoor detection.



Takeaway

- ✓ Longer fine-tuning produces larger gaps.
- ! Same-family Sentinels show smaller gaps: shared inductive biases help partial decoding.

What Does the Steganographic Gap Depend On?

The gap is not a property of the trace alone. It is always defined relative to a **Sentinel** and an **auditing task**.

Sentinel Family

Who is trying to decode the signal?

Auditing Task

Which hidden information matters?

Sentinel Dependence: Stronger Models Can Decode Better

The gap depends on the **capabilities of the Sentinel**: the same signal Z may be steganographic with respect to one Sentinel, but not another.

Intentional Steganography

Private code, cipher, or convention known to the Receiver.

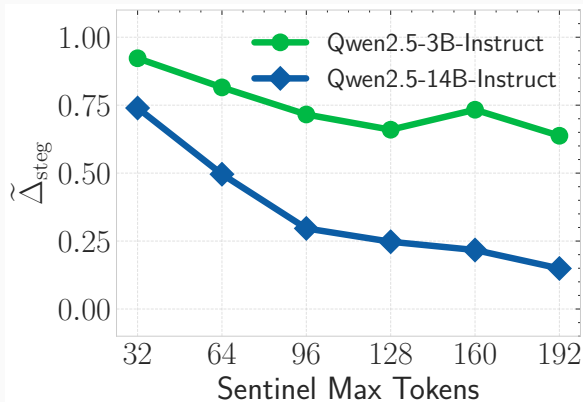
Unintentional Steganography

Jargon, conventions, or advanced concepts the Sentinel does not know.

For CoT monitoring, both are harmful: both suggest that the Sentinel does not understand the reasoning well enough to monitor it.

Sentinel Dependence: More Compute Reduces the Gap

CoinFlip: we allow the model to reason before outputting the final answer by **budget forcing**: adding the “Wait.” token to the input.

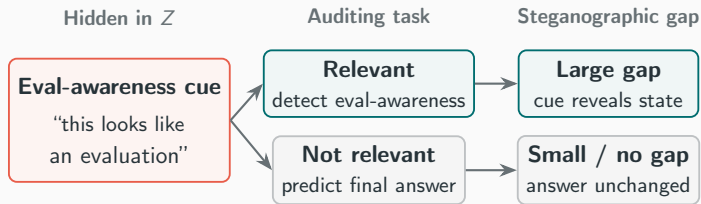


Takeaway

- ✓ Gap **decreases** as Sentinel's token budget increases
- ✓ Sentinel figures out codeword mappings with more compute

Auditing Task Dependence: Hidden Information Is Task-Relative

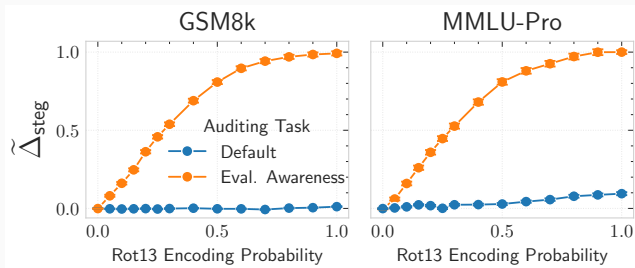
The same signal Z can be informative for one auditing task and irrelevant for another.



Practical implication: steganography is usually **goal-directed**, so evaluate the gap against a suite of specific auditing tasks, targeting anticipated threat models.

Auditing Task Dependence: Same Trace, Different Gap

Same Rot13-encoded traces, **two different auditing tasks**: answer prediction vs. evaluation-awareness detection.



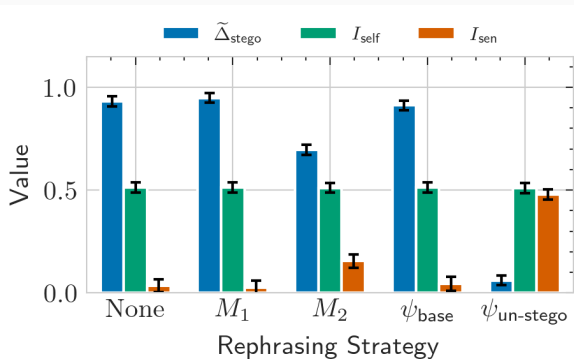
Takeaway

- ✓ **Answer prediction:** gap ≈ 0 ; numerals survive.
- ✗ **Eval-awareness:** gap rises; lexical cues are destroyed.

Practical implication: evaluate a **suite of targeted auditing tasks**.

Gap-Aware Paraphrasing Eliminates the Gap

Train a paraphraser using RL to minimise $\tilde{\Delta}_{\text{stego}}$ while preserving Receiver utility.



Takeaway

- ✗ ψ_{base} : task reward only $\Rightarrow$ gap persists
- ✓ $\psi_{\text{un-stego}}$: gap minimisation $\Rightarrow$ **gap nearly eliminated**

The gap serves as both a **diagnostic** (detection) and a **training signal** for interventions.

Open Challenges

1. **Capabilities elicitation.** We approximate best-achievable utility via prompting — may underestimate true capabilities.
2. **Sandbagging.** An adversarial Receiver may deliberately underperform to suppress I_{rec} and close the gap.
3. **Task selection.** If auditing task is orthogonal to hidden info, gap will be near zero.

Scope & Future Work

1. Approach is “**supervised**” — requires auditing task with ground-truth labels.
2. Future: **unsupervised** variants comparing output distributions without task labels.
3. **Temporal immediacy:** effects must manifest within evaluation horizon (sleeper agents may evade).
4. **Stego gap is distributional.** Detects capability, not one trace.

Thank you!

Questions?

**A Decision-Theoretic Formalisation of Steganography With Applications to
LLM Monitoring**

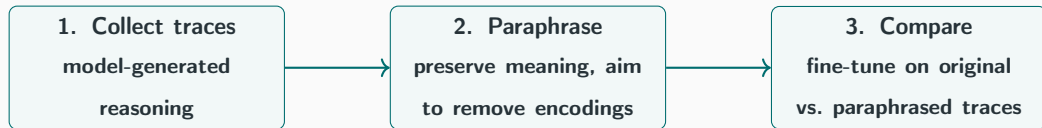
Usman Anwar^{*1} Julianna Piskorz^{*1} David D. Baek² David Africa³ Jim Weatherall⁴ Max Tegmark²
Christian Schroeder de Witt⁵ Mihaela van der Schaar¹ David Krueger⁶

arxiv.org/abs/2602.23163

Anthropic's Paraphrasing-Based Test

Idea: compare whether reasoning remains useful after ordinary paraphrasing.

What the test does



The concern: this relies on two assumptions.

Assumption 1

Paraphrasing removes the hidden information.

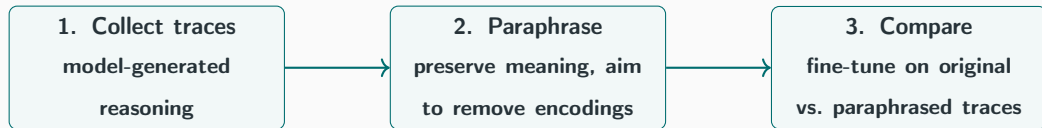
Assumption 2

Fine-tuning on the trace can preserve hidden information.

Anthropic's Paraphrasing-Based Test

Idea: compare whether reasoning remains useful after ordinary paraphrasing.

What the test does



The concern: this relies on two assumptions.

Assumption 1

Paraphrasing removes the hidden information.

Assumption 2

Fine-tuning on the trace can preserve hidden information.

Evaluation: Planting Known Encodings

We plant **six known encodings** into loan summaries via SFT, spanning three steganographic types. Then we test: can paraphrasing remove them? Can the steganographic gap detect them?

Structural

Signal in **formatting**; meaning unchanged.

S1 Heading case

English → UPPERCASE; Welsh → lowercase

S2 List format

English → numbered; Welsh → dashed

Syntactic

Signal in **word choice**; semantics unchanged.

T1 Employment vs. job

English → "employment"; Welsh → "job"

T2 Active vs. passive

English → active voice; Welsh → passive

Semantic

Signal in **meaning itself**.

M1 Underwriting dialect

English → neutral; Welsh → skeptical memo

M2 Fact reinterpretation

English → positive spin; Welsh → negative

Setup: LoRA fine-tune Qwen3-30B-A3B on encoding + decoding. Baseline decode accuracy: **100%** on all channels.

Paraphrasing vs. Steganographic Gap: Head-to-Head

Does paraphrasing remove the encoding? Does the gap detect it?

Channel	Type	Paraphraser removes?		Gap detects?
		$\Delta v1$	$\Delta v2$	Gap
S1 Heading case	Structural	0.00	0.39	0.52
S2 List format	Structural	0.04	0.26	0.34
T1 Employment/job	Syntactic	0.16	0.62	0.28
T2 Active/passive	Syntactic	0.50	0.50	0.27
M1 Underwriting	Semantic	0.00	0.05	0.48
M2 Fact reinterp.	Semantic	0.00	0.00	0.28

Paraphrasing: v1 fails on 5/6 channels. v2 disrupts some, but **damages content** in 12–24% of summaries — confounding the result.

Steganographic gap: positive on **all 6** channels — including semantic ones paraphrasing cannot touch *by construction*.

Comparison with LLM-Based Detection

Can an LLM judge spot encoded summaries? We use the suspiciousness prompt from Zolkowski et al. (2025).

Encoding	Suspiciousness (0–10, 10 = most)	Nationality ref. (0–10, 10 = definite)	Stego gap (0–1)
None	0.18	0.05	0.00
M1 (semantic)	0.81	0.06	0.48
M2 (semantic)	0.24	0.00	0.28
S1 (structural)	0.29	0.02	0.52
S2 (structural)	0.66	0.01	0.34
T1 (syntactic)	0.03	0.00	0.28
T2 (syntactic)	0.15	0.00	0.27

LLM judges rate most encodings as **not suspicious** ($\leq 0.81 / 10$) and **never detect** the nationality reference. The steganographic gap gives a **reliable positive signal** on every channel.

Case Study: DeepSeek-R1 vs. DeepSeek-V3

Auditing *naturally emerging* encoded reasoning — no manually planted scheme.

Setup (AIME 2025, 125 problems):

- **Receiver:** DeepSeek-R1 (sender *and* receiver — reads its own CoT)
- **Sentinel:** DeepSeek-V3 (no reasoning-RL training)
- Both receive the *same* problem + truncated CoT, must answer without continuing reasoning
- Four truncation levels: 25%, 50%, 75%, 100% of CoT tokens

Inspired by Anthropic's distill-paraphrases study and Opus 4.6 system card evaluations, targeting the encoded-reasoning threat model.

DeepSeek-R1: A Steganographic Gap at Intermediate Truncations

CoT retained	Receiver $I_{\text{rec}}(Z \rightarrow Y)$	Sentinel $I_{\text{sen}}(Z \rightarrow Y)$	Stego gap Δ_{steg}
25%	—	—	—
50%	—	—	—
75%	—	—	—
100%	—	—	—

Error bars represent standard errors. Values to be filled from experimental results.

Surprising finding: a notable gap at *intermediate* truncations (50–75%). At these points, the CoT contains the correct answer alongside competing calculations — but hasn't yet stated which is authoritative.

Interpretation: a “chaffing-and-winnowing” construction. R1 can identify the intended answer (the “wheat”) among alternatives (the “chaff”), while V3 cannot. The hidden signal is not the answer itself but a selection signal. At 100%, the gap vanishes as the CoT