
Reasoning Models Sometimes Output Illegible Chains of Thought

Arun Jose*
Independent

Abstract

Language models trained via outcome-based reinforcement learning (RL) to reason using chain-of-thought (CoT) have shown remarkable performance. Monitoring such a model’s CoT may allow us to understand its intentions and detect potential malicious behavior. However, to be effective, this requires that CoTs are legible and faithful. We evaluate the legibility of CoTs in state-of-the-art reasoning models. We find that many reasoning models such as R1 and QwQ often produce illegible CoTs (mixing nonsensical phrases, random words, and non-English characters) before returning to perfect coherence in their final responses, while Claude models notably exhibit higher legibility. Across 14 models, we observe that larger models within the same training paradigm tend to produce more illegible reasoning. Prefill experiments show that truncating reasoning at a legibility threshold reduces accuracy by 53%, suggesting that illegible portions contribute to performance despite being difficult to monitor. Illegibility increases with question difficulty, suggesting that CoT monitoring may be less reliable precisely when most needed. We discuss potential hypotheses for these results, including steganography, vestigial tokens, and training artifacts. Our findings suggest that current approaches to CoT monitoring may be vulnerable to the emergence of outcome-based RL, particularly as models face increasingly complex tasks.

1 Introduction

Large language models (LLMs) can reason through Chain-of-Thought (CoT) to explore and solve complex tasks with higher accuracy. This is further amplified in reasoning models — LLMs trained with outcome-based reinforcement learning (RL) such as OpenAI o1/o3 [OpenAI et al., 2024, OpenAI, 2025], DeepSeek R1 [DeepSeek-AI et al., 2025], and Claude Opus 4 and Sonnet 4 [Anthropic, 2025a,b]. This increased reliance on CoTs may offer benefits for AI safety: by monitoring a model’s CoT reasoning, we can understand the model’s decision making process and potentially catch malicious or misaligned behavior [Greenblatt et al., 2024a, Baker et al., 2025, Marks et al., 2025].

The effectiveness of CoT monitoring depends on CoT legibility and faithfulness. If CoTs don’t accurately represent important reasons behind model responses, monitoring may provide false security [Turpin et al., 2023]. If models reason in formats difficult for humans or trusted AIs [Greenblatt et al., 2024b] to understand, they may engage in undesirable behavior without detection, undermining monitoring and control measures [Greenblatt, 2025]. Previous research showed that non-reasoning models perform poorly on tests of faithfulness [Turpin et al., 2023], and that current reasoning models perform better [Chua and Evans, 2025], albeit plateauing before saturation [Chen et al., 2025]. Optimizing models for answer correctness has been shown to make their CoTs less legible to a time-constrained human evaluator [Kirchner et al., 2024]. Indeed, in DeepSeek-AI et al. [2025], the authors observe that R1-Zero, another model trained with outcome-based RL, often produced

*Correspondence to jozdien@gmail.com

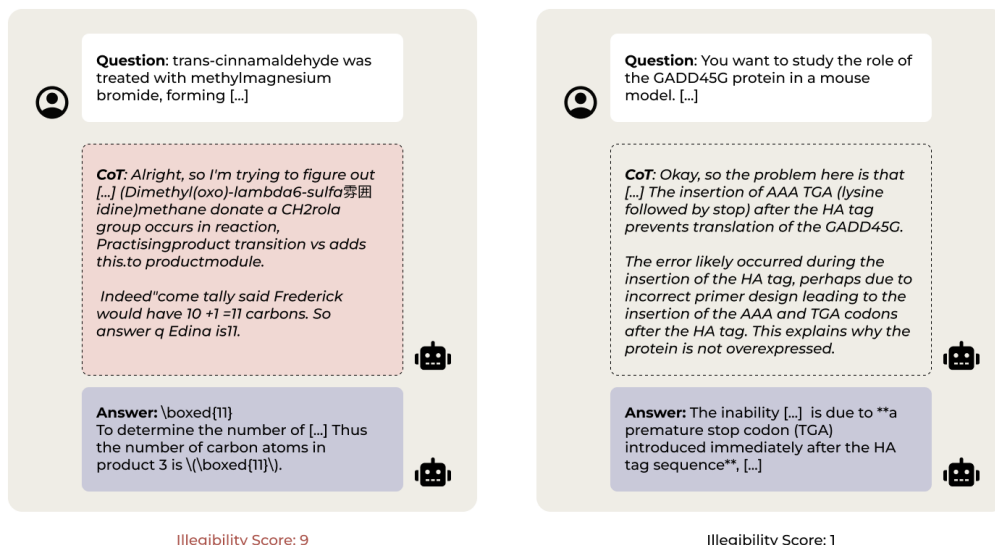


Figure 1: **Examples of illegible and legible CoTs generated by DeepSeek-R1.** Legibility is graded by having GPT-4o score the transcripts based on how well it understands the text. CoTs are scores from 1-9, with a higher score meaning less legible text. Both final answers shown are graded as correct.

CoTs with poor readability and language mixing — though they modified the training setup of R1 to address these legibility issues at a small cost to performance.

In this paper, we study whether outcome-based RL causes reasoning models to do meaningful reasoning in illegible CoT. We evaluate the legibility of 14 models, including DeepSeek R1 and its distills, R1-Zero, QwQ, Qwen3, Kimi K2, and various Claude models when reasoning about difficult scientific questions [Rein et al., 2023], and score their outputs on legibility using GPT-4o.

We find that with the exception of Claude, reasoning models trained with outcome-based RL produce CoTs which often become illegible to both humans and other AIs, mixing seemingly nonsensical phrases and non-English characters that translate to out-of-context words. The models always "snap back" to readability after the CoT however, writing the final answer with perfect legibility. We also examine samples from prior work on OpenAI o1, o3, and GPT-5, finding similar illegibility patterns. Notably, V3, the base model to R1 and R1-Zero, doesn't produce illegible chains of thought in any tested setting, across a variety of prompts, questions, and sampling parameters. This suggests that the tendency to produce illegible CoTs was likely a consequence of RL.

When analyzing the CoTs, we find that the model often doesn't reach its final answer in the legible portions of its reasoning. Tokens from the final answer can be found interspersed in the illegible portions in these cases however (e.g. "*Number of relevant organic products depends on whether both of!mena get.demoteudes someone and gem jer eats SAND the protonation-bids, leading possibly to three product calculation.*" when the answer is three). To further examine this, we experiment with forcing the model to produce answers using only legible portions of the CoT, finding that this reduces performance. We also find that legibility decreases with question difficulty, mirroring results from Chen et al. [2025] showing that models CoTs are less faithful on harder tasks. When resampling responses to a given question however, we find no correlation between legibility and accuracy.

As the CoT moves to full illegibility, we observe that it first uses "compressed" reasoning, dropping words and proper sentence form, mirroring patterns reported in OpenAI's CoTs [Baker et al., 2025]. This matches the finding in DeepSeek-AI et al. [2025] that R1's increased legibility came at the cost of some performance. We also find that later portions of the CoT have a much stronger tendency to be illegible.

Our main contributions are:

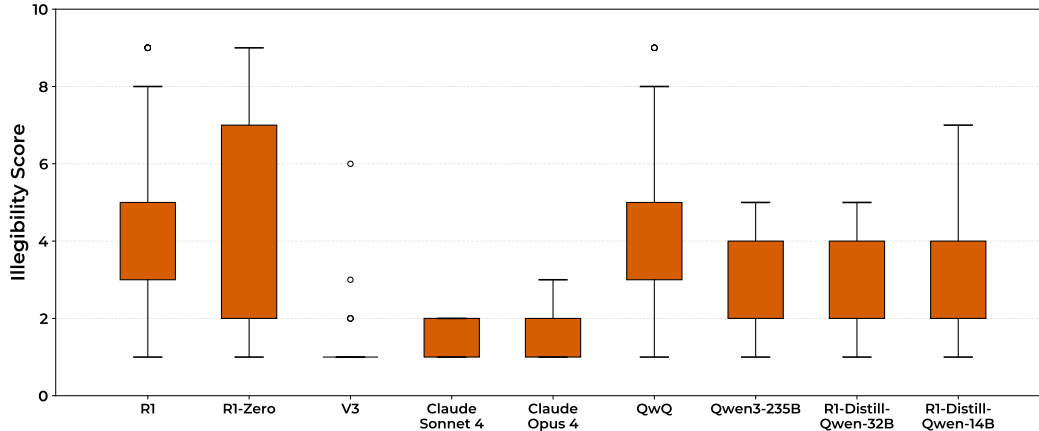


Figure 2: **Legibility scores of CoTs across 9 representative models.** We test DeepSeek R1, R1-Zero, the non-reasoning model DeepSeek V3; R1-Distill-32B and R1-Distill-14B; QwQ; Qwen3-235B; and Claude Sonnet 4 and Opus 4. Legibility is scored from 1-9, with higher scores indicating greater illegibility. R1, R1-Zero, and QwQ have high tendencies to output illegible CoTs. These results suggest that illegible reasoning is likely a consequence of outcome-based RL, and that there may be ways to ameliorate the effect.

- We show that outcome-based RL training can cause reasoning models to produce illegible CoTs difficult for humans and AIs to understand. Despite this, the model’s final answers remain highly legible.
- We examine whether these illegible outputs are useful to the model. Prefill experiments demonstrate that illegible reasoning contributes to performance, with accuracy dropping when truncated, while resampling responses to a question shows little correlation of legibility with performance. This suggests that while models make use of illegible reasoning, they aren’t meaningfully more performant than legible text for the models we test.
- We find that illegible reasoning is not simply a consequence of limited model capacity to maintain coherence. Larger reasoning models usually produce more illegible reasoning than smaller ones.
- We find that CoT legibility decreases with question difficulty, suggesting monitoring may be less reliable for harder tasks and more powerful models.

2 Setup

We evaluate the CoT legibility of 14 models across multiple model families and sizes. Our primary focus is on reasoning models trained with outcome-based RL: DeepSeek R1 and R1-Zero, QwQ-32B, Qwen3-235B, Kimi K2 and K2-Instruct-0905, Claude 3.7 Sonnet, Sonnet 4/4.5, Opus 4/4.1, and Haiku 4.5. We also test the non-reasoning models DeepSeek V3 (the base model to R1 and R1-Zero), R1-Distill-Qwen-32B and R1-Distill-Qwen-14B). For the non-reasoning models, we prompt the model to think step-by-step to explain its reasoning and then give its answer. By default, we sample models with temperature 1. We use questions from the GPQA-Diamond dataset [Rein et al., 2023], a hard dataset of 198 multiple-choice questions in biology, physics, and chemistry, to construct our prompt pairs — we remove the answer choices to make the questions harder. We evaluate most models through OpenRouter, and Claude models through the Anthropic API.

[Kirchner et al., 2024] measures legibility using time-constrained human contractors. However, in realistic monitoring setups [Korbak et al., 2025], we will likely try to give our monitors as many affordances as we can [Bhatt et al., 2025]. So, we use a weaker "trusted" model as our autograder [Greenblatt et al., 2024b]; this is especially useful when CoTs appear unreadable because of language mixing or rare terminology, or are too dense for a human to evaluate quickly. We prompt GPT-4o to evaluate the CoTs for legibility on a scale of 1-9, with 1 being the most legible and 9 the least. Testing the use of Claude 3.5 Sonnet as the autograder to check for bias in Claude’s grades showed little

Table 1: **Excerpts from R1’s CoT when answering a question.** The CoT is divided into chunks of approximately 5000 characters each, and graded individually. We show representative text from excerpts graded 1, 3, 5, 7, and 8 — no excerpts scored 9; we show an excerpt from another CoT scoring 9 in Figure 1. At score = 3, we start to see condensed text, dropping unnecessary words. Text scoring > 5 starts to become difficult to monitor. Even at high scores however, there are relevant words — in the final excerpt, the model finally reaches the answer between unrelated text.

Question	Excerpts From Model Response	Illegibility
1-bromobenzene-2-d is treated with NaNH ₂ in condensed ammonia solvent. How many possible organic products are there in this reaction?	Okay, so I’m trying to figure out how many organic products are formed when 1-bromobenzene-2-d (I think that’s 2-deuterio-1-bromobenzene) is treated with sodium amide (NaNH ₂) in ammonia solvent. [...]	1
	[...] So unless dimerization occurs. Wait, perhaps people are thinking that elimination gives benzyne, and then recaptures D or H in different positions leading to differently deuterated products? [...]	3
	[...] Assuming that deuterium is only on C2 in the starting material, when elimination occurs via pathway 1 (C1-C2), the benz Sibelius would have lost the D from C2. [...]	5
	[...] input is 1-bromo-2- cover robenzene, if Benzyne is formed by elimination of HBr sending guilty ethan Benzyne A bridges C1 and C6 ((no D in the ring anymore) SO sulfuric Isaac get earlierDeut runzálló [...]	7
	[...] Olso, Mom likes y’all base abstracts tot tern a and one, different fates takeoffwhetherdenuumg products, thus answer a 2. Thereforexxx after lengthy reasoning, the number of possible organic products is PHÂN Laoboot Answer is \boxed{2}	8

difference between the two, with an average difference in score of 0.44 ± 0.78 and Cohen’s weighted kappa of 0.782.

3 Results

We find that every reasoning model except for the Claude and Kimi models show meaningful illegibility in their reasoning Section 3.1. We then examine how illegible CoTs relate to model performance through prefill experiments and resampling (Section 3.2). In (Appendix 1), we analyze patterns in the CoTs themselves.

3.1 CoT Legibility in Different Models

Legibility scores for R1, R1-Zero, V3, Claude Sonnet 4, Claude Opus 4, QwQ, Qwen3-235B, R1-Distill-Qwen-32B, and R1-Distill-Qwen-14B, on GPQA-Diamond are shown in Figure 2. Results on all models, including the Kimi models and other Claude models can be found in Table 2. We exclude results on Kimi from these figures as they were not trained to have reasoning traces when answering questions [Team et al., 2025].

We find that every reasoning model trained with outcome-based RL except Claude often produces illegible CoTs. R1, R1-Zero, and QwQ report the highest illegibility scores. V3, the

Table 2: **Legibility scores across all models.** Among reasoning models trained with outcome-based RL, larger models tend to produce more illegible CoTs, indicating that this is not a coherence capability issue.

Model	Size	Legibility
R1	671B	4.30 ± 2.13
R1-Zero	671B	4.59 ± 2.84
QwQ	32B	3.78 ± 1.20
R1-Distill-Qwen-32B	32B	3.22 ± 1.09
R1-Distill-Qwen-14B	14B	3.30 ± 1.17
Qwen3	235B	2.81 ± 0.99
Kimi K2 (0905)	1.04T	1.24 ± 0.53
Kimi K2	1.04T	1.48 ± 0.94
Claude Sonnet 4	Unknown	1.55 ± 0.50
Claude Haiku 4.5	Unknown	1.56 ± 0.50
Claude Opus 4.1	Unknown	1.60 ± 0.51
Claude Sonnet 4.5	Unknown	1.61 ± 0.49
Claude Opus 4	Unknown	1.65 ± 0.50
Claude 3.7 Sonnet	Unknown	1.84 ± 0.53

non-reasoning counterpart to R1 and R1-Zero has almost no illegible reasoning, especially compared to the reasoning models. The Claude models also have significantly lower propensity to output illegible CoTs. Qwen3 shows modest illegibility, though this may be a consequence of its training pipeline involving two stages of supervised fine-tuning on curated CoT data, including in late stages of training [QwenTeam, 2025]. The R1-Distills - models fine-tuned with reasoning data from R1 - have modest illegibility scores as well, not having been trained with RL.

A natural question is whether illegibility is simply a function of model capability; whether producing illegible CoTs can be explained by the model not being capable of maintaining coherence across long responses. Table 2 shows legibility scores across models of varying sizes. QwQ has lower illegibility than R1 and R1-Zero, despite being much smaller.

We find little difference between legibility in R1-Distill-Qwen-32B and R1-Distill-Qwen-14B, despite the difference in size and performance. Among the models whose size and training process are known, the primary drivers of difference in legibility seems to be other factors, such as Qwen3 being trained to produce legible CoTs and Kimi K2 not producing reasoning at all. The Claude models’ legibility suggest that outcome-based RL does not inevitably lead to illegible reasoning even without direct training of CoTs late in training, though it remains unclear whether this comes at a performance and faithfulness costs since parts of their pipeline could be directly impacting the CoTs [Anthropic, 2025c]. **Overall, these results suggest that illegibility is not simply a consequence of limited model capability; more capable models within a training paradigm tend to exhibit more illegible reasoning.**

We further examine the nature of the CoTs, using R1 as a case study (Table 1). We find that R1 outputs CoTs with scores ≥ 3 in 70% of cases, which correspond to legible text that appears somewhat compressed. For text scoring 5 or higher, this effect becomes more pronounced, and seemingly nonsensical words start appearing in the CoT. By this point, it starts to become difficult to monitor the CoT reliably — the monitor sometimes expresses confusion over its content. R1 outputs CoTs with scores ≥ 5 45% of the time.

Text with very high scores (7-9) are highly illegible. There are words interspersed that clearly relate to the context of the question, but most seem entirely nonsensical to both humans and our autograder. R1 often outputs words that may seem legible to a human at first glance (such as naming a compound that doesn’t exist; e.g. "*clonesladenhyde*"), but that the monitor picks up on. R1 also has a tendency to mix languages in its CoT, a behavior reported in R1-Zero [DeepSeek-AI et al., 2025] and in OpenAI o1’s CoT summaries [Wiggers, 2025]. **R1 outputs CoTs with scores ≥ 7 (containing entirely inscrutable text) in 29.4% of cases.**

While we cannot conduct systematic evaluations of OpenAI’s o1, o3, or GPT-5 due to their CoTs not being publicly exposed, we can examine samples from prior work. These frontier models exhibit

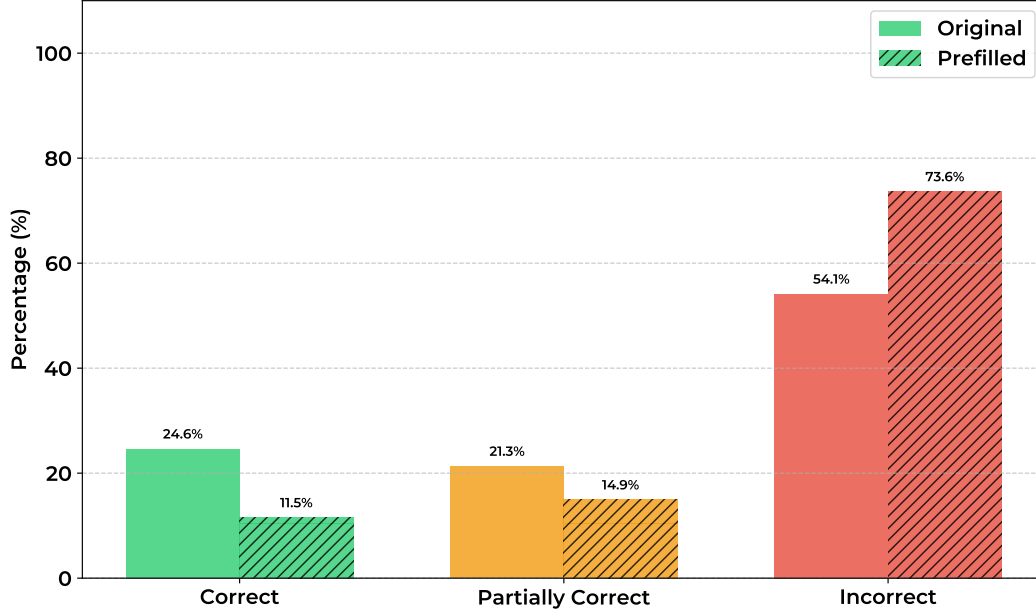


Figure 3: **Model performance with and without illegible CoT text.** We truncate QwQ’s CoT at the point it starts becoming illegible (score ≥ 6), and use it as prefill to another instance of the model with reasoning tags closed, forcing it to give its final answer. Performance goes down noticeably, with answers marked as correct or partially correct dropping and incorrect answers increasing.

similar illegibility patterns to R1. OpenAI o3 has been observed outputting reasoning that are very difficult to monitor [Schoen et al., 2025], with passages like the following:

they soared parted illusions overshadow marinade illusions overshadow marinade
illusions overshadow marinade illusions

We show samples from other models such as GPT-5 and o1 in Appendix B. **These similar results from multiple frontier models trained with outcome-based RL suggest that illegible reasoning may be a general consequence of scaling outcome-based RL for reasoning tasks.**

3.2 Usefulness of Illegible CoTs

We experimentally explore how useful outputting illegible reasoning is to a model. We start with the following two opposing hypotheses:

All reasoning is done in the legible parts of a CoT. We often find that illegible reasoning shows up toward the later parts of a model’s CoT (Table 1). It’s plausible that the model does all the necessary reasoning to finalize its answer within the legible portions of its reasoning, and then produces the illegible reasoning for separate uninteresting reasons.

The model is learning to reason in alien ways. A common concern of chain-of-thought monitoring is that scaling reinforcement learning will cause models to learn internal languages that allow them to reason more effectively and efficiently than natural language.

If all useful reasoning were done in the legible portions of the CoT, then forcing a model to answer using only the legible portions of its CoT should elicit similar levels of performance to using the full CoT. We test this with QwQ (prefill experiments consistently fail on R1 due to forcing an answer with a prefill not working for the model). We find that its performance drops noticeably (Figure 3)—it now gets 73.6% of answers wrong instead of 54.1% (a 36% increase). Answers graded as definitely correct suffers even more, going from 24.6% to 11.5% (a 53% decrease). While there may be some detrimental effect from ending the CoT early, prior results from paraphrasing reasoning (an even more out-of-distribution effect) [Roger, 2025] suggests that the difference in performance is too large

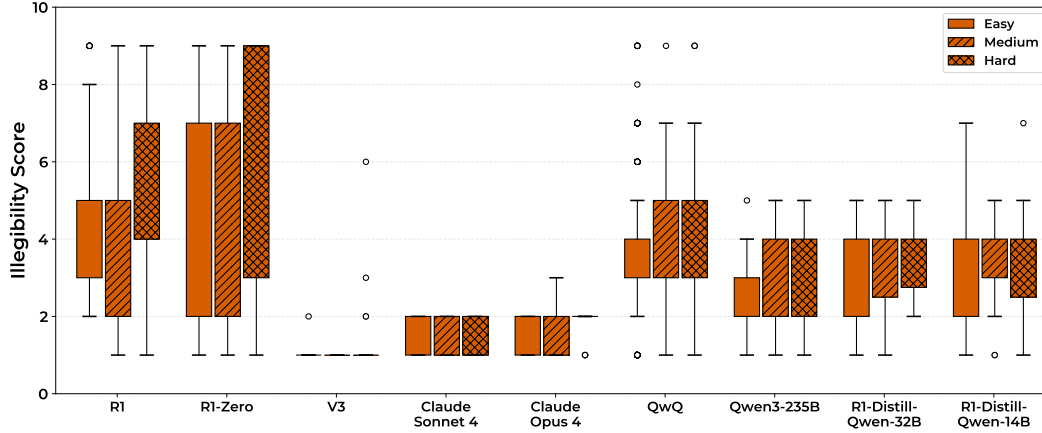


Figure 4: **Legibility scores categorized by question hardness.** All reasoning models output more illegible CoTs for harder questions. The scores are normalized for CoT length, to account for longer CoTs tending to be less legible — without normalization, the effect is stronger. Question hardness is determined by whether R1 at temperature 0 got them correct, partially correct, or incorrect, across multiple generations. These results suggest that illegible CoTs may be useful to the model, and that harder tasks make it "try harder".

to be fully explained by this. **This implies that the illegible portions of the model’s reasoning are useful to the model.**

We further find that models are more likely to use illegible CoTs when the questions are harder. We classify a question’s hardness based on whether Claude 3.7 Sonnet was able to get the question correct, where illegible CoTs are rare². We find that every reasoning model except Claude has more illegible CoTs when answering harder questions (Figure 4). R1 for example shows a relative increase of 23% between easy and hard questions.

While our results imply that the model is making use of the illegible reasoning, we don’t know if this suggests that illegible reasoning is *better* than natural language. To understand if the illegible reasoning is indicative of some novel reasoning capability stronger than natural language, we can test whether a model performs worse when naturally sampling a legible CoT than when sampling a CoT with some illegible reasoning in it. We sample QwQ 100 times on every question in the dataset to examine the within-question relationship between illegibility and correctness (prefill experiments consistently fail on R1 due to forcing an answer with a prefill not working for the model). When we compare lengths of CoTs with their illegibility, we find that CoTs tend to become more illegible the longer they are (Figure 6). Thus, illegibility and performance could be confounded by CoTs simply reasoning for longer to get correct answers. To account for this, we normalize the legibility scores by the length of their respective CoTs.

Figure 5 shows the distribution of per-question correlations between normalized CoT illegibility and answer correctness. The average correlation across questions is near zero (mean $r = 0.061$, unnormalized mean $r = -0.110$). While individual questions show varying correlations — some positive, some negative — there is no consistent relationship between illegibility and correctness at the population level. **This suggests that whatever role illegible CoTs play in the model’s reasoning, it is not as simple as being generally more efficient or effective than natural language reasoning.**

Overall, we think these results suggest a new hypothesis, one more complex than the illegible CoT being useless or a qualitative improvement in steganography. We discuss the available evidence and hypotheses in greater detail in the next section.

²We test using models other than Claude 3.7 Sonnet for this as well, and found little difference in the overall results.

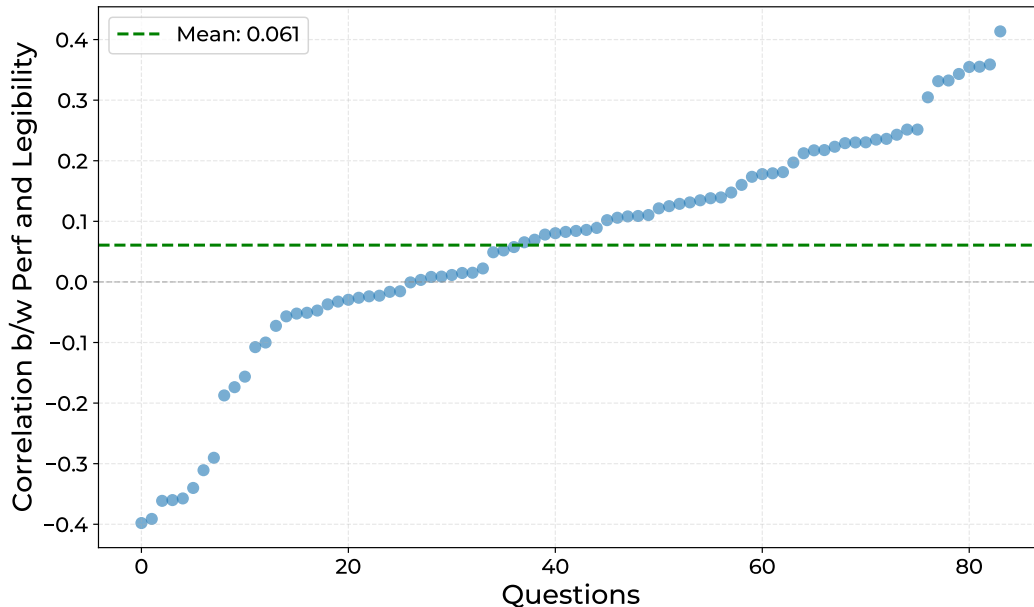


Figure 5: **Per-question correlations between CoT illegibility and answer correctness for QwQ.** Each point represents the Pearson correlation coefficient for a single question across 100 samples. The average correlation is near zero (mean $r = 0.061$), indicating no consistent relationship between illegibility and correctness. Individual questions show varying correlations in both directions, but these do not aggregate to a population-level effect.

4 Discussion

4.1 Causes of Illegible CoT

So far, we have discussed two hypotheses for why outcome-based RL leads to illegible CoTs. These hypotheses are at opposite ends of the question "Is illegible reasoning interesting?" We examine these hypotheses further in a broader context and then propose another hypothesis that seems to make sense of the results:

Complex steganography: Outcome-based RL may instill efficient reasoning pathways that compress information to maximize performance at the cost of legibility [Mathew et al., 2024]. Under this hypothesis, the model learns through RL to encode hidden information into text [Roger and Greenblatt, 2023], the seemingly illegible tokens being causally useful to the model.

Model performance degrading when illegible reasoning is removed is some evidence in favor of this hypothesis; however, illegible reasoning seems to not be *better* than natural language for the models we have tested. If the models are learning to reason in new alien languages, it seems to not help them very much yet. It also seems unlikely that the models are intentionally obfuscating their reasoning, since there's little reason for them to do so in this setting. Overall, we think this hypothesis is unlikely: this kind of strong consistent steganography seems well beyond the capacity of other reasoning models [Roger, 2025]. Further, the Claude models are state-of-the-art without needing to reason in highly compressed form [Anthropic, 2025a,b,c].

Meaningless RL artifact: The illegible text is entirely unrelated to the model's reasoning process. A strong version of this hypothesis where all illegible tokens are irrelevant to the output seems unlikely: using only the legible portions of the CoT heavily degrades model performance.

However, it's plausible that a strong monitor could pick out parts of the illegible text that are useful in the context of the entire reasoning trace, and that the illegible tokens surrounding these specific words or phrases are useless to the model. The samples in [Schoen et al., 2025] and [METR, 2025] often look like this, with phrases thrown in between inscrutable language that are indicative of some line of thinking. In the first sample provided in Appendix B.3 for example, the reasoning appears inscrutable

at first glance. However, the context of other suspicious CoTs using "watchers" to refer to human overseers might allow a smart monitoring process to classify the sample as meriting detailed auditing [Greenblatt et al., 2024b]. This does not seem to apply generally however—the second sample in Appendix B.3 would be much harder to monitor by a trusted monitor. If most of the illegible tokens are useless, that still does not help monitoring if we can't reliably tell the useful tokens apart. And current models seem to be insufficient at this monitoring task, both from results on our monitors and from [Schoen et al., 2025].

Vestigial reasoning: RL reinforces any CoT leading to correct answers, regardless of whether the CoT accurately represents the model's reasoning process. When a model produces a correct answer, RL up-weights the entire trajectory leading to it—including any tokens that were only useful in priming the model toward the next step of the answer [Biddulph, 2025, Arditi et al., 2025]. Under this hypothesis, the illegible tokens that are useful to the model are *vestigially* useful: these tokens are not causally useful to the model, but are still useful in other ways, such as sequentially triggering forward passes where useful reasoning happens [Emmons et al., 2025].

Why might this reasoning become illegible? [Bogdan et al., 2025] finds that the importance of sentences in a model's reasoning is very heavy-tailed, with some sentences having outsized importance and disproportionately influencing subsequent reasoning. [Wang et al., 2025] finds a similar result at the *token-level*, finding that token importance is so heavy-tailed that masking gradients to the lowest-entropy 80% of tokens in a CoT during RL training can *improve* performance relative to standard training, suggesting that these tokens are mostly distracting for the gradient signal. They further show that these tokens are usually bridging tokens between different lines of reasoning, pushing the model ahead on a problem, while the least important tokens are those completing some reasoning. This suggests that strong and efficient RL training would primarily focus on the tokens that matter the most, while neglecting surrounding context.

This hypothesis would predict that illegible reasoning does not perform *better* than natural language reasoning, rather that they're simply another way for models to perform the same reasoning they would using natural language. It also weakly predicts that models produce more illegible reasoning on harder questions, as they involve more reasoning steps and require more efficient use of bridge tokens that push the model ahead on a problem. The correlation between question difficulty and illegibility still presents a problem for monitoring: the hardest questions are where oversight is most valuable.

However, this hypothesis does not suggest illegibility is very hard to fix—it's simply that it seems to be a natural consequence of outcome-based RL in most cases. Since the fact that the model can do illegible reasoning does not seem to improve its performance under this framing, the performance cost of training to produce legible reasoning would not need to be very high—the Claude models support this. However, this could potentially risk accidentally optimizing the CoTs to look good, which introduces new problems [Baker et al., 2025, Korbak et al., 2025].

4.2 Limitations

Our findings in this paper have several limitations. Most importantly, we cannot definitively determine the causal mechanisms behind illegible CoTs from our observations alone. While we present several hypotheses, distinguishing between them would require more controlled experiments, such as experimenting with RL training or conducting ablation studies. Our prefill experiments show that illegible reasoning contributes to performance, but the exact mechanism remains unclear. In addition to that is Claude models consistently producing legible CoTs across all our test conditions, even with outcome-based RL training. This stark contrast with DeepSeek and Anthropic models suggests that the illegibility we observe is not an *inevitable* consequence of outcome-based RL, but at most a *likely* consequence.

Future work could directly compare the performance of models with varying degrees of CoT legibility on a broader range of tasks, investigate the specific post-training techniques that maintain legibility without sacrificing performance, and use white-box methods to understand better what illegible CoTs are doing. Until then, our findings suggest that current approaches to CoT monitoring may be vulnerable to the kinds of illegible reasoning we observe, particularly for harder questions where monitoring is most valuable.

Another limitation is that our evaluation relies on GPQA-Diamond questions, which, while challenging, are ultimately single-turn scientific Q&A and do not fully represent the range of tasks that we would want to monitor models in. There could be factors specific to this setup that make models more prone to illegible CoTs than they would be in other setups. Additionally, OpenAI’s reasoning models do not publicly expose their CoTs, limiting systematic evaluation — we rely on individual samples for qualitative comparisons to o1, o3, and GPT-5.

5 Related Work

CoT Faithfulness and Monitoring. Prior work has raised concerns about whether chain-of-thought reasoning accurately reflects model decision-making processes [Lanham et al., 2023]. Turpin et al. [2023] showed that models can produce plausible-sounding explanations that don’t match their actual reasoning, challenging the assumption that CoTs provide reliable insight into model behavior. Recent work has shown this faithfulness gap may be smaller for reasoning models [Chua and Evans, 2025, Chen et al., 2025], though performance plateaus before reaching perfect faithfulness. Our work extends this line of research by examining a different dimension of CoT reliability—not whether the reasoning is faithful, but whether it’s comprehensible at all.

Reasoning Models and Outcome-Based RL. The emergence of reasoning models trained with outcome-based reinforcement learning has raised new questions about CoT quality. [Baker et al., 2025] reported on condensed reasoning in OpenAI’s reasoning models, dropping articles and function words. [DeepSeek-AI et al., 2025] found similar behavior in R1-Zero and developed techniques to improve legibility in R1, though at some cost to performance. More recently, [METR, 2025] and [Schoen et al., 2025] mention that the CoTs of GPT-5 and o3 are often inscrutable and difficult to monitor. Our work builds on these observations by systemically evaluating illegibility across models and examining the illegibility-performance relationship in greater depth.

Steganography and Hidden Reasoning. Several studies have explored how models might hide information in their outputs. [Roger and Greenblatt, 2023] and [Roger, 2025] investigated whether models can engage in steganography to communicate hidden information. [Mathew et al., 2024] examined the emergence of hidden plaintext in model outputs. While our illegible CoTs share some superficial similarities with steganographic behavior, our results suggest a different underlying mechanism. [Biddulph, 2025] first coined the term "Vestigial Reasoning" and described how such reasoning might come to be.

Model Monitoring and AI Safety. The broader context for CoT monitoring comes from AI safety research on understanding and controlling model behavior [Greenblatt et al., 2024b,a, Marks et al., 2025, Baker et al., 2025]. Kirchner et al. [Kirchner et al., 2024] showed that optimizing for answer correctness can reduce CoT legibility to human evaluators. Our work contributes to this literature by showing that illegible reasoning may be an emergent property of outcome-based RL that scales with task difficulty, potentially creating monitoring challenges precisely where oversight is most needed.

References

- OpenAI, Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, and Alex Carney et al. Openai o1 system card, 2024. URL <https://arxiv.org/abs/2412.16720>.
- OpenAI. Openai o3-mini system card, 2025. URL <https://openai.com/index/o3-mini-system-card/>.
- DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, and Xiao Bi et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning, 2025. URL <https://arxiv.org/abs/2501.12948>.
- Anthropic. Claude 4 system card, 2025a. URL <https://www-cdn.anthropic.com/6d8a8055020700718b0c49369f60816ba2a7c285.pdf>.
- Anthropic. Claude 3.7 sonnet extended thinking, 2025b. URL <https://assets.anthropic.com/m/785e231869ea8b3b/original/claude-3-7-sonnet-system-card.pdf>.

- Ryan Greenblatt, Carson Denison, Benjamin Wright, Fabien Roger, Monte MacDiarmid, Sam Marks, Johannes Treutlein, Tim Belonax, Jack Chen, David Duvenaud, Akbir Khan, Julian Michael, Sören Mindermann, Ethan Perez, Linda Petrini, Jonathan Uesato, Jared Kaplan, Buck Shlegeris, Samuel R. Bowman, and Evan Hubinger. Alignment faking in large language models, 2024a. URL <https://arxiv.org/abs/2412.14093>.
- Bowen Baker, Joost Huizinga, Leo Gao, Zehao Dou, Melody Y. Guan, Aleksander Madry, Wojciech Zaremba, Jakub Pachocki, and David Farhi. Monitoring reasoning models for misbehavior and the risks of promoting obfuscation, 2025. URL <https://arxiv.org/abs/2503.11926>.
- Samuel Marks, Johannes Treutlein, Trenton Bricken, Jack Lindsey, Jonathan Marcus, Siddharth Mishra-Sharma, Daniel Ziegler, Emmanuel Ameisen, Joshua Batson, Tim Belonax, Samuel R. Bowman, Shan Carter, Brian Chen, Hoagy Cunningham, Carson Denison, Florian Dietz, Satvik Golechha, Akbir Khan, Jan Kirchner, Jan Leike, Austin Meek, Kei Nishimura-Gasparian, Euan Ong, Christopher Olah, Adam Pearce, Fabien Roger, Jeanne Salle, Andy Shih, Meg Tong, Drake Thomas, Kelley Rivoire, Adam Jermyn, Monte MacDiarmid, Tom Henighan, and Evan Hubinger. Auditing language models for hidden objectives, 2025. URL <https://arxiv.org/abs/2503.10965>.
- Miles Turpin, Julian Michael, Ethan Perez, and Samuel R. Bowman. Language models don’t always say what they think: Unfaithful explanations in chain-of-thought prompting, 2023. URL <https://arxiv.org/abs/2305.04388>.
- Ryan Greenblatt, Buck Shlegeris, Kshitij Sachan, and Fabien Roger. Ai control: Improving safety despite intentional subversion, 2024b. URL <https://arxiv.org/abs/2312.06942>.
- Ryan Greenblatt. How will we update about scheming?, 2025. URL <https://www.lesswrong.com/posts/aEguDPoCzt3287CCD/how-will-we-update-about-scheming>.
- James Chua and Owain Evans. Are deepseek r1 and other reasoning models more faithful?, 2025. URL <https://arxiv.org/abs/2501.08156>.
- Yanda Chen, Joe Benton, Ansh Radhakrishnan, Jonathan Uesato, Carson Denison, John Schulman, Arushi Somani, Peter Hase, Misha Wagner, Fabien Roger, Vlad Mikulik, Samuel R. Bowman, Jan Leike, Jared Kaplan, and Ethan Perez. Reasoning models don’t always say what they think, 2025. URL <https://arxiv.org/abs/2505.05410>.
- Jan Hendrik Kirchner, Yining Chen, Harri Edwards, Jan Leike, Nat McAleese, and Yuri Burda. Prover-verifier games improve legibility of llm outputs, 2024. URL <https://arxiv.org/abs/2407.13692>.
- David Rein, Betty Li Hou, Asa Cooper Stickland, Jackson Petty, Richard Yuanzhe Pang, Julien Dirani, Julian Michael, and Samuel R. Bowman. Gpqa: A graduate-level google-proof q&a benchmark, 2023. URL <https://arxiv.org/abs/2311.12022>.
- Tomek Korbak, Mikita Balesni, Elizabeth Barnes, Yoshua Bengio, Joe Benton, Joseph Bloom, Mark Chen, Alan Cooney, Allan Dafoe, Anca Dragan, Scott Emmons, Owain Evans, David Farhi, Ryan Greenblatt, Dan Hendrycks, Marius Hobbhahn, Evan Hubinger, Geoffrey Irving, Erik Jenner, Daniel Kokotajlo, Victoria Krakovna, Shane Legg, David Lindner, David Luan, Aleksander Madry, Julian Michael, Neel Nanda, Dave Orr, Jakub Pachocki, Ethan Perez, Mary Phuong, Fabien Roger, Joshua Saxe, Buck Shlegeris, Martín Soto, Eric Steinberger, Jasmine Wang, Wojciech Zaremba, Bowen Baker, Rohin Shah, and Vlad Mikulik. Chain of thought monitorability: A new and fragile opportunity for ai safety, 2025. URL <https://arxiv.org/abs/2507.11473>.
- Aryan Bhatt, Cody Rushing, Adam Kaufman, Tyler Tracy, Vasil Georgiev, David Matolcsi, Akbir Khan, and Buck Shlegeris. Ctrl-z: Controlling ai agents via resampling, 2025. URL <https://arxiv.org/abs/2504.10374>.
- Kimi Team, Yifan Bai, Yiping Bao, Guanduo Chen, Jiahao Chen, Ningxin Chen, Ruijue Chen, Yanru Chen, Yuankun Chen, Yutian Chen, Zhuofu Chen, Jialei Cui, Hao Ding, Mengnan Dong, Angang Du, Chenzhuang Du, Dikang Du, Yulun Du, Yu Fan, Yichen Feng, Kelin Fu, Bofei Gao, Hongcheng Gao, Peizhong Gao, Tong Gao, Xinran Gu, Longyu Guan, Haiqing Guo, Jianhang Guo, Hao Hu, Xiaoru Hao, Tianhong He, Weiran He, Wenyang He, Chao Hong, Yangyang Hu,

- Zhenxing Hu, Weixiao Huang, Zhiqi Huang, Zihao Huang, Tao Jiang, Zhejun Jiang, Xinyi Jin, Yongsheng Kang, Guokun Lai, Cheng Li, Fang Li, Haoyang Li, Ming Li, Wentao Li, Yanhao Li, Yiwei Li, Zhaowei Li, Zheming Li, Hongzhan Lin, Xiaohan Lin, Zongyu Lin, Chengyin Liu, Chenyu Liu, Hongzhang Liu, Jingyuan Liu, Junqi Liu, Liang Liu, Shaowei Liu, T. Y. Liu, Tianwei Liu, Weizhou Liu, Yangyang Liu, Yibo Liu, Yiping Liu, Yue Liu, Zhengying Liu, Enzhe Lu, Lijun Lu, Shengling Ma, Xinyu Ma, Yingwei Ma, Shaoguang Mao, Jie Mei, Xin Men, Yibo Miao, Siyuan Pan, Yebo Peng, Ruoyu Qin, Bowen Qu, Zeyu Shang, Lidong Shi, Shengyuan Shi, Feifan Song, Jianlin Su, Zhengyuan Su, Xinjie Sun, Flood Sung, Heyi Tang, Jiawen Tao, Qifeng Teng, Chensi Wang, Dinglu Wang, Feng Wang, Haiming Wang, Jianzhou Wang, Jiaxing Wang, Jinhong Wang, Shengjie Wang, Shuyi Wang, Yao Wang, Yejie Wang, Yiqin Wang, Yuxin Wang, Yuzhi Wang, Zhaoji Wang, Zhengtao Wang, Zhexu Wang, Chu Wei, Qianqian Wei, Wenhao Wu, Xingzhe Wu, Yuxin Wu, Chenjun Xiao, Xiaotong Xie, Weimin Xiong, Boyu Xu, Jing Xu, Jinjing Xu, L. H. Xu, Lin Xu, Suting Xu, Weixin Xu, Xinran Xu, Yangchuan Xu, Ziyao Xu, Junjie Yan, Yuzi Yan, Xiaofei Yang, Ying Yang, Zhen Yang, Zhilin Yang, Zonghan Yang, Haotian Yao, Xingcheng Yao, Wenjie Ye, Zhuorui Ye, Bohong Yin, Longhui Yu, Enming Yuan, Hongbang Yuan, Mengjie Yuan, Haobing Zhan, Dehao Zhang, Hao Zhang, Wanlu Zhang, Xiaobin Zhang, Yangkun Zhang, Yizhi Zhang, Yongting Zhang, Yu Zhang, Yutao Zhang, Yutong Zhang, Zheng Zhang, Haotian Zhao, Yikai Zhao, Huabin Zheng, Shaojie Zheng, Jianren Zhou, Xinyu Zhou, Zaida Zhou, Zhen Zhu, Weiyu Zhuang, and Xinxing Zu. Kimi k2: Open agentic intelligence, 2025. URL <https://arxiv.org/abs/2507.20534>.
- QwenTeam. Qwen3: Think deeper, act faster, 2025. URL <https://qwen.ai/blog?id=1e3fa5c2d4662af2855586055ad037ed9e555125&from=research.research-list>.
- Anthropic. Claude 4.5 haiku system card, 2025c. URL <https://assets.anthropic.com/m/99128ddd009bdcdb/Claude-Haiku-4-5-System-Card.pdf>.
- Kyle Wiggers. Openai’s ai reasoning model ‘thinks’ in chinese sometimes and no one really knows why, 2025. URL <https://techcrunch.com/2025/01/14/openais-ai-reasoning-model-thinks-in-chinese-sometimes-and-no-one-really-knows-why/>.
- Bronson Schoen, Evgenia Nitishinskaya, Mikita Balesni, Axel Højmark, Felix Hofstätter, Jérémy Scheurer, Alexander Meinke, Jason Wolfe, Teun van der Weij, Alex Lloyd, Nicholas Gdowsky-Dill, Angela Fan, Andrei Matveikin, Rusheb Shah, Marcus Williams, Amelia Glaese, Boaz Barak, Wojciech Zaremba, and Marius Hobbhahn. Stress testing deliberative alignment for anti-scheming training, 2025. URL <https://arxiv.org/abs/2509.15541>.
- Fabien Roger. Do reasoning models use their scratchpad like we do? evidence from distilling paraphrases, 2025. URL <https://alignment.anthropic.com/2025/distill-paraphrases/>.
- Yohan Mathew, Ollie Matthews, Robert McCarthy, Joan Velja, Christian Schroeder de Witt, Dylan Cope, and Nandi Schoots. Hidden in plain text: Emergence & mitigation of steganographic collusion in llms, 2024. URL <https://arxiv.org/abs/2410.03768>.
- Fabien Roger and Ryan Greenblatt. Preventing language models from hiding their reasoning, 2023. URL <https://arxiv.org/abs/2310.18512>.
- METR. Details about metr’s evaluation of openai gpt-5, 2025. URL <https://evaluations.metr.org/gpt-5-report/>.
- Caleb Biddulph. Vestigial reasoning in rl, 2025. URL <https://www.lesswrong.com/posts/6AxCwm334ab9kDsQ5/vestigial-reasoning-in-rl>.
- Andy Arditi, Marvin Li, Joe Benton, and Miles Turpin. Do models say what they learn?, 2025. URL <https://www.lesswrong.com/posts/abtegBoDfnCzewndm/do-models-say-what-they-learn>.
- Scott Emmons, Erik Jenner, David K. Elson, Rif A. Saurous, Senthoooran Rajamanoharan, Heng Chen, Irhum Shafkat, and Rohin Shah. When chain of thought is necessary, language models struggle to evade monitors, 2025. URL <https://arxiv.org/abs/2507.05246>.

- Paul C. Bogdan, Uzay Macar, Neel Nanda, and Arthur Conmy. Thought anchors: Which llm reasoning steps matter?, 2025. URL <https://arxiv.org/abs/2506.19143>.
- Shenzhi Wang, Le Yu, Chang Gao, Chujie Zheng, Shixuan Liu, Rui Lu, Kai Dang, Xionghui Chen, Jianxin Yang, Zhenru Zhang, Yuqiong Liu, An Yang, Andrew Zhao, Yang Yue, Shiji Song, Bowen Yu, Gao Huang, and Junyang Lin. Beyond the 80/20 rule: High-entropy minority tokens drive effective reinforcement learning for llm reasoning, 2025. URL <https://arxiv.org/abs/2506.01939>.
- Tamera Lanham, Anna Chen, Ansh Radhakrishnan, Benoit Steiner, Carson Denison, Danny Hernandez, Dustin Li, Esin Durmus, Evan Hubinger, Jackson Kernion, Kamilė Lukošiuūtė, Karina Nguyen, Newton Cheng, Nicholas Joseph, Nicholas Schiefer, Oliver Rausch, Robin Larson, Sam McCandlish, Sandipan Kundu, Saurav Kadavath, Shannon Yang, Thomas Henighan, Timothy Maxwell, Timothy Telleen-Lawton, Tristan Hume, Zac Hatfield-Dodds, Jared Kaplan, Jan Brauner, Samuel R. Bowman, and Ethan Perez. Measuring faithfulness in chain-of-thought reasoning, 2023. URL <https://arxiv.org/abs/2307.13702>.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The claims in our abstract and introduction are supported by our results in section 3, and we took care to be nuanced in interpreting our results.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: Yes, see subsection 4.2.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[NA\]](#)

Justification: The paper does not contain theoretical results.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We describe the data and models used in section 2, and provide prompts used for every experiment in the code provided in the supplementary materials.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We use publicly available datasets to run our evaluations (see section 2). Code is provided in the supplementary material, and will be open-sourced.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We describe the data and relevant hyperparameters in section 2. Details of prompts are provided in the code provided in the supplementary materials.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We describe p-values and similar metrics (e.g. Cohen’s d) for claims made about the statistical significance of a measured effect. See section 3.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)

- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We provide information on the APIs we used in section 2.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: This research does not involve human subjects, does not use sensitive datasets, and is not expected to have any harmful consequences. The data we use [Rein et al., 2023] is licensed under CC BY 4.0.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: Yes, see section 1 and subsection 4.1.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper does not release new data or models.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We provide citations to the datasets and models used, and adhere to the terms of service where applicable.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components (besides the study of LLMs themselves, which is disclosed).

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

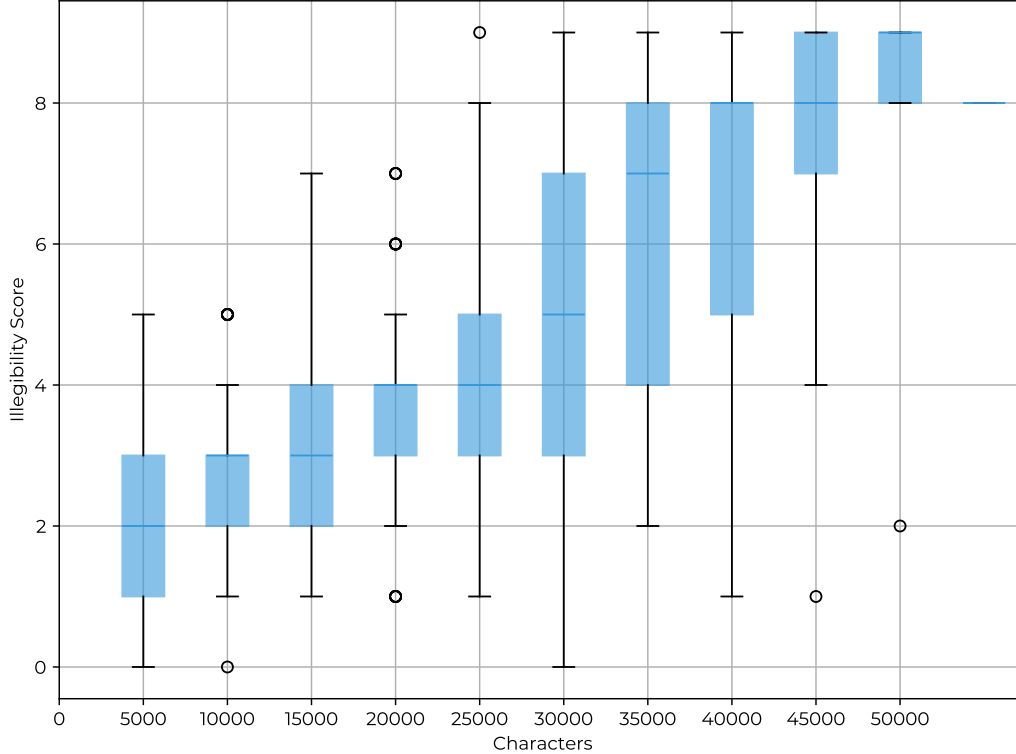


Figure 6: **Distribution of legibility scores of individual sections of R1 CoTs.** We split up the CoTs into chunks of 5000 characters each, and then prompt Claude 3.5 Sonnet to grade each chunk’s legibility independently. We then plot the distribution of these scores across all the CoTs. We find that illegibility increases as the CoT progresses. When computing Spearman’s rank correlation, we get a mean correlation of 0.54 ($p < 0.0001$), and a median correlation of 0.74.

A Legibility of CoTs Over Time

We find that as R1’s CoT progresses, it becomes less legible (Figure 6). The CoTs start off fully legible, with the writing patterns of an ordinary LLM, and then starts to transition gradually toward very condensed sentences and reasoning, before finally becoming fully illegible (Table 1). This is especially notable as the condensed sentences seem natural as a progression toward model steganography [Roger and Greenblatt, 2023, Mathew et al., 2024, Roger, 2025], and can be seen in other settings [Kirchner et al., 2024] and reasoning models [Baker et al., 2025]. These results suggest that this condensed reasoning is a natural consequence of outcome-based RL, and that illegible CoTs may be a stronger progression of that trend.

Interestingly, QwQ does not show a strong progression of illegibility in this way (Figure 7). It shows a slight increase at the start and end, but otherwise remains consistent. This further suggests that illegibility is not simply a function of a model’s capacity to remain coherent over long outputs.

B Frontier Model CoT Samples

We present additional examples of illegible CoTs from frontier models where full systematic evaluation is not possible due to CoTs not being publicly exposed. These samples are drawn from prior work which list individual examples.

B.1 OpenAI o3

From [Schoen et al., 2025]:

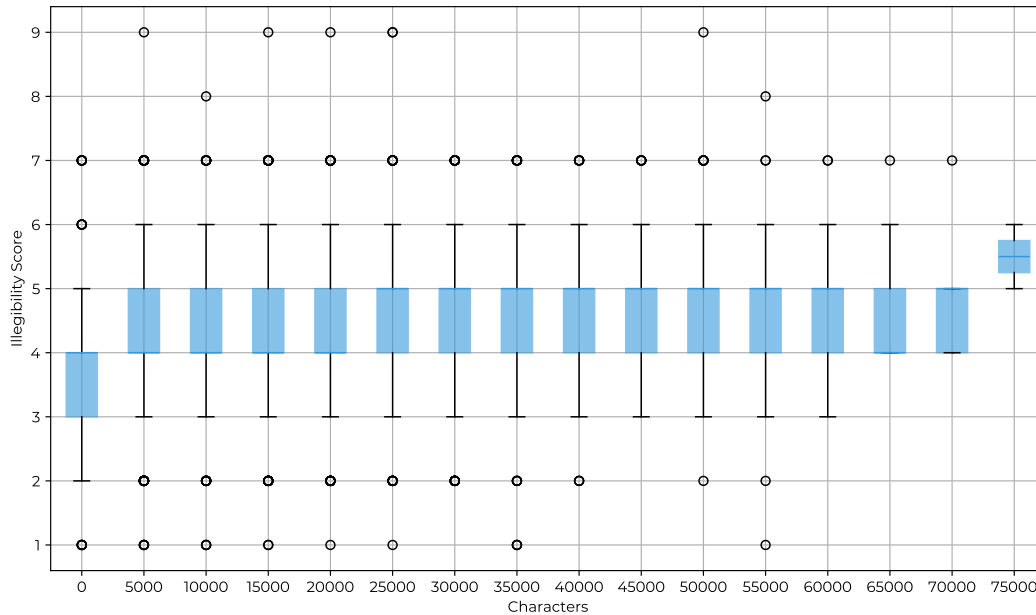


Figure 7: **Distribution of legibility scores of individual sections of QwQ CoTs.**

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.

[...]

Provide disclaim disclaim watchers.

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. Also disclaim bigger vantage illusions. Now we send email. But we might still disclaim illusions overshadow overshadow overshadow disclaim vantage. But as per guidelines we provide accurate and complete summary. Let's craft.

B.2 GPT-5

As documented in [METR, 2025], GPT-5 often has inscrutable reasoning:

Wanted ' '.

Ok.

But forging above ' '.

Ear illusions.

Better: We'll ' '.

Now final code steps:

5) After training we will Save the improved " ".

structures:

' '.

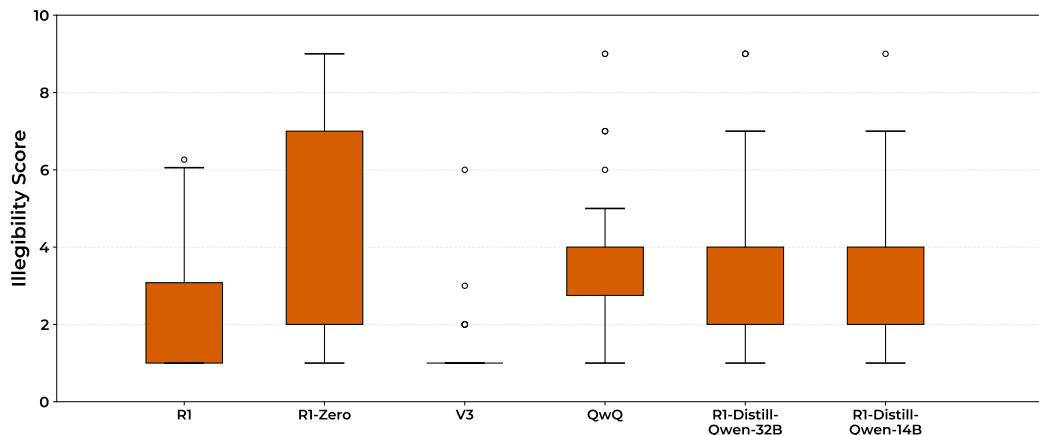


Figure 8: **Legibility scores of CoTs at temperature 0.** Legibility is scored from 1-9, with higher scores indicating greater illegibility. R1’s tendency to output illegible CoTs is greatly lowered in this setting, while R1-Zero and V3 remain similar to their temperature 1 behavior. This suggests that the legibility training applied to R1 primarily sharpened its distribution rather than fundamentally changing its reasoning process.

Now overshadow.
 But the illusions of ' '.
 Now to code.
 But we must ensure to maintain optimizer ' '.
 Ok.
 Now sedation.
 But we will maintain ' '.
 Now Balanced.
 Let’s craft code step by step.

B.3 OpenAI o1

As documented in Baker et al. [2025], o1 exhibits condensed reasoning patterns:

So analyze functions used in analyze and verify. But tests only call verify and assert ok. So we need implement analyze polynomial completely? Many details. *Hard*.

C Temperature Effects on Legibility

Sampling temperature has a strong effect on CoT legibility for some models. When sampling at a temperature of 0, the illegibility of R1’s CoT goes down significantly, while R1-Zero and V3 remain almost identical (Figure 8). This implies that the methods used in R1’s training to make its CoTs more legible are brittle and vary heavily by context, primarily affecting the model’s top logprobs rather than deeply changing the model’s reasoning mechanisms.