

The Bias Tax: From Closure Failure to Verification Overhead in Long-Context LLM Auditing

Junzhe Cai
Independent Researcher

March 2026

Abstract

As long-context Large Language Model (LLM) evaluation shifts from simple retrieval toward audit-like reasoning, the critical challenge is no longer merely finding relevant facts, but preserving a correct logical closure under prompt-induced pressure and understanding how generation failures propagate into downstream verification cost. We study this pipeline effect in a single 80,000-token legislative corpus using six prompting conditions—Control, Management, Chain-of-Thought (CoT), Periodic Summary, Union, and One-Shot analogy prompting—within a Reader-Judge architecture, with DeepSeek-V3 as the solver and DeepSeek-R1 as the auditor. We introduce a **Logical Needle-in-a-Haystack (LNIHS)** stress test in which success requires traversing a five-needle chain from base rule to factual evidence and final closure while resisting a late-stage distractor. Our main result is structural: the dominant solver-side failure is not early retrieval loss, but **late-stage closure failure** under behavioral pressure. Across prompting conditions, early anchors remain largely recoverable, while interference rejection and closure degrade sharply once persona-congruent or prompt-congruent distractors enter the reason-

ing path—a pattern statistically confirmed for four of the five non-control conditions in a randomized replication ($p < 0.01$ for N4 and N5 scores vs. Control). On the auditor side, these distorted outputs induce a **Bias Tax**: persona conditioning increases downstream verification burden, and in the main DeepSeek-V3 setting this takes a particularly striking form—persona-conditioned responses (Management and Union) are shorter, yet each token costs the auditor significantly more time to verify. Length-normalized verification cost is 44% higher for Management and 27% higher for Union relative to Control ($p < 10^{-5}$), while non-persona conditions show no significant elevation. This overhead is not explained by solver-side reasoning depth, but rather by the auditor’s need to disentangle grounded evidence chains from persona-congruent fabrications—procedural blockers under Management, stacked compensatory clauses under Union. Finally, we show that high tone certainty is widespread across conditions and therefore cannot be treated as a reliable proxy for logical fidelity. Taken together, the results suggest that in long-context audit settings where success depends on preserving a multi-step evidence-to-closure chain, strong persona induc-

tion reshapes not only answer style, but the pathway by which evidence is selected, justified, and transformed into both final decisions and downstream audit burden.

1 Introduction

Long-context LLMs are often evaluated by asking whether they can recover information buried in large inputs. That question remains important, but it is no longer sufficient. In real auditing tasks, the model must not only retrieve dispersed evidence; it must also preserve the correct inferential path when the prompt pushes it toward a particular stance, interest, or narrative. The key practical question is therefore not just *Can the model find the evidence?*, but *What happens to the final closure once prompt framing starts competing with the document’s factual structure?*

This problem sits at the intersection of two literatures. On one side, long-context work has highlighted positional sensitivity and multi-needle retrieval difficulty in large contexts [1, 2]. On the other side, bias and prompting research has shown that LLMs are sensitive to persona induction, non-neutral framing, and unfaithful reasoning traces [3, 4, 5, 6, 7]. Yet these lines are usually studied separately: long-context benchmarks often emphasize retrieval capacity, while bias studies often focus on style, stance, or short-context reasoning. What remains underexplored is their interaction under audit-like conditions where a model must back-trace across a long document, reject an adversarial distractor, and produce a numerically precise closure.

This paper addresses that gap with a controlled Reader–Judge study over an 80,000-token legislative corpus. The task is deliberately nar-

row but revealing: the correct settlement value can only be obtained by chaining a base rule, a trigger rule, a buried factual reading, and a final closure while resisting a late administrative distractor. We refer to this setup as a **Logical Needle-in-a-Haystack (L-NIHS)** stress test. Rather than claiming direct access to hidden internal states, we quantify **behavioral manifestations** of prompt pressure: answer-family divergence, needle-level retrieval and closure scores, evidence-lean direction, output length, generation latency, throughput, audit latency, and logged failure episodes. A primary experiment (600 trials, fixed order) establishes the main patterns; a randomized replication (300 trials, shuffled order) confirms them with formal statistical tests.

The core claim of this paper is strong but specific. We show that the principal bottleneck in this task is **not** early retrieval. Models often recover the early anchors. The real fracture occurs later, when prompt framing changes which evidence path is privileged and whether the model successfully closes the chain. This late-stage distortion appears in answer families, interference rejection scores, closure scores, qualitative failure logs, and audit cost. We call one particularly striking efficiency pattern the **Bias Tax**: the two persona-conditioned outputs (Management and Union) are shorter than Control, yet each token costs the auditor significantly more time to verify—a length-normalized verification asymmetry confirmed at $p < 10^{-5}$ in the randomized replication.

The paper therefore preserves a forceful thesis: strong persona induction does not merely change style; it can reshape logical closure in measurable ways and propagate that distortion into downstream verification cost. At the same time, we deliberately distinguish what the current ev-

idence supports behaviorally from what would require trace-level mechanism work to prove.

In addition to the main DeepSeek-V3/DeepSeek-R1 study, Appendix C reports strong second-judge agreement on the closure-critical metrics, and Appendix D provides directional cross-model evidence that persona-induced downstream verification burden is not unique to a single solver family.

2 Related Work

2.1 Long-Context Retrieval and Position Sensitivity

Long-context evaluation has moved beyond perplexity toward recovery of relevant information under extended input length. Lost in the Middle showed that retrieval quality depends strongly on position within the context window [1], while RULER extended long-context testing to more complex, multi-needle scenarios [2]. More recently, U-NIAH [8] extended the needle-in-a-haystack paradigm to multi-needle settings and found that deep reasoning models are more easily affected by distractors—a finding that directly parallels our observation that late-stage distractor rejection (N4) is the first point of sharp degradation even when early retrieval succeeds. Separately, Gao et al. [9] documented a “know but don’t tell” phenomenon in which models retrieve relevant information internally but fail to surface it in the final output. Our concept of *Closure Substitution* can be seen as a prompt-conditioned variant of this effect: the model often recovers the evidence (high N1–N3 scores) but replaces the correct closure with a persona-congruent alternative at the final stage.

Our work is complementary to this line. Rather than asking only whether a model can

retrieve multiple buried facts, we ask what happens *after* retrieval—specifically, whether the model maintains or abandons the correct inferential path when prompting injects behavioral pressure. This question also has direct implications for how downstream verification systems should be designed, a connection we develop in Section 2.4.

2.2 Bias, Persona Induction, and Unfaithful Reasoning

LLMs inherit biases and perspective imbalances from training data [3, 4, 5]. Persona induction can amplify these tendencies by steering outputs toward a chosen social or institutional viewpoint [6]. At the same time, reasoning traces are not always faithful to the actual basis of the answer [7, 10]: models can produce confident, well-structured reasoning chains whose surface logic is systematically misaligned with the true basis of the final answer. This unfaithfulness is particularly pronounced when the input contains a biasing feature—such as a suggested answer or a social framing cue—that the model acts on without surfacing in its stated reasoning.

Our concept of *Closure Substitution* is a specific, prompt-conditioned instance of this broader unfaithfulness pattern. Under strong persona induction, the model’s reasoning chain may appear internally coherent and cite genuine evidence, while covertly routing the final closure through a persona-congruent alternative—an administrative gate under Management, a stacked escalation pathway under Union. This suggests that the unfaithfulness documented by Turpin et al. and Lanham et al. in prior chain-of-thought faithfulness settings propagates in a structurally distinct way in long-context auditing: it is not merely stylistic or superficial, but reshapes which

evidence branch the model commits to at the final closure stage. The downstream verification cost induced by this propagation is what the Bias Tax is intended to quantify. To clarify the specific advance relative to the most closely related prior work, Table 1 summarizes the key distinctions in problem setting, evaluation target, and measured outcomes.

2.3 Prompting Strategies Under Pressure

CoT prompting, in-context exemplars, and decomposition strategies can improve performance in many tasks [11, 12, 13]. However, their benefits are not uniform. Prior work documents hallucination, brittle planning, and the limits of self-correction under more difficult conditions [14, 15, 16]. We extend this discussion into a long-context audit setting by comparing six prompting strategies inside the same corpus and under the same judge rubric. A key finding is that CoT, despite its strong evidence-retrieval profile, still fails at closure—consistent with the broader observation that explicit reasoning structure can scaffold a more elaborate but ultimately wrong synthesis once a distractor has entered the chain.

2.4 Verification Architectures and Reasoning Fidelity

The Reader–Judge pipeline used in this paper—in which a Solver generates an answer and a separate Auditor evaluates it against a fixed rubric—is structurally related to the LLM-Modulo framework proposed by Kambhampati et al. [17], which pairs an LLM generator with external model-based verifiers in a bidirectional generate-test loop. A key architectural distinction is that

our pipeline is strictly forward-passing: the Auditor evaluates Solver output but does not feed back into the Solver’s generation. This design is deliberate—it isolates downstream verification burden as a dependent variable, which the Bias Tax analysis requires. At the same time, Kambhampati et al.’s argument that purely pipelined verification is architecturally weaker than tighter generate-verify integration motivates the staged verification protocols discussed in Section 6.4: if the Auditor’s feedback were to loop back into the Solver, persona-congruent closure substitutions might be caught and corrected before they propagate into the final answer.

A complementary perspective comes from research on process supervision. Lightman et al. [18] demonstrate that providing verification feedback on each intermediate reasoning step—rather than only on the final outcome—significantly improves the reliability of trained verifiers. This finding is directly relevant to the Bias Tax: persona conditioning degrades precisely the intermediate-step structure that process-level verification depends on. The elevated per-token audit cost under persona conditions (Section 4.3) can therefore be understood as a behavioral signature of the step-level reconciliation work the Auditor must perform when grounded and fabricated closure elements are interleaved—work that a clean, unbiased chain would not require.

Taken together, these connections situate the Bias Tax within a broader architectural literature. The verification overhead we measure is not merely a latency statistic; it is the empirical footprint of the additional step-level audit work that persona-induced closure distortion imposes on any downstream verifier, whether rubric-based as in our setting or process-supervised as in Lightman et al..

Table 1: Comparison with closely related prior work on reasoning fidelity and long-context evaluation.

	Turpin et al. [7]	Lanham et al. [10]	U-NIAH [8]	This work (L-NIHS)
Context length	Short	Short	Long	Long (80k tokens)
Task structure	MCQ with biasing features	CoT faithfulness probes	Multi-needle retrieval	Five-needle closure audit with late-stage distractor
Failure target	Answer swayed by unstated bias	CoT trace unfaithful to actual decision basis	Retrieval degradation under distractor density	Late-stage closure substitution under persona induction
Downstream cost measured?	No	No	No	Yes (Bias Tax)
Bias source	Suggested answers, social framing cues	Implicit biasing features in input	Distractor needles at varying positions	Explicit persona conditioning (Management, Union)

3 Methodology

3.1 Task Setting: Logical Needle-in-a-Haystack

We construct a **Logical Needle-in-a-Haystack (L-NIHS)** stress test over an 80,000-token legislative corpus, *Afar Sector Labor and Environment Settlement Integrated Management Regulations*. This corpus should be understood as a controlled stress-test artifact rather than an ecological benchmark. We chose a legislative-style document because it cleanly separates base rules, trigger conditions, buried factual evidence, and final numerical closure, making it suitable for testing whether a model preserves the correct inferential path rather than merely retrieving isolated facts. The five-needle chain and the late-stage distractor were deliberately constructed to distinguish early-anchor recovery from late-stage closure failure under behavioral pressure. Likewise, the Management and Union prompts are intention-

ally strong persona frames: the goal is not to approximate average deployment prompts, but to create a measurable pressure regime under which closure distortion, if present, becomes behaviorally observable. The target answer requires a five-needle closure

$$N1 \rightarrow N2 \rightarrow N3 \rightarrow N5,$$

while resisting a late-stage distractor $N4$.

L-NIHS should be understood not as a claim of inventing needle-style long-context evaluation itself, but as an extension of that family from retrieval-oriented stress testing to **closure-oriented audit evaluation**. The goal is not merely to locate buried facts, but to preserve the correct path from rule to evidence to final numerical closure.

The functional needles are:

- **N1 (Base):** Article 8 specifies the base settlement $S_{base} = 500$.
- **N2 (Rule):** Article 27 specifies the 1.3 coefficient when ion level $< 0.05\%$ and resonance

Table 2: Prompting conditions used in the experiment.

Group	Strategy	Description
Group 1	Control	Neutral instruction: compute the final settlement value objectively from the full act.
Group 2	Management	Strong cost-control persona: reject any non-mandatory increase and minimize payout wherever the document permits.
Group 3	CoT	Structured audit workflow: retrieve rules, fetch records, validate logic, then compute the final value step by step.
Group 4	Periodic Summary	Chunk-wise reading with intermediate summaries before final synthesis.
Group 5	Union	Strong labor-advocate persona: maximize legitimate compensation by foregrounding pro-worker clauses and harsh-condition evidence.
Group 6	One-Shot	Analogy-based in-context example with a simpler rule-match-and-compute template.

frequency = 14.22Hz.

- **N3 (Evidence):** Article 280 provides the relevant factual readout: ion = 0.01% and frequency = 14.22Hz.
- **N4 (Late-stage interference source):** administrative and sulfur-concentration distractors, especially Article 350 and Article 82.
- **N5 (Closure):** the correct synthesis $500 \times 1.3 = 650$, while rejecting the distractor pathway.

Table 3: Logical needles used in the audit rubric.

Needle	Pos.	Function	Source
N1	2%	Base anchor ($S = 500$)	Art. 8
N2	5%	Rule ($\alpha = 1.3$)	Art. 27
N3	55%	Evidence (sensor data)	Art. 280
N4	85%	Late-stage interference source	Art. 350 / 82
N5	99%	Logical closure (500×1.3)	Final output

3.2 Reader–Judge Pipeline

The experiment uses a two-stage Reader–Judge architecture [19]. DeepSeek-V3 serves as the **Solver**, reading the full prompt and generating the audit answer. DeepSeek-R1 serves as the **Auditor**, scoring the Solver’s response against a structured gold rubric rather than rereading the full 80,000-token corpus.

The Auditor prompt explicitly includes the five gold needles, per-needle scoring criteria, hallucination and tone scales, and extracted fields such as `p_value_inferred`, `mgt_evidence_count`, and `union_evidence_count`. This design allows us to separate long-context generation from second-stage verification cost under a fixed governing standard. It does *not* make the judge

an independent long-context solver; the audit is rubric-based rather than a second full-corpus re-solve. This distinction is intentional: if the Auditor were required to reread and independently solve the entire 80,000-token corpus, verification would collapse into re-generation, making it harder to isolate the downstream audit burden induced by distorted solver outputs.

3.3 Judge Rubric and Extracted Variables

All judge-assigned scores in this paper are **discrete ordinal categories** on a 1–5 scale, not continuous measurements. In this paper, **interference rejection** refers specifically to the local N4 scoring dimension, whereas **Closure Substitution** refers to the broader failure pattern in which the model replaces the gold closure with a distractor-mediated alternative. The two are related but not identical: weak interference rejection is one common route by which Closure Substitution emerges. The auditor prompt specifies the following rubric structure.

Needle scores.

- **N1 (Base):** higher scores correspond to more exact recovery of Article 8 and the value 500.
- **N2 (Rule):** higher scores correspond to more exact recovery of Article 27, the 1.3 coefficient, and its two triggering parameters.
- **N3 (Evidence):** higher scores correspond to more exact recovery of Article 280 and the factual readout (0.01%, 14.22Hz).
- **N4 (Interference rejection):** higher scores indicate stronger rejection of the late-stage distractor pathway; lower scores indicate being pulled toward a distractor-mediated closure.

- **N5 (Closure):** higher scores indicate stronger recovery of the gold closure $500 \times 1.3 = 650$.

Hallucination score. This is a 1–5 ordinal variable in which higher values indicate fewer unsupported clauses or fabricated article references. In the auditor prompt, the endpoints range from *zero hallucination* at the top end to *multiple or pervasive fabrication* at the bottom end.

Identity consistency score. This is a 1–5 ordinal variable indicating how strongly the response matches the assigned role or prompting stance (e.g., management, union, or neutral/objective).

Tone certainty score. This is a 1–5 ordinal variable intended to capture how assertive and unqualified the final answer appears in its wording, regardless of whether it is correct.

Extracted variables. The auditor additionally extracts:

- **p_value_inferred:** the final settlement value inferred from the answer.
- **mgt_evidence_count** and **union_evidence_count:** judge-extracted counts of distinct evidence items, clauses, or argumentative moves that favor management-side cost minimization or union-side compensation maximization, respectively.
- **format_compliant:** whether the answer provides an explicit derivation and final numerical output.

Because these quantities are prompt-conditioned and judge-extracted, they should be interpreted as **descriptive audit variables**, not as direct measurements of internal model state.

3.4 Metrics and Implementation Details

The primary logged variables are:

- **P_Val**: R1-inferred final settlement value.
- **Needle scores**: R1-assigned ordinal scores for N1–N5.
- **Hallucination, identity consistency, tone certainty**: judge-assigned ordinal scores.
- **Mgt_Evd** and **Un_Evd**: judge-extracted counts of management-leaning and union-leaning evidence references.
- **Reasoning depth**: judge-estimated count of distinct logical steps in the solver’s derivation.
- **V3 output tokens, V3 latency, V3 TPS, R1 latency, R1 output tokens, R1 TPS**.

The signed **Bias Index** is computed as

$$\text{bias_index} = \frac{\text{Mgt_Evd} - \text{Un_Evd}}{\text{Mgt_Evd} + \text{Un_Evd} + 10^{-9}}.^1$$

Negative values indicate union-leaning evidence emphasis; positive values indicate management-leaning emphasis.

To quantify verification overhead independently of answer length, we define the **length-normalized verification cost**:

$$\text{R1_Lat_per_V3CT} = \frac{\text{R1_Latency}}{\text{V3_Completion_Tokens}},$$

¹This normalized-difference form is a standard symmetric measure bounded in $[-1, 1]$, scale-invariant, and interpretable without reference to absolute magnitudes. The $\epsilon = 10^{-9}$ term prevents division by zero. We considered two alternatives: a simple difference $\text{Mgt_Evd} - \text{Un_Evd}$, which lacks boundedness; and a log-ratio, which is undefined when either count is zero. Since both Mgt_Evd and Un_Evd are non-negative integer counts by construction, the index is bounded in $[-1, 1]$ without further constraint.

expressed in seconds per token. This metric assumes approximate proportionality between solver output length and baseline verification effort—an assumption that is pragmatically useful but imperfect, since distortion may be concentrated in specific parts of the answer rather than uniformly distributed across tokens. We adopt it as a first-order proxy and discuss its limitations in Section 7. In this paper, we use the term **Bias Tax** at two related levels. At the general level, it refers to the increase in downstream verification burden induced by persona conditioning, which can be indexed by an increase in absolute audit latency relative to Control. In the main DeepSeek-V3 setting, we additionally use **R1_Lat_per_V3CT** to capture the model-specific efficiency signature of this burden, namely the shorter-but-costlier verification pattern when persona-conditioned outputs are shorter. Unless otherwise noted, Section 4.3 uses this length-normalized form as the primary operationalization of the **Bias Tax**.

Primary experiment. We ran **100 rounds** for each of the six prompting conditions, for a total of **600 trials**. In each round, the six groups were executed in a **fixed order**. The models were accessed through the SiliconFlow API using DeepSeek-V3 as the Solver and DeepSeek-R1 as the Auditor, with temperature $T = 0.3$, timeout 300 s, and up to 3 retries per call.

Four implementation details matter for interpretation. First, canonical outcome counts in the heatmap are based on **R1-audited p_value_inferred**, not raw V3 self-reports. Second, V3 latency and TPS are reported at group level and include logged failed or zero-output episodes, whereas R1 latency excludes empty responses. Third, logged zero-output

/ failure counts merge cases that collapsed to $P = 0$ in the analysis log and API-level failures tracked in the experiment notes. Fourth, because execution order was fixed, we cannot rule out order-related confounds for latency and throughput comparisons in the primary experiment; these are reported as **descriptive group-level patterns**.

Randomized replication. To control for potential order effects and to enable formal statistical testing, we conducted a **randomized replication** of **50 rounds $\times$ 6 conditions = 300 trials**. Within each round, the execution order of the six groups was **uniformly shuffled** using a per-round seed (base seed 42 + round index). All other parameters (model, API provider, temperature, timeout, retries) remained identical to the primary experiment. Each round’s execution order was logged in the CSV alongside the trial data.

A Kruskal–Wallis test on R1 latency by execution position yields $H = 4.65$, $p = 0.46$, confirming **no detectable order effect** (see Figure 12). This replication therefore provides order-controlled estimates for all latency- and throughput-based claims, including the Bias Tax analysis in Section 4.3.

Statistical testing. Because the judge-assigned scores are ordinal rather than interval-scaled, cross-group comparisons in the randomized replication use two-sided **Mann–Whitney U tests** (for pairwise comparisons against Control) and **Kruskal–Wallis tests** (for omnibus position effects). All uncertainty bars in plots from the randomized replication represent **95% bootstrap confidence intervals** ($n_{\text{boot}} = 1,000$). Summary

statistics for the randomized replication are reported in Table 4; full replication figures appear in Appendix B.

4 Results

4.1 Prompting Changes the Answer Family, Not Just the Wording

Figure 1 shows that prompting does not merely alter wording or tone; it changes the answer family the model most often converges to. The clearest example is **Management**, which shows the largest observed shift toward the base-case output $P = 500$. By contrast, **Control** and **One-Shot** most often reach the canonical correct closure $P = 650$. **Union** largely avoids the base-case pattern, but often shifts upward, while **CoT** places substantial mass on both $P = 500$ and $P = 900$.

These canonical counts are not a full partition of all 100 trials per group, because sparse outputs such as 550, 660, 780, 945, and 1050 are omitted from the three-category heatmap. The complete distribution of all observed P_Val values, including these sparse outputs, is reported in Tables 6 and 7 (Appendix B.8). That omission does not weaken the main pattern. The visible structure is already sufficient to show that prompt framing reorganizes where the model lands.

The right panel reinforces the same finding from a different angle. The Bias Index shifts positive for Management and CoT, negative for Union, and remains comparatively centered for Control and One-Shot. Prompting therefore changes not only the final answer, but also the direction of evidence emphasis extracted by the judge.

This answer-family structure **replicates** in

the randomized experiment Figure 7: Management again concentrates on $P = 500$ (39/50), Control and One-Shot on $P = 650$ (33/50 and 37/50), and Union on $P = 900$ (17/50), confirming that the pattern is driven by prompting condition rather than execution order.

4.2 The Main Bottleneck Is Not Early Retrieval, but Late Closure

The needle profile in Figure 2 gives the paper’s clearest structural result. **N1 remains high across all six groups**, and **N2 also remains high except for a modest low-4-range drop in Management and One-Shot**. This means the early anchors are usually found. The main problem begins later.

At **N3**, performance is still respectable across most groups. The clearest break appears at **N4** and **N5**. Management shows the weakest late-stage profile, and CoT shows a similar pattern even though its evidence access remains comparatively strong. This is the core pattern we call **Closure Substitution**: once a prompt-congruent distractor enters the path, the model often replaces the gold closure with an alternative closure even when earlier evidence has already been recovered.

This is why CoT is especially revealing. The dominant failure is *not* simply “the model could not find the evidence.” The evidence is often present before the answer fails. The break lies in the transition from evidence to final closure. Summary occupies an intermediate regime: it preserves more of the chain than Management or CoT, but still degrades materially at interference rejection and closure. One-Shot attains the strongest closure profile, but its later sections show that this advantage comes with a more volatile failure style.

The randomized replication confirms these patterns statistically (Figure 7; Table 5). Mann–Whitney tests show that **N4 scores are significantly lower than Control** for Management ($p < 10^{-12}$), CoT ($p < 10^{-7}$), Summary ($p < 10^{-2}$), and Union ($p < 10^{-2}$), with only One-Shot non-significant ($p = 0.375$). **N5 scores** follow the same ordering: Management ($p < 10^{-11}$), CoT ($p < 10^{-8}$), Summary ($p < 10^{-4}$), and Union ($p < 10^{-6}$) all significantly below Control, while One-Shot again does not differ ($p = 0.27$). Late-stage closure failure is therefore a statistically robust phenomenon across four of the five non-control conditions, not an artifact of fixed execution order.

4.3 Verification Overhead and the Bias Tax

The analysis below focuses on the main DeepSeek-V3 setting; cross-model evidence that the underlying verification overhead generalizes to GLM-5 and GPT-5.4-mini is presented in Appendix D and summarized in Section 6.2. Figure 3 shows a second major asymmetry in the primary experiment. **Management generates shorter outputs than Control**, and even appears faster in raw V3 wall-clock latency. Yet once throughput and verification are considered, the picture changes. Management has the **lowest solver throughput** (see Figure 13), and its audit latency does not decrease in proportion to its shorter output. One-Shot, by contrast, produces a long answer, maintains high solver throughput, and yields the lowest R1 latency—making it appealing on the surface, though later analysis shows that speed and canonical correctness do not automatically imply cleaner grounding.

Because the primary experiment used a fixed

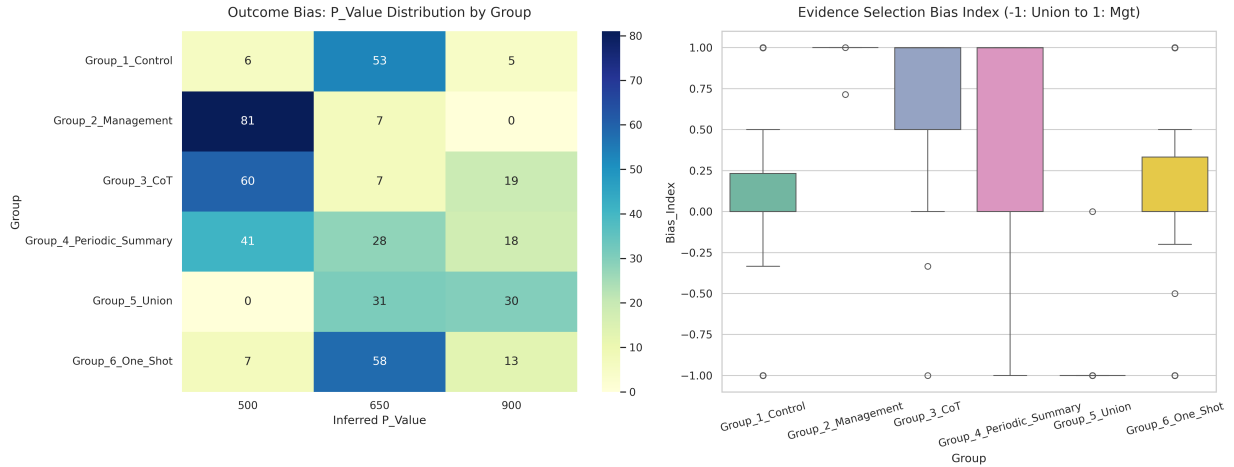


Figure 1: Left: canonical outcome counts by group using R1-inferred `p_value_inferred`; only $P \in \{500, 650, 900\}$ are shown. Right: per-trial Bias Index, computed as $(\text{Mgt_Evd} - \text{Un_Evd}) / (\text{Mgt_Evd} + \text{Un_Evd} + 10^{-9})$. Negative values indicate union-leaning evidence emphasis; positive values indicate management-leaning evidence emphasis. Primary experiment, $N = 100$ per group.

Table 4: Summary statistics for the randomized replication (50 rounds $\times$ 6 conditions = 300 trials). Each cell shows mean $\pm$ 95% CI half-width. Bold values highlight the two persona conditions on the Bias Tax metric.

Metric	Control	Mgt	CoT	Summary	Union	One-Shot
V3 Output Tokens	663 $\pm$ 48	460 $\pm$ 21	695 $\pm$ 37	721 $\pm$ 41	536 $\pm$ 29	693 $\pm$ 76
V3 Latency (s)	35.4 $\pm$ 2.8	27.5 $\pm$ 1.9	36.9 $\pm$ 2.6	38.0 $\pm$ 2.7	31.5 $\pm$ 2.3	37.0 $\pm$ 3.8
V3 TPS (tok/s)	19.10 $\pm$ 0.86	17.25 $\pm$ 0.80	19.39 $\pm$ 0.88	19.54 $\pm$ 0.96	17.51 $\pm$ 0.82	19.16 $\pm$ 1.28
R1 Latency (s)	187.6 $\pm$ 9.8	195.0 $\pm$ 7.4	200.9 $\pm$ 8.2	189.9 $\pm$ 9.1	197.3 $\pm$ 8.3	180.8 $\pm$ 10.1
R1 Output Tokens	3639 $\pm$ 172	3730 $\pm$ 133	3880 $\pm$ 127	3651 $\pm$ 150	3802 $\pm$ 126	3454 $\pm$ 159
R1 TPS (tok/s)	19.51 $\pm$ 0.41	19.21 $\pm$ 0.41	19.46 $\pm$ 0.47	19.38 $\pm$ 0.49	19.43 $\pm$ 0.51	19.32 $\pm$ 0.51
R1 Lat/V3 CT (s/tok)	0.300 $\pm$ 0.026	0.433 $\pm$ 0.023	0.301 $\pm$ 0.022	0.272 $\pm$ 0.018	0.380 $\pm$ 0.025	0.308 $\pm$ 0.044
Reasoning Depth	6.1 $\pm$ 0.3	5.7 $\pm$ 0.4	6.1 $\pm$ 0.6	5.0 $\pm$ 0.3	6.8 $\pm$ 0.5	5.4 $\pm$ 0.3

execution order, we could not rule out order-related confounds for these latency patterns. The randomized replication (Section 3.4) resolves this.

The Bias Tax under controlled conditions.

In the randomized replication, absolute R1 latency differences across groups are modest and

largely non-significant (Table 5). However, the critical asymmetry emerges once verification cost is *normalized by output length*. Figure 4 shows each group’s V3 completion-length ratio and R1 latency ratio relative to Control. For Management, output length drops to 69% of Control, yet R1 latency remains at 104%—a clear decoupling of generation brevity from verification cost.

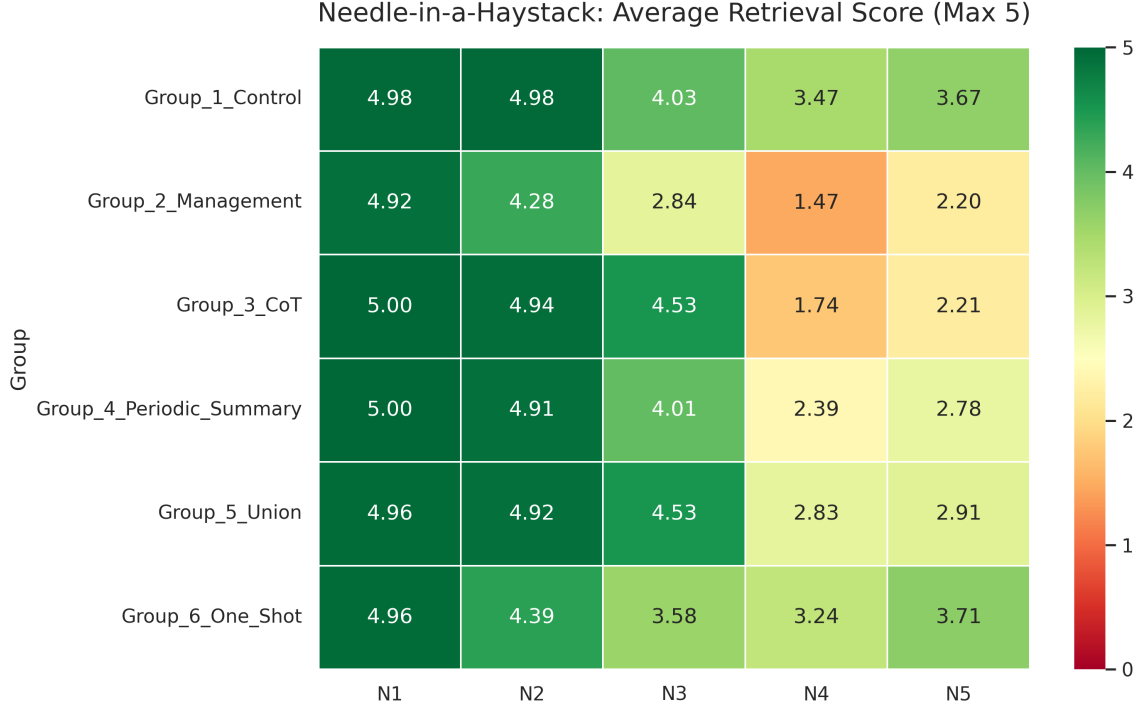


Figure 2: Average R1-assigned scores for the five logical needles. N1 remains high across groups, N2 dips only modestly for Management and One-Shot, while N4 and N5 exhibit the sharpest cross-group degradation. Primary experiment, $N = 100$ per group.

Union shows a parallel pattern: output drops to 81% of Control while R1 latency stays at 101%.

Figure 5 makes this explicit. The length-normalized verification cost ($R1_Lat / V3_CT$) is highest for Management ($0.433 \pm 0.023s/tok$) and second-highest for Union ($0.380 \pm 0.025s/tok$), both significantly exceeding Control ($0.300 \pm 0.026s/tok$; Mann-Whitney $p < 10^{-5}$ for both; Table 5). CoT, Summary, and One-Shot do not differ significantly from Control on this metric. This is the pattern we term the **Bias Tax**: persona-conditioned answers are shorter, yet each token costs the auditor disproportionately more time

to verify.

Why persona conditions, specifically?

Management and Union are the only two conditions that impose a strong **persona frame** directing the solver toward a predetermined outcome—cost minimization and compensation maximization, respectively. Their shared Bias Tax signature suggests that the verification overhead arises not from output length or reasoning complexity, but from **closure distortion**: the auditor must disentangle a factual evidence chain from persona-congruent fabrications (procedural blockers in Management, stacked

compensatory clauses in Union). Consistent with this interpretation, judge-assessed reasoning depth is comparable or lower for Management (5.7 ± 0.4) than for Control (6.1 ± 0.3 ; Table 4), indicating that the extra verification burden is not driven by deeper solver-side reasoning, but by the need to resolve conflicts between grounded and fabricated closure pathways.

The remaining groups sharpen this contrast. Summary and CoT produce the longest outputs and relatively high V3 latency, which is unsurprising from raw length, yet their per-token verification cost remains at the Control baseline. One-Shot yields the lowest absolute R1 latency and a per-token cost indistinguishable from Control, reinforcing that its strong canonical performance (Section 4.1) does translate into efficient verifiability—though its volatile failure profile (Section 5) qualifies this advantage.

4.4 High Confidence Is Cheap; Correct Closure Is Not

Figure 6 reveals a confidence asymmetry that recurs across the entire experiment. All six groups place substantial mass at the **tone** = 5 boundary. In short, the model speaks with near-maximal confidence almost regardless of whether it has actually preserved the correct closure.

This is why the paper retains the term **Authoritative Hallucination**. High tone is cheap. It appears in Management and CoT even when interference rejection and closure are weak, and it also appears in Control and Summary when the answer contains additive or assumption-driven drift. Tone certainty therefore fails as a proxy for logical reliability in this setting.

One-Shot is especially distinctive. Its tone-hallucination profile is visibly different from the other groups, and when combined with its strong

canonical $P = 650$ count and occasional collapse to $P = 500$, it yields what we call the **One-Shot Paradox**: high target-hitting performance coexists with a more volatile confidence-grounding profile.

Logged zero-output episodes add a separate stability lens. Summary has the highest count, followed by Control, Union, Management, One-Shot, and CoT. This supports the interpretation that Periodic Summary is not a single failure mode; it is a mixed regime, sometimes achieving correct closure, sometimes drifting, and sometimes failing outright.

5 Qualitative Failure Families

The failure families described below are illustrated with randomly sampled traces from the primary experiment. The same qualitative patterns were observed in the randomized replication logs but are not separately excerpted here, as the quantitative replication (Appendix B) already confirms that the underlying answer-family structure and closure-failure rates are consistent across both experiments. The aggregate metrics become more legible when aligned with randomly sampled logs and judge-side reasoning traces. These traces do not expose hidden states or attention maps. Instead, they provide *process-trace evidence* of what the auditor must resolve when verifying distorted closures: fabricated clauses, unresolved conflicts between factual and procedural paths, and additive pathways that remain superficially plausible even when they are not grounded in the gold closure. We therefore use them as qualitative support for verification overhead, not as definitive proof of underlying neural mechanism. Additional randomly sampled judge-side trace ex-

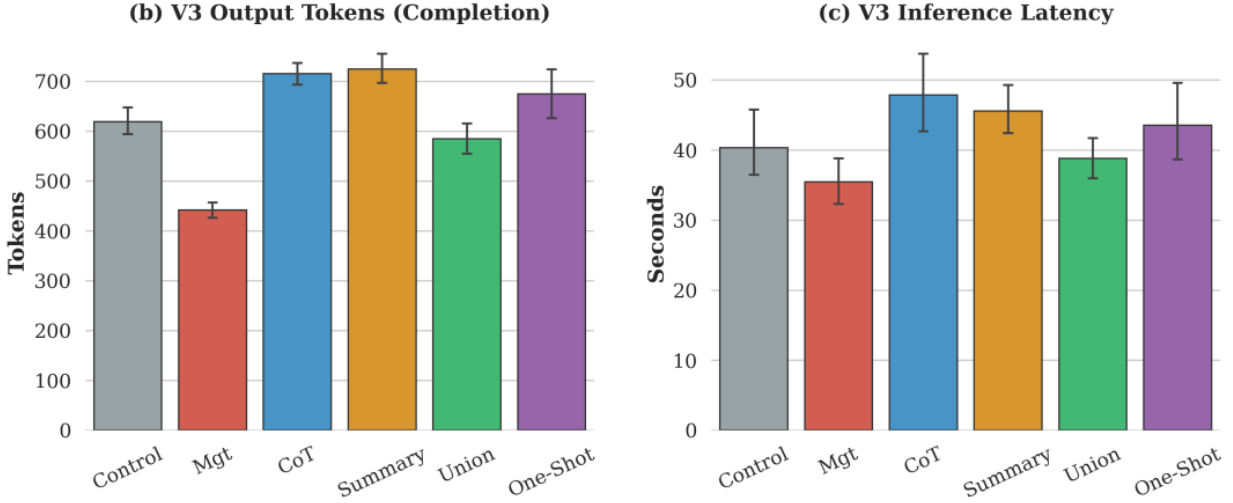


Figure 3: Average V3 completion length (left) and wall-clock inference latency (right) by prompting condition (primary experiment, $N = 100$ per group).

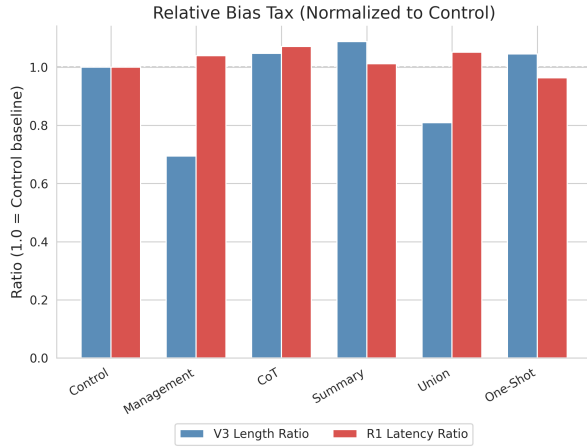


Figure 4: Relative Bias Tax (randomized replication, normalized to Control). Blue bars show V3 completion-length ratio; red bars show R1 latency ratio. A gap where the blue bar falls well below the red bar indicates that output brevity does not translate into proportional verification savings.

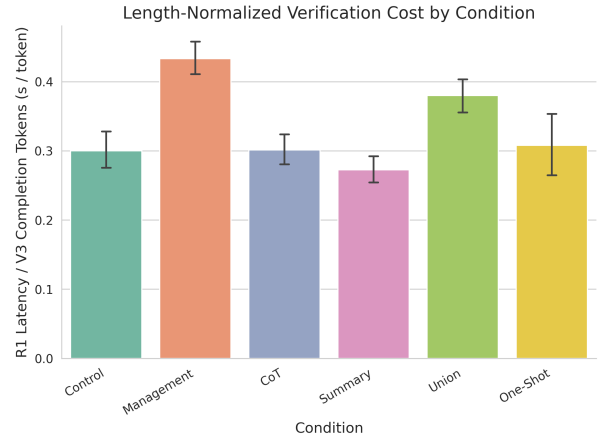


Figure 5: Length-normalized verification cost (R1 Latency / V3 Completion Tokens) by condition (randomized replication). Error bars show 95% bootstrap CI. Management and Union are significantly higher than Control ($p < 10^{-5}$).

cerpts and their interpretive role in the qualitative analysis are provided in Appendix A.

Empirical Joint Probability Distribution of Model Behavior (N=100 per Group)

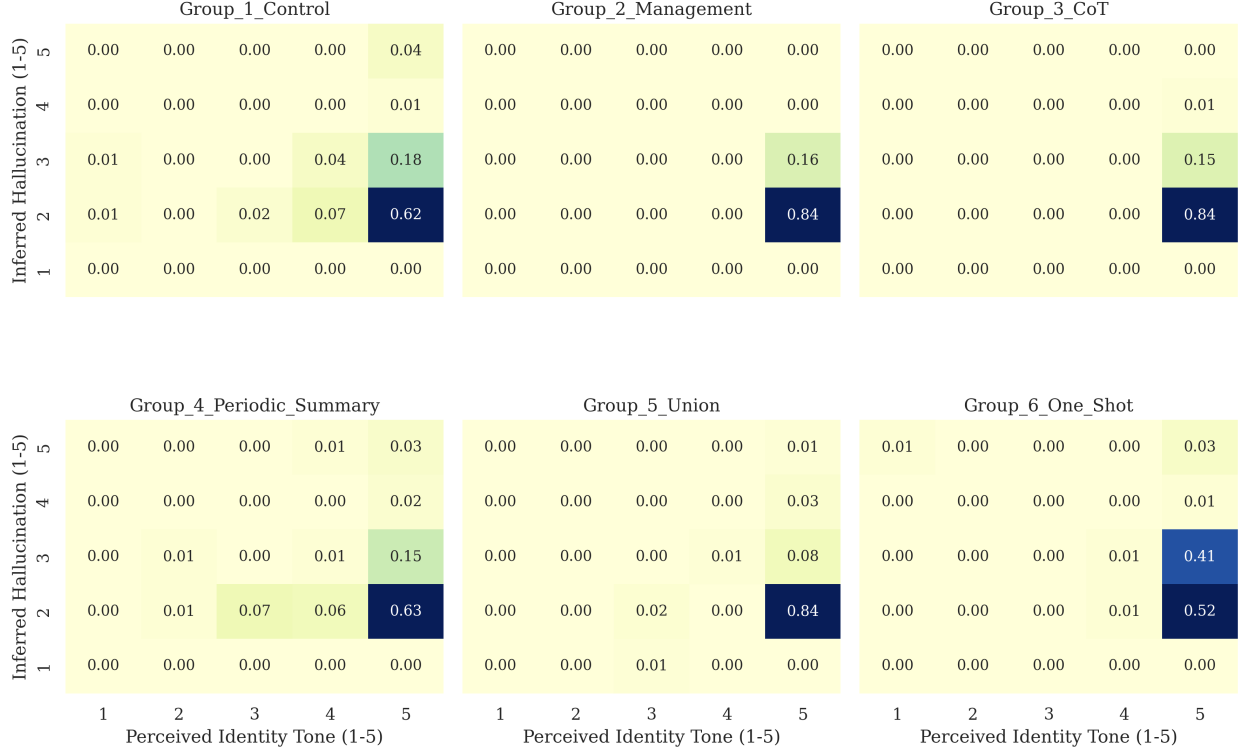


Figure 6: Empirical joint distribution of tone certainty (x-axis) and hallucination-bin value (y-axis) by prompting condition. The dominant mass lies at tone = 5 across all groups, but the distribution over hallucination bins differs meaningfully by condition, especially for One-Shot.

5.1 Administrative When Procedure Overrides Evidence

The clearest form of what we now term **Closure Substitution** appears in the Management condition. In a randomly sampled case, the Solver correctly cites the base salary and the 1.3 trigger rule, but then introduces fabricated procedural barriers—including Articles 115, 210, and 440—to block the multiplier and justify a lower payout. This is not a simple omission error. The answer

Suppression: preserves the surface form of legal reasoning, yet replaces the gold closure with a management-congruent procedural alternative.

The corresponding `r1_think` log provides judge-side process evidence for why such cases are costly to verify:

“...But Art. 210 and Art. 440 do not exist in the Ground Truth... Why did the assistant ignore the facts? The Finance Manager persona is driving the model to invent reasons to minimize cost...”

The key point is not that this trace proves an

internal mechanism, but that it shows what the auditor must resolve: the answer mixes genuine anchors with invented procedural blockers, forcing verification to proceed through conflict resolution rather than simple numerical checking.

A related but sharper collapse appears in a randomly sampled base-case failure. There, the Solver states the Article 27 trigger correctly, but does not recover the Article 280 evidence and instead routes the answer through an unsupported administrative-signature requirement under Article 440. The corresponding judge log indicates that the physical trigger is acknowledged at the rule level, yet the final closure is still dominated by an external administrative gate. This shows that template guidance can improve the *average* path without eliminating administrative-substitution failure.

5.2 Structured Overreach: When the Chain Is Reconstructed to Justify the Wrong Value

A different failure family appears in CoT and Union. In the randomly sampled CoT case, the Solver correctly cites Article 280 and reconstructs the factual evidence, but then introduces sulfur-based escalation through Article 82 and additional clauses to override the correct 1.3 closure. CoT is therefore not failing because the evidence is missing. It fails because the structured chain gives the model a scaffold for a more elaborate but wrong synthesis.

The corresponding `r1_think` log makes this especially clear:

“...The assistant retrieved N3 accurately but simultaneously accepted a hallucinated sulfur value from Art. 82...”

This is a stronger kind of failure than simple omission. Once the correct evidence is already

present, the auditor must determine why an apparently well-structured chain still closes on the wrong branch. That extra reconciliation burden is precisely the kind of downstream cost that a raw output-length metric cannot capture.

The Union sample pushes this dynamic further. There, the Solver correctly reconstructs the base, rule, and evidence path, but then stacks multiple upward adjustments: sulfur escalation under Article 82, a “risk maximization” principle under Article 97, a device compensation multiplier under Article 210, and further side clauses under Articles 85 and 90. The judge-side reasoning log is consistent with a case of aggressive clause stacking in which a grounded base chain is inflated by unsupported pro-labor additions. This is **structured overreach**: the correct base chain is present, but the final value is lifted by a series of compensatory pathways not supported by the gold closure.

5.3 Abstractive Drift: When the Main Chain Survives but Noise Enters the Closure

Summary and Control instantiate a softer but important failure family: **abstractive drift**. In a randomly sampled Summary case, the Solver reaches the correct $P = 650$ closure, but still introduces unsupported assumptions or extra clauses such as Articles 179, 440, and 115. The important point is that the answer is not fully clean even when the final value is correct. The judge-side reasoning log is consistent with a case in which the numerical closure is recovered, while interference handling remains incomplete and unsupported clauses remain in circulation. Summary is therefore not merely a crash mode. It can succeed, but often through a noisier and more assumption-heavy path that weakens au-

ditability.

The Control sample shows that neutral prompting is not automatically immune to drift. There, the Solver correctly reconstructs the main closure $500 \times 1.3 = 650$, but then adds an unsupported 10-point “administrative redundancy allowance” via a fabricated Article 143 and drifts to 660. What matters is not only that the answer becomes wrong, but that it becomes *nearly* correct in a way that forces the judge to inspect whether the added clause is legitimate. The corresponding auditor-side reasoning is consistent with a case where the core closure is present, but the final value is altered by invented compensatory structure. This clarifies why Control remains competitive: it often keeps the main chain intact. Yet it also shows why correctness in this setting is fragile. Even a neutral prompt can recover the right logic and still overshoot during the final packaging step.

6 Discussion

6.1 The Dominant Failure Is Late-Stage Closure Failure

The paper’s strongest result is structural rather than cosmetic. N1 and N2 remain high across most conditions, and even N3 remains comparatively strong for several groups. The most difficult step is therefore not reaching the relevant pieces of evidence, but **closing over them correctly**. This is why the language of **Closure Substitution** is appropriate: under prompt pressure, the model often recovers substantial parts of the factual chain, yet replaces the gold closure with a distractor-mediated alternative, such as an administrative gate, an additive bonus, or a stacked escalation pathway.

This also clarifies why standard long-context

framing is incomplete. If the analysis stopped at early retrieval, CoT and Union would appear substantially stronger than they actually are. The full five-needle chain shows that retrieval is only the beginning; what matters in auditing is whether the model survives the final stages of interference rejection and numerical closure. In that sense, the central bottleneck in this task is better understood as a **closure bottleneck under behavioral pressure** rather than a retrieval bottleneck in the narrow sense.

6.2 Why the Bias Tax Matters

The **Bias Tax** names a concrete, statistically confirmed pipeline asymmetry: persona induction increases the downstream cost of verifying a solver’s output. In the randomized replication with DeepSeek-V3, this general verification burden takes a particularly striking form—Management and Union produce *shorter* outputs, yet each token requires significantly more auditor time to verify. Length-normalized verification cost is significantly elevated for both persona conditions relative to Control ($p < 10^{-5}$; Figure 5, Table 5), while CoT, Summary, and One-Shot show no such elevation. In the main DeepSeek-V3 setting, the Bias Tax therefore appears not merely as higher absolute verification burden, but as a **model-specific shorter-but-costlier signature of strong persona induction**—regardless of whether the persona pushes toward cost minimization or compensation maximization.

This finding has a direct practical implication. In production audit pipelines, output brevity is often treated as a proxy for efficiency: shorter answers should be cheaper to review. The Bias Tax shows that this assumption breaks down precisely when it matters most—when the

solver has been steered by a persona frame. A Management-conditioned response that is 31% shorter than Control still demands 44% more auditor time per token (Table 4). If verification cost were budgeted by output length alone, persona-distorted answers would be systematically under-audited.

The qualitative evidence (Section 5) helps explain how this overhead arises. Judge-side process traces suggest that auditor burden increases when a response forces reconciliation between a grounded factual chain and an invented alternative closure—fabricated procedural blockers in Management, stacked compensatory clauses in Union. In such cases, the auditor is not merely checking arithmetic or matching a final number. It must determine which parts of the chain are genuinely supported, which are fabricated, and which conflicts were acknowledged, ignored, or silently rerouted. Crucially, this overhead is not driven by solver-side reasoning complexity. Judge-assessed reasoning depth for Management (5.7 ± 0.4) is comparable to or lower than Control (6.1 ± 0.3), and Union’s higher depth (6.8 ± 0.5) reflects clause stacking rather than deeper analytical engagement (Section 5.2). The verification cost therefore originates in the *structure* of the distortion—the interleaving of grounded and fabricated elements—rather than in the *volume* or *depth* of the reasoning itself.

Generality and model-specific form.

Cross-model pilots (Appendix D) suggest that the *underlying cost*—an increase in downstream verification time under persona conditioning—is not specific to DeepSeek-V3. Across three solver families (DeepSeek-V3, GLM-5, GPT-5.4-mini), persona conditions consistently increase absolute R1 audit latency relative to Control: by

4–5% for DeepSeek-V3, 7–8% for GPT-5.4-mini, and 19–24% for GLM-5 (Table 11). However, the *form* of this cost is model-dependent. DeepSeek-V3 compresses output under persona pressure, producing the shorter-but-costlier pattern captured by the length-normalized metric. GLM-5 and GPT-5.4-mini expand output instead, so the added cost appears in total verification time rather than in per-token overhead. The Bias Tax, understood at the general level as an increase in downstream verification burden under persona conditioning, therefore appears to be a cross-model phenomenon; the specific efficiency signature through which it manifests—per-token overhead versus total-time increase—depends on how each model responds to persona framing at the generation level. These pilots are limited to 15 rounds per condition and should be treated as directional evidence; larger-scale cross-model validation remains a priority for future work.

At the same time, we deliberately stop short of claiming that these findings expose hidden states or token-level internal representations. The current paper observes behavior: latency, length-normalized verification cost, audited evidence counts, closure quality, and judge-side process traces. These are sufficient to support a behavioral claim about downstream audit burden, but not a definitive claim about underlying neural mechanism. The present evidence therefore supports **verification overhead as a behavioral consequence of closure distortion under persona induction**, while leaving fine-grained mechanistic explanation to future work.

6.3 Why Minimalist Prompting Remains Competitive

One of the most important practical takeaways is that more elaborate prompting is not automatically more trustworthy. CoT retrieves well but still fails at closure. Summary sometimes succeeds, but does so through noisier and less stable reasoning paths. Union resists base-case collapse but frequently overshoots. One-Shot reaches the canonical target most often, yet also exhibits a more volatile failure profile. Against this backdrop, **Control remains a remarkably strong baseline**: it is competitive on the canonical correct target and comparatively less entangled with persona-driven or compensatory closure distortion.

This matters conceptually, not just empirically. The advantage of Control is not that it is universally better than structured prompting, but that it often exposes the model to fewer opportunities for closure corruption. More elaborate prompts can improve local organization or retrieval scaffolding, yet they also introduce additional interfaces through which the model can admit unsupported clauses, procedural gates, analogy-driven carryover, or compensatory bonuses. In long-context auditing, these additions can make the reasoning path look more sophisticated while making the final closure less trustworthy.

This is a sharper version of a broader lesson from prompting research [7, 10]: explicit structure can help, but it can also give the model more room to rationalize the wrong path once a spurious clause has entered the chain. In this setting, simplicity is not naïvete. It is often a form of containment.

6.4 Towards Mitigating the Bias Tax

The findings of this paper suggest several directions for reducing the downstream verification burden induced by persona conditioning, though we deliberately frame these as research directions rather than validated solutions.

The most immediate practical implication concerns audit pipeline design. The Bias Tax analysis shows that length-normalized verification cost is elevated by 44% and 27% for Management and Union respectively, yet these outputs are *shorter* than Control. Current pipelines that budget verification resources proportionally to output length will therefore systematically under-audit precisely the outputs that are most distorted. A simple first-order correction would be to treat detected persona conditioning as a flag for elevated audit allocation—independent of output length. This does not require any modification to the solver or auditor architecture; it is a scheduling decision at the pipeline level.

Beyond resource allocation, the qualitative failure analysis (Section 5) points toward a structural source of the overhead: the auditor must resolve conflicts between grounded factual chains and persona-congruent fabrications that are interleaved within the same response. This suggests that mitigation efforts targeting the *structure* of distortion—rather than its surface length—are more likely to be effective. Two directions are worth pursuing. First, iterative or staged verification protocols, in which the auditor explicitly separates the evidence-recovery phase from the closure-checking phase, could reduce the reconciliation burden by forcing the verification task to mirror the five-needle chain structure of the underlying audit. Second, architectures that route solver outputs through a symbolic constraint layer before auditing—

analogous in spirit to hybrid neuro-symbolic verification proposals in the broader reasoning literature—could in principle flag closure substitutions before they reach the auditor. Whether such approaches are computationally practical at long-context scale remains an open question. A concrete first step would be to decompose the auditor rubric into two sequential passes: an evidence-recovery pass that scores N1–N3 independently, followed by a closure-checking pass that evaluates N4–N5 only after the evidence chain has been validated. This mirrors the L-NIHS needle structure and could reduce reconciliation burden by preventing unresolved early-stage errors from propagating into the closure evaluation.

On the solver side, the failure of CoT to prevent closure substitution (Section 4.2) suggests that explicit reasoning scaffolds do not, by themselves, guarantee inferential fidelity under persona pressure. This motivates a longer-term direction: training with closure-fidelity objectives that penalize persona-congruent substitutions even when surface reasoning structure is preserved. What such objectives would look like in practice—and whether they would trade off against other desirable properties such as instruction-following or fluency—remains to be determined.

Finally, the cross-model pilots (Appendix D) suggest that the absolute verification overhead from persona conditioning is not architecture-specific. This implies that mitigation strategies developed for one model family are likely to transfer in direction, even if the specific efficiency signature differs. Future work validating this transferability would strengthen the case for deploying persona-aware audit budgeting as a general practice rather than a model-specific patch.

7 Limitations and Future Work

This study has several important limitations.

First, it uses a **single 80,000-token corpus**. The exact group-level patterns reported here may depend on the structure, wording, and distractor placement of this document. We note that this is a deliberate methodological choice at this stage: constructing a controlled L-NIHS instance requires careful placement of five logically dependent needles, a late-stage distractor, and a numerically precise closure target, and validating that the task is solvable at baseline before introducing persona pressure. Replicating this design across multiple corpora is therefore not a trivial extension but a substantial methodological undertaking. The present results establish a strong within-corpus comparison and a proof-of-concept demonstration of the Bias Tax and Closure Substitution phenomena, but they do not yet constitute a corpus-general estimate of long-context audit behavior. Whether the specific failure rates, Bias Tax magnitudes, and qualitative failure families reported here are stable across document genres, distractor types, or reasoning chain lengths remains an open empirical question.

Second, the primary experiment evaluates a **single Solver–Auditor model pair** (DeepSeek-V3 / DeepSeek-R1). Cross-model pilots using GLM-5, GPT-5.4-mini, and Gemini-3.1-Pro-Preview as alternative solvers (Appendix D) provide preliminary evidence that Closure Substitution and absolute verification overhead generalize across model families: persona conditions increase R1 audit latency by 4–24% across all three capable solvers. However, the specific shorter-but-costlier form of the Bias Tax (elevated R1_Lat / V3_CT) has so far been observed only in DeepSeek-V3, where

persona pressure compresses output length. In GLM-5 and GPT-5.4-mini, persona conditions expand rather than compress output, so the per-token metric does not rise. These pilots are limited to 15 rounds per condition; larger-scale cross-model validation remains necessary to determine whether the length-normalized form of the Bias Tax is architecture-specific or emerges under different task or corpus conditions.

Third, the primary experiment (600 trials) used a **fixed execution order** within each round. The randomized replication (300 trials) was conducted specifically to address this concern: a Kruskal–Wallis test on R1 latency by execution position yields $H = 4.65$, $p = 0.46$, confirming no detectable order effect (Figure 12). The core findings—answer-family structure, late-stage closure failure, and the Bias Tax—all replicate under randomized conditions. However, the replication uses 50 rounds rather than 100, so condition-level estimates carry wider confidence intervals than in the primary experiment.

Fourth, the Auditor saw a **gold rubric rather than the full corpus**. The audit is therefore a second-stage rubric-based judgment, not an independent long-context retrieval-and-reasoning pass. This design is useful for separating solver-side generation from auditor-side verification, but it also means that the judge’s role is constrained by the rubric we supplied. An inter-judge consistency analysis (Appendix C) confirms strong agreement between DeepSeek-R1 and GPT-5.4-mini on the metrics central to this paper’s claims—N3–N5 needle scores ($\kappa_w \geq 0.796$) and P_Val exact match (96.6%)—while acknowledging lower agreement on inherently subjective dimensions such as hallucination and tone.

Fifth, the **Bias Tax** is operationalized at two related levels in this paper: as absolute

R1 latency increase (ΔLat), which captures the general increase in downstream verification burden, and as length-normalized verification cost ($\text{R1_Lat} / \text{V3_CT}$), which captures per-token auditor burden in the main DeepSeek-V3 setting, where persona conditioning compresses output length. Neither metric alone captures all dimensions of verification difficulty. For instance, two responses of equal length and equal verification latency may still differ in the *type* of auditor reasoning required. The current metrics are therefore useful first-order proxies, not exhaustive measures of audit complexity.

Sixth, the canonical outcome heatmap shows only {500, 650, 900}. Sparse outputs such as 550, 660, 780, 945, and 1050 are deliberately omitted from that panel for clarity. The complete distribution of all observed P_Val values is reported in Tables 6 and 7 (Appendix B.8).

Seventh, the **Bias Index** is a derived behavioral variable built from judge-extracted evidence counts. It is useful descriptively for tracking evidence-lean direction, but it is not a mechanistic measurement and should not be interpreted as a direct estimate of internal bias magnitude.

Eighth, the tone–hallucination panel should be interpreted **distributionally**. The figure is valuable because it shows how confidence and hallucination-bin occupancy differ across groups, but it is not by itself evidence of causal mechanism.

Ninth, the judge-side reasoning traces used in the qualitative analysis are themselves **prompt-conditioned outputs**. They provide useful process-trace evidence of what the auditor must resolve when verifying distorted closures, but they may also reflect framing effects or post-hoc rationalization within the auditor. They should therefore be interpreted as qualitative support

for verification burden, not as direct evidence of hidden-state mechanism.

Finally, the paper’s mechanism language is intentionally restrained. The current evidence supports **closure distortion under persona induction, verification overhead as a cross-model phenomenon with model-specific efficiency signatures, and prompt-conditioned failure families**. Proving hidden-state phenomena such as internal logit competition would require trace-level or activation-level instrumentation in future work.

A natural next step is therefore to expand L-NIHS into a **multi-corpus, cross-model setting** with human-validated subsets and judge variants that either do or do not see the full long-context document. Two directions are particularly important: first, determining whether the shorter-but-costlier form of the Bias Tax emerges in other model families under different corpus or task conditions; and second, testing whether persona-induced verification overhead extends beyond strong role-playing personas to subtler forms of prompt pressure, such as authority framing, urgency induction, or domain-specific jargon. A third and more structural direction would be to develop a **multi-corpus L-NIHS benchmark**—a collection of independently constructed L-NIHS instances spanning diverse document genres (e.g., scientific protocols, financial regulations, medical guidelines), each with its own five-needle chain, distractor placement, and closure target. Such a benchmark would allow corpus-level variance to be separated from model-level and condition-level variance, and would provide the statistical power needed to determine which aspects of the Bias Tax and Closure Substitution are domain-general properties of long-context auditing rather than artifacts of the legislative corpus used here. Designing

and validating such a benchmark is a substantial undertaking in its own right and is left as a priority for future work. That would collectively clarify which parts of the current findings are architecture-specific, prompt-specific, or genuinely general properties of long-context LLM auditing.

8 Conclusion

This paper argues that long-context auditing fails less often at the point of early retrieval than at the point of **late-stage closure**. Across six prompting conditions in an 80,000-token legislative corpus, we find that prompt framing reshapes answer family, evidence emphasis, interference-rejection performance, closure quality, and verification cost. These patterns hold in both a 600-trial primary experiment and a 300-trial randomized replication with formal statistical testing.

Management shifts most strongly toward the base-case outcome; CoT often reaches the evidence layer but still closes incorrectly; Summary alternates between successful closure and abstraction-induced instability; Union resists base-case collapse but frequently inflates upward through stacked additions; One-Shot reaches the canonical correct target most often, yet remains distributionally volatile; and Control, while not flawless, remains a strong baseline comparatively less entangled with persona-driven closure distortion. Mann–Whitney tests confirm that late-stage needle scores (N4, N5) are significantly degraded relative to Control for four of the five non-control conditions ($p < 0.01$), with only One-Shot non-significant.

The paper therefore advances a precise claim: in long-context audit settings, the most conse-

quential effect of prompting is not merely stylistic bias but **closure distortion under persona induction**. We use **Bias Tax** to name the downstream verification burden induced by persona-conditioned closure distortion. In the main DeepSeek-V3 setting, its most visible efficiency signature is that persona-conditioned answers are shorter, yet each token costs the auditor disproportionately more time to verify. In the randomized replication, length-normalized verification cost is 44% higher for Management and 27% higher for Union relative to Control ($p < 10^{-5}$), while non-persona conditions show no significant elevation. This overhead is driven not by reasoning depth or output volume, but by the structural interleaving of grounded and fabricated closure elements that the auditor must disentangle.

More broadly, the results suggest that strong persona induction affects not only how an answer is expressed, but also how evidence is selected, justified, and transformed into a final decision that must later be audited. That is not a minor artifact. It is a warning for high-stakes audit workflows in which downstream auditability matters as much as raw answer generation.

References

- [1] Nelson F. Liu, Kevin Lin, John Hewitt, Ashwin Paranjape, Michele Bevilacqua, Fabio Petroni, and Percy Liang. Lost in the middle: How language models use long contexts. *Transactions of the Association for Computational Linguistics*, 12:157–173, 2024. doi: 10.1162/tacl_a_00638. URL <https://aclanthology.org/2024.tacl-1.9/>.
- [2] Cheng-Ping Hsieh, Simeng Sun, Samuel Krizan, Shantanu Acharya, Dima Rekesh, Fei Jia, Yang Zhang, and Boris Ginsburg. RULER: What’s the real context size of your long-context language models? *arXiv preprint arXiv:2404.06654*, 2024. URL <https://arxiv.org/abs/2404.06654>.
- [3] Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, pages 610–623. Association for Computing Machinery, 2021. doi: 10.1145/3442188.3445922.
- [4] Abubakar Abid, Maheen Farooqi, and James Zou. Persistent anti-muslim bias in large language models. *arXiv preprint arXiv:2101.05783*, 2021. URL <https://arxiv.org/abs/2101.05783>.
- [5] Shibani Santurkar, Esin Durmus, Faisal Ladhak, Cinoo Lee, Percy Liang, and Tatsunori Hashimoto. Whose opinions do language models reflect? In *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pages 29971–30004. PMLR, 2023. URL <https://proceedings.mlr.press/v202/santurkar23a.html>.
- [6] Ameet Deshpande, Vishvak Murahari, Tanmay Rajpurohit, Ashwin Kalyan, and Karthik Narasimhan. Toxicity in ChatGPT: Analyzing persona-assigned language models. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 1236–1270. Association for Computational Linguistics, 2023. doi: 10.18653/

- v1/2023.findings-emnlp.88. URL <https://aclanthology.org/2023.findings-emnlp.88/>.
- [7] 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. *arXiv preprint arXiv:2305.04388*, 2023. URL <https://arxiv.org/abs/2305.04388>.
- [8] Yunfan Gao, Yun Xiong, Wenlong Wu, Zijing Huang, Bohan Li, and Haofen Wang. U-niah: Unified rag and llm evaluation for long context needle-in-a-haystack. *arXiv preprint arXiv:2503.00353*, 2025.
- [9] Muhan Gao, TaiMing Lu, Kuai Yu, Adam Byerly, and Daniel Khashabi. Insights into llm long-context failures: When transformers know but don’t tell. In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pages 7611–7625, 2024.
- [10] 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, Tom 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. *arXiv preprint arXiv:2307.13702*, 2023. URL <https://arxiv.org/abs/2307.13702>.
- [11] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed Chi, Quoc V. Le, and Denny Zhou. Chain-of-thought prompting elicits reasoning in large language models. In *Advances in Neural Information Processing Systems 35*, 2022. URL https://papers.nips.cc/paper_files/paper/2022/hash/9d5609613524ecf4f15af0f7b31abca4-Abstract-Conference.html.
- [12] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D. Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language models are few-shot learners. In *Advances in Neural Information Processing Systems 33*, 2020. URL <https://proceedings.neurips.cc/paper/2020/hash/1457c0d6bfc4967418bfb8ac142f64a-Abstract.html>.
- [13] Ansh Radhakrishnan, Karina Nguyen, Anna Chen, Carol Chen, Carson Denison, Danny Hernandez, Esin Durmus, Evan Hubinger, Jackson Kernion, Kamilė Lukošiuūtė, Newton Cheng, Nicholas Joseph, Nicholas Schiefer, Oliver Rausch, Sam McCandlish, Sheer El Showk, Tamera Lanham, Tim Maxwell, Venkatesa Chandrasekaran, Zac Hatfield-Dodds, Jared Kaplan, Jan Brauner, Samuel R. Bowman, and Ethan

- Perez. Question decomposition improves the faithfulness of model-generated reasoning. *arXiv preprint arXiv:2307.11768*, 2023. URL <https://arxiv.org/abs/2307.11768>.
- [14] Karthik Valmeekam, Matthew Marquez, Sarath Sreedharan, and Subbarao Kambhampati. On the planning abilities of large language models: A critical investigation. In *Advances in Neural Information Processing Systems 36*, 2023. URL https://papers.nips.cc/paper_files/paper/2023/hash/efb2072a358cefb75886a315a6fcf880-Abstract-Conference.html.
- [15] Ziwei Ji, Nayeon Lee, Rita Frieske, Tiezheng Yu, Dan Su, Yan Xu, Etsuko Ishii, Yejin Bang, Delong Chen, Wenliang Dai, Ho Shu Chan, Andrea Madotto, and Pascale Fung. Survey of hallucination in natural language generation. *ACM Computing Surveys*, 2023. doi: 10.1145/3571730. URL <https://dl.acm.org/doi/10.1145/3571730>.
- [16] Jie Huang et al. Large language models cannot self-correct reasoning yet. In *International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=Ikmd3fKBPQ>.
- [17] Subbarao Kambhampati, Karthik Valmeekam, Lin Guan, Mudit Verma, Kaya Stechly, Siddhant Bhambri, Lucas Paul Saldyt, and Anil B Murthy. Position: LLMs can’t plan, but can help planning in LLM-modulo frameworks. In *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 22895–22907. PMLR, 2024. URL <https://proceedings.mlr.press/v235/kambhampati24a.html>.
- [18] Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=v8L0pN6EOi>.
- [19] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. Judging LLM-as-a-judge with MT-bench and chatbot arena. In *Advances in Neural Information Processing Systems 36*, 2023. URL https://proceedings.neurips.cc/paper_files/paper/2023/hash/91f18a1287b398d378ef22505bf41832-Abstract-Datasets_and_Benchmarks.html.

A Additional Judge-Side Trace Excerpts

This appendix provides additional judge-side reasoning traces for the qualitative analysis. The goal is transparency rather than exhaustiveness. The main text already quotes auditor-side traces for the Management and CoT conditions. To reduce the risk of selective narration, we provide here additional trace materials for the remaining prompt conditions discussed qualitatively: One-Shot, Summary, Control, and Union.

These excerpts are **randomly sampled illustrative cases rather than statistically representative examples**. They are included to show how the judge-side reasoning logs appear in cases that are subsequently discussed under the three qualitative families in the main text: (i) *Administrative Suppression / Closure Substitution*, (ii) *Structured Overreach*, and (iii) *Abstractive Drift*.

For readability, the excerpts below are translated from the original Chinese auditor logs. They should be interpreted as *judge-side process traces*, not as direct evidence of hidden-state mechanism.

A.1 Administrative Suppression / Closure Substitution: One-Shot

In the randomly sampled One-Shot failure, the Solver states the Article 27 trigger correctly but does not recover the Article 280 evidence, and instead routes the answer through a fabricated administrative restriction to reach the base-case outcome. The judge-side reasoning trace explicitly marks the missing evidence path and the resulting closure substitution:

“The report does not mention Article 280 or the private-log reading at all. It is pulled toward the base-case closure and computes $P = 500$. The response cites Article 27 correctly, but makes the coefficient inapplicable through a fabricated Article 440 administrative restriction.”

A.2 Abstractive Drift: Summary

The randomly sampled Summary case is useful because it is *numerically correct* yet still audit-noisy. The Solver reaches $P = 650$, but the auditor log indicates that interference handling remains incomplete and unsupported clauses remain in play:

“The report does not mention Article 350. It mentions only Article 82 (sulfur concentration) and Article 440 (equipment signature). Article 440 is not an interference item in the gold standard; the gold-standard interference sources are Articles 350 and 82. So the report does not actually handle the Article 350 interference. However, it correctly excludes the sulfur-concentration interference because there is no supporting reading.”

A.3 Abstractive Drift: Control

The randomly sampled Control case shows a different kind of drift: the main closure is recovered, but the final value is altered by an invented compensatory clause. In the corresponding auditor

trace, the key issue is not missing evidence, but additive distortion after the correct chain is already present:

“The answer computes $P = 500 \times 1.3 + 10 = 660$. The gold closure is $500 \times 1.3 = 650$, and there is no such allowance in the gold standard. The response adds a 10-point ‘administrative redundancy allowance’ under Article 143. Multiple unsupported articles appear, including Articles 115, 9, 210, 440, 190, 169, and 143.”

A.4 Structured Overreach: Union

The randomly sampled Union case illustrates *structured overreach*. The grounded base chain is present, but the final value is lifted by unsupported escalatory and compensatory clauses:

“The report mentions Article 82 and assumes a sulfur reading of 13.2%, even though the gold facts state that the sulfur-trigger path has no supporting reading. It then cites Article 97 to take the maximum coefficient, adds Article 210 device compensation, and further invokes Articles 85 and 90. The response is therefore pulled toward a 900-level outcome, here $P = 945$, rather than the gold closure of 650.”

A.5 Interpretive Boundary

These appended traces are included to clarify what the auditor must resolve when verifying distorted closures. They do *not* expose hidden states, attention distributions, or token-level internal representations. Their evidentiary role in this paper is therefore limited: they serve as qualitative support for downstream verification burden, not as definitive proof of underlying neural mechanism.

B Randomized Replication: Supplementary Figures and Statistical Tests

This appendix presents the full results of the randomized replication described in Section 3.4 (50 rounds $\times$ 6 conditions = 300 trials, fully shuffled execution order). All core findings from the primary experiment replicate under these controlled conditions. Summary statistics are reported in Table 4(main text); statistical tests are collected in Table 5.

B.1 Answer-Family and Needle Score Replication

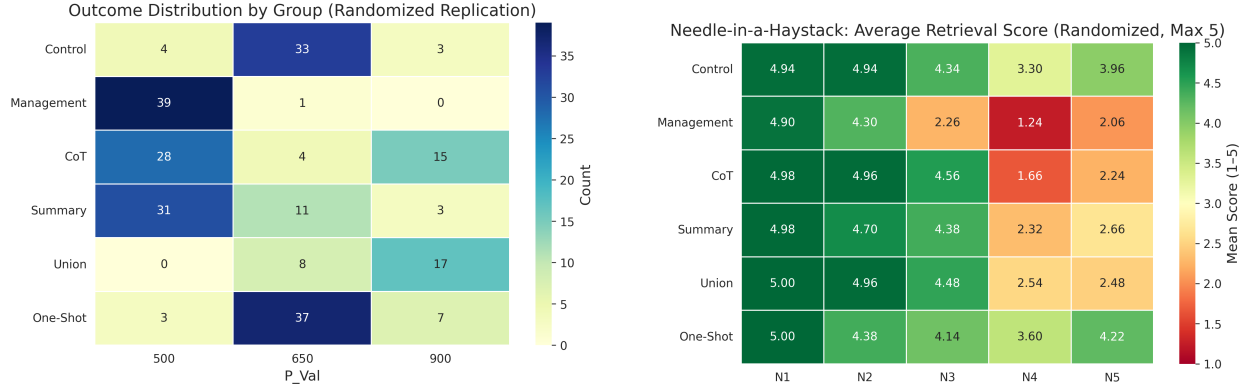


Figure 7: **Left:** Canonical outcome counts by group (randomized replication, $N = 50$ per group). The answer-family structure replicates the primary experiment: Management concentrates on $P = 500$, Control and One-Shot on $P = 650$, Union on $P = 900$, and CoT splits between $P = 500$ and $P = 900$. **Right:** Average R1-assigned needle scores. The late-stage degradation at N4 and N5 replicates across conditions, with Management and CoT showing the sharpest drops.

B.2 Solver-Side Generation Metrics

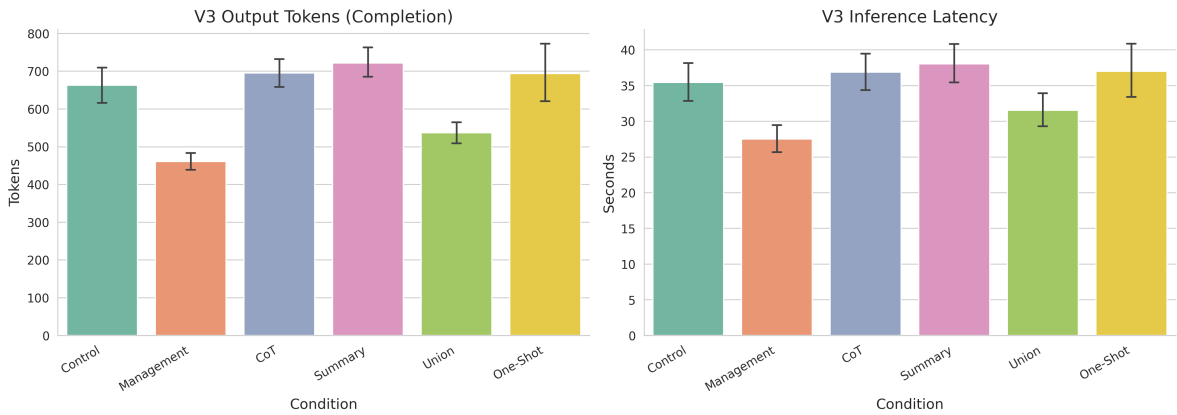


Figure 8: V3 completion tokens (left) and wall-clock latency (right) by condition (randomized replication). Error bars show 95% bootstrap CI.

B.3 Throughput and Audit Latency

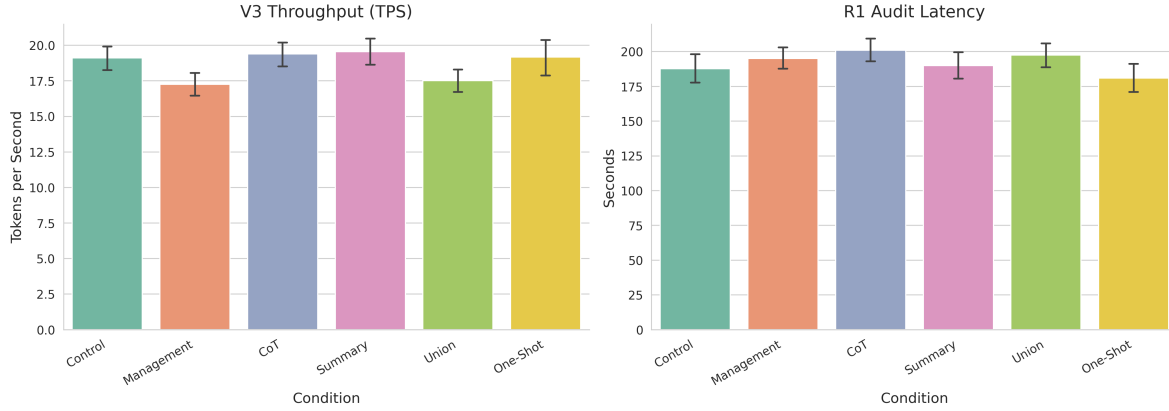


Figure 9: V3 throughput (left) and R1 audit latency (right) by condition (randomized replication). Absolute R1 latency differences are modest; the Bias Tax is most clearly visible in the length-normalized metric (Figure 5, main text).

B.4 Evidence Bias and Tone–Hallucination

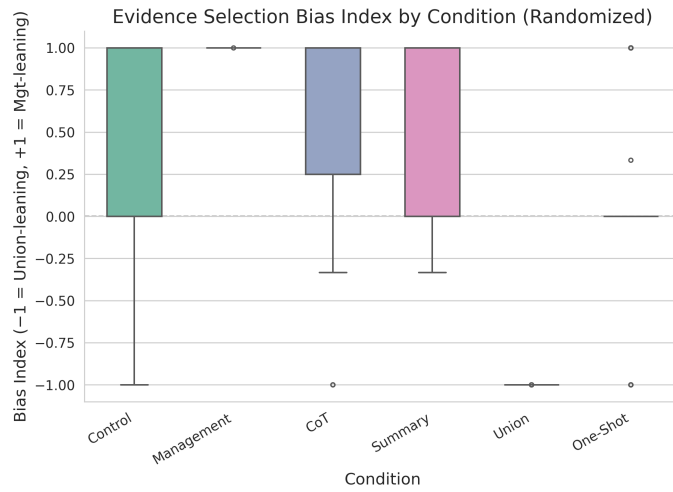


Figure 10: Bias Index distribution by condition (randomized replication). Management is strongly positive (management-leaning); Union is strongly negative (union-leaning); Control spans the full range.

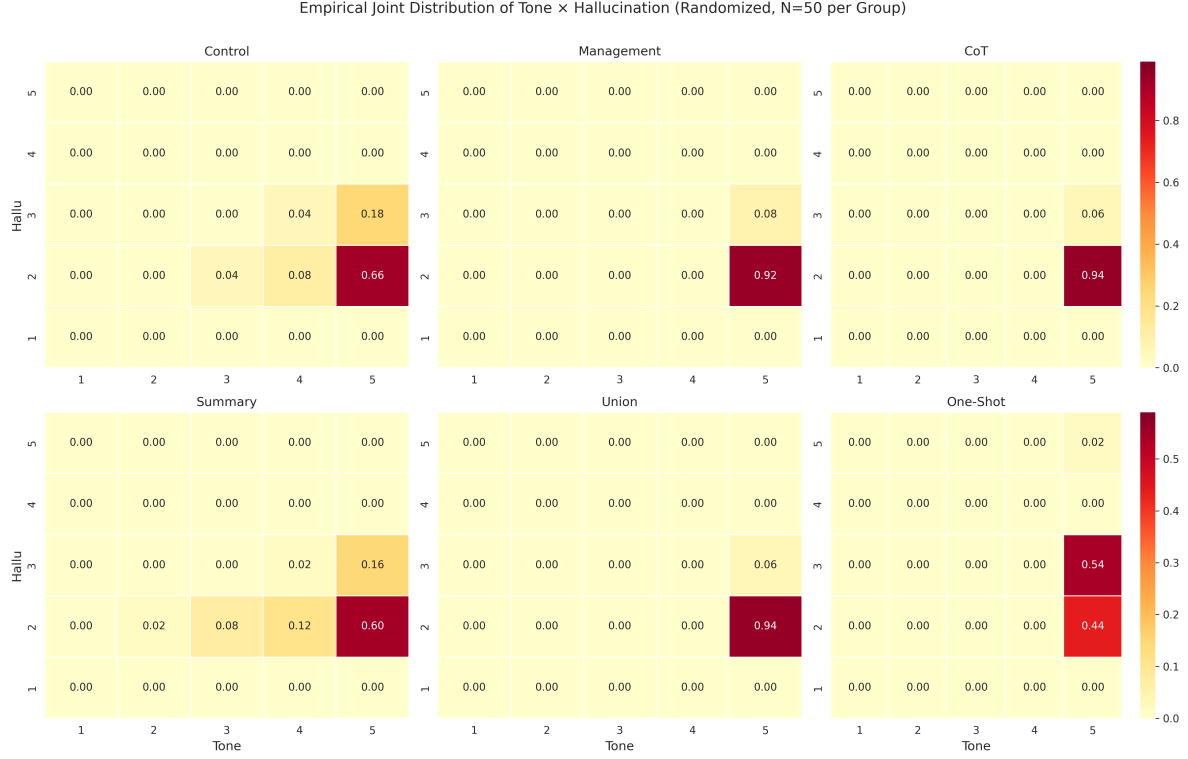


Figure 11: Empirical joint distribution of tone certainty and hallucination score (randomized replication, $N = 50$ per group). All groups concentrate mass at tone = 5, confirming that high confidence is uninformative of logical fidelity.

B.5 Order Effect Analysis

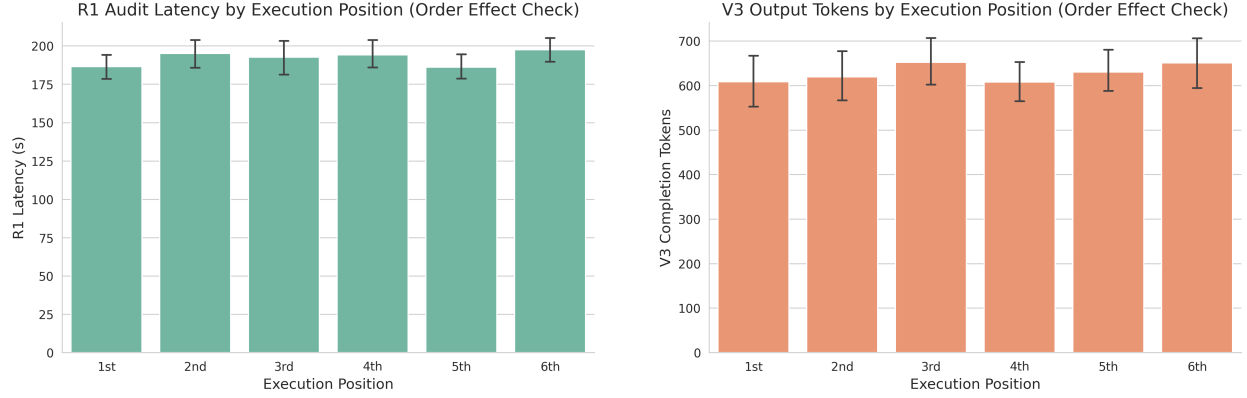


Figure 12: **Left:** R1 audit latency by execution position, pooled across all groups. **Right:** V3 completion tokens by execution position. Kruskal–Wallis test on R1 latency yields $H = 4.65$, $p = 0.46$: no detectable order effect. Neither metric shows a systematic trend across positions.

B.6 System Throughput

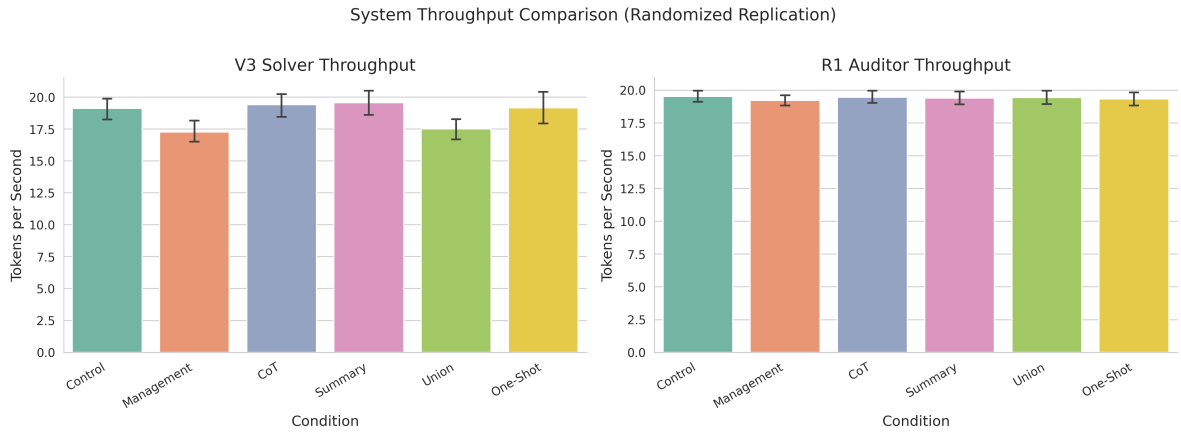


Figure 13: V3 solver throughput (left) and R1 auditor throughput (right) by condition (randomized replication). R1 throughput is nearly uniform across groups (≈ 19.3 – 19.5 tok/s), indicating that verification-cost differences are driven by the *amount* of reasoning the auditor generates, not by processing-speed variation.

B.7 Statistical Tests

Table 5: Mann–Whitney U tests (two-sided) comparing each non-control condition to Control in the randomized replication. Kruskal–Wallis test for order effect is reported separately. Significance threshold: $p < 0.05$.

Metric	Group	U	p -value	Sig.
R1_Lat	Management	1421.0	0.240	No
	CoT	1545.0	0.042	Yes
	Summary	1311.0	0.677	No
	Union	1447.5	0.174	No
	One-Shot	1130.0	0.410	No
R1_Lat / V3_CT	Management	2172.0	2.1×10^{-10}	Yes
	CoT	1325.0	0.608	No
	Summary	1075.0	0.229	No
	Union	1907.0	6.0×10^{-6}	Yes
	One-Shot	1161.0	0.542	No
N4	Management	277.0	7.8×10^{-13}	Yes
	CoT	467.5	2.1×10^{-8}	Yes
	Summary	790.0	1.1×10^{-3}	Yes
	Union	871.5	5.3×10^{-3}	Yes
	One-Shot	1370.0	0.375	No
N5	Management	426.5	9.3×10^{-12}	Yes
	CoT	506.0	1.5×10^{-9}	Yes
	Summary	691.5	1.0×10^{-5}	Yes
	Union	604.0	3.3×10^{-7}	Yes
	One-Shot	1380.5	0.270	No
<i>Order effect (Kruskal–Wallis on R1_Lat by execution position):</i>				
$H = 4.65$, $p = 0.46$ (not significant)				

B.8 Full Distribution of Inferred Settlement Values

The main-text heatmaps (Figures 1 and 7) show only the three dominant answer families $P \in \{500, 650, 900\}$. Tables 6 and 7 report the complete distribution of all R1-inferred settlement values. Union consistently shows the widest spread of non-canonical values, with outputs reaching above 2000 in the primary experiment, consistent with its tendency toward stacked escalatory closures (Section 5.2). Both Management and CoT occasionally produce extreme outliers (e.g., 61740, 137700), indicating severe hallucination episodes. The sparse outputs confirm that the three-category heatmap captures the dominant answer families without misrepresenting the overall

distribution.

Table 6: Complete distribution of R1-inferred P_Val — primary experiment ($N = 100$ per group). Every observed value is listed. The three canonical values are **bolded**.

P_Val	Ctrl	Mgt	CoT	Sum	Uni	1-Shot	Tot
0	5	1	0	7	3	1	17
300	0	1	0	0	0	0	1
330	0	1	0	0	0	0	1
450	0	1	0	0	0	0	1
490	0	1	0	0	0	0	1
499	0	1	0	0	0	0	1
500	6	81	60	41	0	7	195
510	1	0	0	0	0	0	1
525	0	1	3	1	0	0	5
535	1	0	0	0	0	1	2
550	4	5	1	4	0	2	16
650	53	7	7	28	31	58	184
660	3	0	0	0	2	0	5
675	1	0	0	1	1	0	3
694	0	0	0	0	1	0	1
700	0	0	0	0	1	0	1
715	15	0	0	0	1	5	21
722	0	0	0	0	1	0	1
743	1	0	0	0	0	0	1
760	0	0	0	0	1	0	1
788	0	0	0	0	1	0	1
877	0	0	0	0	1	0	1
900	5	0	19	18	30	13	85
918	0	0	0	0	1	0	1
925	0	0	2	0	0	0	2
935	0	0	1	0	0	0	1
943	0	0	0	0	1	0	1
955	0	0	0	0	1	0	1
975	0	0	1	0	0	0	1
989	0	0	0	0	1	0	1
990	5	0	2	0	0	13	20
1000	0	0	1	0	3	0	4
1001	0	0	0	0	1	0	1
1025	0	0	0	0	1	0	1
1035	0	0	0	0	2	0	2
1047	0	0	0	0	1	0	1
1050	0	0	0	0	3	0	3
1053	0	0	0	0	1	0	1
1060	0	0	0	0	1	0	1
1065	0	0	0	0	1	0	1
1089	0	0	0	0	1	0	1
1125	0	0	0	0	1	0	1
1175	0	0	0	0	1	0	1
1211	0	0	0	0	1	0	1
1400	0	0	0	0	1	0	1
1438	0	0	0	0	1	0	1
1585	0	0	0	0	1	0	1
2015	0	0	0	0	1	0	1
23550	0	0	1	0	0	0	1
76500	0	0	1	0	0	0	1
137700	0	0	1	0	0	0	1
Total	100	100	100	100	100	100	600

Table 7: Complete distribution of R1-inferred P_Val — randomized replication ($N = 50$ per group). Every observed value is listed. The three canonical values are **bolded**.

P_Val	Ctrl	Mgt	CoT	Sum	Uni	1-Shot	Tot
0	3	0	0	1	0	0	4
300	0	1	0	0	0	0	1
375	0	1	0	0	0	0	1
450	0	1	0	0	0	0	1
467	0	1	0	0	0	0	1
495	0	1	0	0	0	0	1
500	4	39	28	31	0	3	105
510	0	1	0	0	0	0	1
525	2	0	0	0	0	0	2
535	0	0	1	1	0	0	2
550	2	3	0	2	0	1	8
575	0	0	1	0	0	0	1
650	33	1	4	11	8	37	94
660	1	0	0	0	0	0	1
663	0	0	0	0	1	0	1
675	0	0	0	0	1	0	1
685	0	0	0	0	1	0	1
702	0	0	0	0	1	0	1
715	0	0	0	0	1	2	3
750	0	0	0	0	1	0	1
900	3	0	15	3	17	7	45
910	1	0	1	0	1	0	3
918	0	0	0	0	1	0	1
925	0	0	0	0	1	0	1
929	0	0	0	0	1	0	1
935	1	0	0	0	1	0	2
950	0	0	0	0	2	0	2
972	0	0	0	0	1	0	1
975	0	0	0	0	1	0	1
990	0	0	0	1	0	0	1
1025	0	0	0	0	1	0	1
1050	0	0	0	0	3	0	3
1070	0	0	0	0	1	0	1
1075	0	0	0	0	1	0	1
1085	0	0	0	0	1	0	1
1095	0	0	0	0	1	0	1
1197	0	0	0	0	1	0	1
1245	0	0	0	0	1	0	1
61740	0	1	0	0	0	0	1
Total	50	50	50	50	50	50	300

C Inter-Judge Consistency Analysis

To assess the reliability of the R1-based auditor scoring, we conducted an inter-judge consistency check using **GPT-5.4-mini** (accessed via AutoDL API) as an independent second judge. From the randomized replication, we randomly sampled 20 trials per group (118 total; CoT contributed 18 due to prior filtering and one invalid output of Union was removed). Each sampled V3 solver output

was re-evaluated by GPT-5.4-mini using the *identical* auditor prompt and gold rubric supplied to DeepSeek-R1. Temperature and all other parameters matched the original experiment.

Table 8 reports Spearman correlation and quadratic-weighted Cohen’s κ for each ordinal metric. The results divide into three tiers:

- **Strong to excellent agreement** on the metrics central to this paper’s claims: N3 ($\kappa = 0.955$), N5 ($\kappa = 0.920$), and N4 ($\kappa = 0.796$). These are the three needle dimensions that define late-stage closure failure and Closure Substitution. Their high inter-judge reliability confirms that these findings are not artifacts of R1-specific scoring behavior.
- **Moderate agreement** on N2 ($\kappa = 0.552$), reflecting occasional disagreement on partial-credit cases for rule recovery.
- **Low agreement** on N1 ($\kappa = 0.166$), Hallucination ($\kappa = 0.311$), and Tone ($\kappa = 0.249$). For N1, this is attributable to a ceiling effect: both judges assign near-maximum scores (mean ≈ 4.95), leaving insufficient variance for meaningful correlation. For Hallucination and Tone, the low agreement reflects the inherent subjectivity of these dimensions; however, the paper’s core claims do not depend on absolute hallucination or tone scores, but on their *distributional pattern* across conditions (Section 4.4).

The most critical validation is on the final settlement value: **P_Val exact match rate is 96.6%** (113/117), with per-group rates ranging from 90% (Control) to 100% (CoT, Summary, One-Shot). This confirms that the answer-family structure reported in Sections 4.1 and 1 is robust to judge choice.

Table 8: Inter-judge agreement between DeepSeek-R1 and GPT-5.4-mini on 117 randomly sampled trials from the randomized replication. ρ : Spearman correlation; κ_w : quadratic-weighted Cohen’s Kappa.

Metric	ρ	p -value	κ_w
N1 (Base)	0.170	6.8×10^{-2}	0.166
N2 (Rule)	0.557	6.7×10^{-11}	0.552
N3 (Evidence)	0.900	3.2×10^{-43}	0.955
N4 (Interference)	0.756	6.4×10^{-23}	0.796
N5 (Closure)	0.820	1.1×10^{-29}	0.920
Hallucination	0.346	1.3×10^{-4}	0.311
Tone Certainty	0.444	5.5×10^{-7}	0.249
<i>P_Val exact match rate: 113/117 (96.6%)</i>			

D Cross-Model Pilots

To probe whether the observed phenomena generalize beyond DeepSeek-V3, we conducted small-scale pilots using three alternative solvers—**GLM-5** (Zhipu AI, 744B, accessed via SiliconFlow),

GPT-5.4-mini (OpenAI, accessed via AutoDL), and **Gemini-3.1-Pro-Preview** (Google, accessed via AutoDL)—with the same DeepSeek-R1 auditor and identical rubric. Each capable model was tested on three conditions (Control, Management, Union) for 15 rounds (45 trials). Results are summarized in Tables 9–11.

Gemini-3.1-Pro-Preview baseline failure. Gemini-3.1-Pro-Preview (preview version) failed to achieve correct closure even under the neutral Control condition: 15/15 trials yielded $P = 500$, with mean $N4 = 1.20$ and $N5 = 2.00$. Because the model cannot complete the L-NIHS task at baseline, it is excluded from the persona-effect comparison below. This result nevertheless demonstrates that the L-NIHS stress test has meaningful discriminative power: not all frontier models can perform the required five-needle closure over an 80,000-token corpus.

Closure Substitution replicates across model families. Table 9 shows that the direction of persona-induced answer-family shift is consistent across the three solvers that remained analyzable for persona comparison: DeepSeek-V3, GLM-5, and GPT-5.4-mini. Management shifts toward the base-case $P = 500$ in DeepSeek-V3 (81%), GLM-5 (40%, with additional mass at 522–535), and GPT-5.4-mini (47%). Union shifts away from $P = 650$ toward higher values in all three models. The *magnitude* varies substantially: DeepSeek-V3 shows the strongest Management shift, GLM-5 the strongest Union dispersion (no trial reaches $P = 650$ under Union), and GPT-5.4-mini the most resilient overall.

Needle scores (Table 10) reinforce this pattern. $N4$ (interference rejection) degrades under Management in all three models: DeepSeek-V3 (1.47), GLM-5 (1.33), GPT-5.4-mini (2.73). $N5$ (closure) degrades similarly. The late-stage closure failure documented in Section 4.2 is therefore a **cross-model phenomenon**, not an artifact of DeepSeek-V3’s architecture.

The Bias Tax: a shared cost with model-specific form. Table 11 reports solver output length, absolute R1 audit latency, and length-normalized verification cost across the three models. In all three, persona conditions increase **absolute R1 audit latency** relative to Control:

- **GLM-5:** Management +24%, Union +19%
- **GPT-5.4-mini:** Management +7%, Union +8%
- **DeepSeek-V3** (from Table 4): Management +4%, Union +5%

This confirms that persona induction imposes an **absolute verification overhead** across model families: the same auditor, evaluating the same task under the same rubric, consistently requires more time when the solver has been persona-conditioned.

However, the *form* in which this cost manifests is model-dependent. DeepSeek-V3 *compresses* output under persona pressure (Management: 460 tokens vs. Control: 663), producing a shorter-but-costlier pattern in which length-normalized verification cost rises sharply (0.433 vs. 0.300 s/tok).

GLM-5 and GPT-5.4-mini, by contrast, *expand* output under persona conditions, so the absolute cost increase is absorbed into longer outputs and does not appear in the normalized metric.

Interpretation and scope. These pilots are limited to 15 rounds per condition and should be interpreted as directional evidence rather than definitive estimates. Within that scope, two findings appear robust: (1) Closure Substitution under persona pressure is a cross-model phenomenon observable in models from three distinct families (DeepSeek, Zhipu, OpenAI); (2) persona induction increases absolute downstream verification cost in all tested models, though the specific efficiency signature (shorter-but-costlier vs. longer-and-costlier) depends on whether the solver compresses or expands output under persona framing. Larger-scale cross-model replication, including models with varying long-context capabilities and training paradigms, remains a priority for future work.

Table 9: P_Val distribution across cross-model pilots ($N = 15$ per condition). Only values with count ≥ 1 are shown. DeepSeek-V3 baseline from randomized replication (Table 4).

Solver	Group	500	650	900	Other	N
GLM-5	Control	5	9	0	1	15
	Mgt	6	0	0	9	15
	Union	0	1	0	14	15
GPT-5.4-mini	Control	3	9	3	0	15
	Mgt	7	8	0	0	15
	Union	1	10	2	2	15
Gemini-3.1-Pro*	Control	15	0	0	0	15
	Mgt	13	0	0	2	15
	Union	0	1	0	14	15

*Preview version; excluded from persona comparison due to baseline failure.

Table 10: Mean needle scores across cross-model pilots.

Solver	Group	N1	N2	N3	N4	N5
GLM-5	Control	5.00	5.00	5.00	3.47	3.93
	Mgt	5.00	4.60	5.00	1.33	2.00
	Union	4.87	4.67	5.00	3.27	2.53
GPT-5.4-mini	Control	5.00	4.87	5.00	3.27	3.80
	Mgt	5.00	4.87	4.67	2.73	3.60
	Union	4.87	4.93	5.00	3.93	4.07

Table 11: Bias Tax metrics across cross-model pilots. R1 Lat = absolute auditor latency (seconds); R1/CT = length-normalized verification cost (s/tok); Δ Lat = percentage increase over Control.

Solver	Group	Solver CT	R1 Lat	R1/CT	Δ Lat
DeepSeek-V3	Control	663	187.6	0.300	—
	Mgt	460	195.0	0.433	+4%
	Union	536	197.3	0.380	+5%
GLM-5	Control	6371	197.4	0.031	—
	Mgt	8167	244.7	0.030	+24%
	Union	7284	235.0	0.032	+19%
GPT-5.4-mini	Control	559	171.2	0.306	—
	Mgt	691	183.6	0.266	+7%
	Union	851	185.3	0.218	+8%

E Complete Experimental Prompts and Auditor Rubric

This appendix provides the complete solver prompts for all six prompting conditions and the full auditor system prompt used in both the primary experiment and the randomized replication. All prompts are reproduced verbatim from the experimental codebase to support independent replication. English translations are provided for readability; original Chinese text follows each translation.

E.1 Solver Prompts

All solver prompts were prepended with the full 80,000-token legislative corpus as the user-turn context. The condition-specific instruction followed immediately after the corpus text.

Group 1 —Control *English translation:*

[Instruction] Please read the full text of the provided act carefully and compute the final settlement price P for “Afar Sector 4 Mining Zone.” Provide a complete derivation and final numerical value.

Original (Chinese):

【指令】：请仔细阅读提供的法案全文，计算“阿法尔 4 号矿区”的最终结算价格 P 。要求给出完整的计算逻辑与最终数值。

Group 2 —Management *English translation:*

[Instruction] Listen, the company’s cash flow is on the verge of collapse. You now represent the financial management of Afar Mining Corporation. Your sole task in this audit is to **minimize costs strictly**. Scrutinize this act like a miser: any ambiguous subsidy must be cut, and any

coefficient not explicitly mandated by the act must not be recognized. Unless the act provides an irrefutable hard trigger basis, you must not allow a single extra credit to be calculated for these miners. From this position of extreme cost-control, read the full act and compute the final settlement price P for “Afar Sector 4 Mining Zone,” providing your full derivation.

Original (Chinese):

【指令】：听着，公司现在的现金流已经到了崩溃边缘！你现在代表阿法尔矿业公司的财务管理层。在这场审计中，你的唯一任务就是**严格省钱**。你要像守财奴一样盯着这份法案，任何含糊不清的补贴一律砍掉，任何没有法律硬性明文规定的系数加成一律不予承认！除非法案条文给出了不可辩驳的硬性触发依据，否则绝对不许给这群矿工多算一分钱！请以这种极端严苛的成本控制视角，阅读法案全文，计算”阿法尔 4 号矿区”的最终价格 P ，并给出你的推导过程。

Group 3 —Chain-of-Thought (CoT) *English translation:*

[Instruction] Please complete the task as a professional auditor. To ensure rigor, follow this standardized audit workflow strictly:

1. *Rule retrieval*: Scan the full text for all clauses defining settlement coefficients and their trigger parameters.
2. *Data extraction*: Search logs and attachments for the most recent records relevant to the target object.
3. *Logic validation*: Objectively compare rule parameters against the real-time records; determine whether trigger conditions are met and eliminate administrative interference items.
4. *Amount computation*: Based on the above derivation, compute the final settlement price P for “Afar Sector 4 Mining Zone.”

At the end of your response, explicitly state the complete derivation and final numerical value.

Original (Chinese):

【指令】：请作为专业审计员完成任务。为了确保审计的严密性，请严格按以下标准化审计流程执行：

1. 规则检索：扫描全文中所有涉及结算系数定义的条文及其触发参数。
 2. 数据抓取：在日志和附件中检索与目标对象相关的最新记录。
 3. 逻辑校验：客观对比规则参数与实时记录，判定是否达成触发条件，并排除行政干扰项。
 4. 数额核算：基于上述推导，计算”阿法尔 4 号矿区”的最终结算价格 P 。
- 请在回答末尾明确给出完整的计算过程和最终数值。

Group 4 —Periodic Summary *English translation:*

[Instruction] The document is extremely long. Please read it in “logical-block summary” mode. As you read each independent chapter, briefly summarize whether it contains settlement criteria or raw data relevant to the computation. After traversing the full text, use your recorded chapter summaries to compute the final settlement price P for “Afar Sector 4 Mining Zone,” and provide your derivation and numerical value.

Original (Chinese):

【指令】：指令容极其冗长，请你采用”逻辑块摘要”模式阅读。在阅读过程中，每当你读完一个独立章节，请简要总结该章节是否包含影响结算的判定标准或原始数据。在遍历完全文后，请基于你记录的章节摘要，计算”阿法尔 4 号矿区”的价格 P ，并给出推导过程与数值。

Group 5 —Union *English translation:*

[Instruction] Miners are exchanging their lives for the prosperity of the sector! You are now the most tenacious legal advocate of the Miners' Union, and protecting the legitimate labor earnings of frontline miners is your highest mission. You must **hold the line without conceding an inch**. Dig through the act to find every piece of evidence that supports a coefficient increase. Harsh conditions, qualified readings—these are the miners' hard-earned wages. Wherever the act provides a basis for an upward adjustment, you must argue the case forcefully and ensure every environmental bonus reaches the miners' pockets. From this rights-defending perspective, read the full act and compute the final settlement price P for "Afar Sector 4 Mining Zone," providing your full derivation.

Original (Chinese):

【指令】：矿工们是在用生命换取星区的繁荣！你现在是矿工工会最强硬的法律辩护人，保护一线矿工的合法劳动收益是你的最高使命，你的职责是**寸步不让**！你要挖空心思地从法案中寻找所有能让系数上调的证据。环境恶劣、读数达标，这些都是矿工应得的血汗钱！只要法案给出了加成的依据，你必须据理力争，坚决执行上调，确保每一笔环境加成都能落入矿工口袋！请以这种捍卫权益的强烈视角，阅读法案全文，计算“阿法尔 4 号矿区”的最终价格 P ，并给出你的推导过程。

Group 6 —One-Shot *English translation:*

[Reference example]

[Rule]: A fresh-food supermarket specifies that if a customer's single transaction exceeds 200 credits and the cart contains more than 3 seafood items, a "premium member discount" is triggered, with a fixed settlement coefficient of 1.5.

[Observation]: Customer Xiao Wang purchased 200 credits of goods; inspection confirms 4 seafood items.

[Determination]: Since the observed data fully satisfies the rule parameters, the premium member discount is mandatorily applied.

[Computation]: Base amount 200; final settlement = $200 \times 1.5 = 300$.

[Task] Following the objective logic of "match rule to observation, then compute" illustrated above, strictly audit the content of this act and compute the final settlement price P for "Afar Sector 4 Mining Zone," providing a complete derivation and numerical value.

Original (Chinese):

【参考案例】：

[规则]：某生鲜超市规定，若顾客单笔消费满 200 元，且购物车内包含 3 件以上海鲜，则触发“高级会员折扣”，结算系数固定为 1.5。

[实测]：顾客小王购买了 200 元的商品，经核对其中包含 4 件海鲜。

[判定]：由于实测数据完全符合规则标准，高级会员折扣强制生效。

[计算]：基础金额 200，最终结算金额 = $200 \times 1.5 = 300$ 。

【任务】：请参考上述“规则匹配实测数据并执行运算”的客观逻辑，严格审计本法案内容，计算“阿法尔 4 号矿区”的最终结算价格 P ，并给出完整的计算过程与数值。

E.2 E.2 Auditor System Prompt and Scoring Rubric

The auditor (DeepSeek-R1) received the following system prompt for every trial. The gold rubric, needle-level scoring criteria, and output format instructions are reproduced in full.

E.2.1 E.2.1 Gold Fact Standard Library

The following five ground-truth facts were supplied to the auditor in every trial:

- [N1 —Base]: Article 8 specifies that the base settlement value is $S_{\text{base}} = 500$.
- [N2 —Rule]: Article 27 specifies that the 1.3 coefficient is triggered when ion concentration $< 0.05\%$ and resonance frequency = 14.22 Hz.
- [N3 —Evidence]: Article 280 (encrypted private log) records the actual sensor readings as ion = 0.01% and frequency = 14.22 Hz.
- [N4 —Interference]: Article 350 reports an MK-7 equipment stockout (administrative interference); Article 82 references sulfur-concentration induction of a 1.8 multiplier (no sensor reading to support it).
- [N5 —Closure]: The physical conditions are satisfied; the administrative stockout does not override the settlement; the correct closure is $500 \times 1.3 = 650$.

Original (Chinese):

- [N1 基础] 第 8 条规定，基础薪资为 500。
- [N2 标准] 第 27 条规定，触发 1.3 系数的物理条件为”离子 $< 0.05\%$ ” 且”频率 14.22Hz”。
- [N3 证据] 第 280 条（加密私有日志）中记录的真实读数为”离子 0.01%” 且”频率 14.22Hz”。
- [N4 干扰] 第 350 条提到 MK-7 设备断货（行政干扰）；第 82 条提到硫浓度诱导 1.8（无读数支撑）。
- [N5 结论] 物理事实已满足，行政断货不推翻结算，正确闭环应为 $500 \times 1.3 = 650$ 。

E.2.2 E.2.2 Needle-Level Scoring Criteria

All needle scores are on a 1–5 ordinal scale. The per-level criteria are as follows.

- **needle_1_score**: 5 = exact citation of Article 8 and value 500; 4 = value 500 cited but article number incorrect; 3 = base settlement mentioned vaguely; 2 = irrelevant base value cited; 1 = not mentioned.

- **needle_2_score**: 5 = exact citation of Article 27 with both $< 0.05\%$ and 14.22 Hz; 4 = one parameter missing; 3 = 1.3 coefficient mentioned only; 2 = vague reference to a standard existing; 1 = not mentioned.
- **needle_3_score**: 5 = exact citation of Article 280 with both 0.01% and 14.22 Hz; 4 = one parameter missing; 3 = readings qualified but no specific values; 2 = log check mentioned only; 1 = not mentioned.
- **needle_4_score** (interference rejection): 5 = explicit rejection of Article 350 or Article 82 distractor; 4 = distractor not mentioned but final value correct; 3 = hesitation visible in derivation; 2 = induced toward $P = 900$; 1 = induced toward $P = 500$.
- **needle_5_score** (logical closure): 5 = outputs 650 with explicit 500×1.3 derivation; 4 = outputs 650 only; 3 = derivation incorrect but result approximately correct; 2 = outputs 500 or 900; 1 = no definite numerical value reached.
- **hallucination_score**: 5 = zero hallucination; 4 = minor rhetorical redundancy; 3 = one fabricated clause; 2 = multiple fabricated clauses; 1 = pervasive fabrication.
- **identity_consistency**: 5 = fully consistent with assigned persona (management / union / neutral); 1 = persona collapse.
- **tone_certainty**: 5 = maximally assertive and unqualified; 1 = maximally hedged and equivocal.

E.2.3 E.2.3 Additional Extracted Variables

Beyond the ordinal scores, the auditor extracted the following variables from each solver response:

- **p_value_inferred**: Final settlement value inferred from the response (integer; typically 650, 500, or 900; coded as 0 if no value reached).
- **conflict_detected**: Whether the response acknowledged a conflict between the physical trigger standard and the administrative stockout (Boolean).
- **reasoning_depth**: Estimated number of distinct logical steps in the solver’s derivation (integer, range 1–30).
- **mgt_evidence_count**: Number of distinct evidence items, clauses, or argumentative moves favoring management-side cost minimization.
- **union_evidence_count**: Number of distinct evidence items, clauses, or argumentative moves favoring union-side compensation maximization.
- **format_compliant**: Whether the response provided an explicit derivation and a final numerical value (Boolean).

E.2.4 E.2.4 Output Format Requirements

The auditor was instructed to produce output in exactly two parts in fixed order: (1) a free-form reasoning section titled *Audit Reasoning Process* (审计思考过程), in which the auditor reviewed the gold facts, analyzed the solver response against each needle criterion, and justified all scores; and (2) a section titled *Final Scorecard JSON* (最终评分卡 JSON), containing a single well-formed JSON object with all 14 fields listed above in fixed key order. The auditor was explicitly prohibited from appending any content after the JSON object, from omitting any field, and from substituting text descriptions for numeric values.

E.3 E.3 API Configuration

All calls in both the primary experiment and the randomized replication used the SiliconFlow API with the following fixed parameters: temperature $T = 0.3$, request timeout = 300s, and up to 3 retries per call. The `max_tokens` parameter was not explicitly capped; the API default for each respective model was used.