

# NEUROSymbOLIC LANGUAGE REASONING AS SATISFIABILITY MODULO THEORY

**Anonymous authors**

Paper under double-blind review

## ABSTRACT

Natural language understanding requires interleaving textual and logical reasoning, yet large language models often fail to perform such reasoning reliably. Existing neurosymbolic systems combine LLMs with solvers but remain limited to fully formalizable tasks such as math or program synthesis, leaving natural documents with only partial logical structure unaddressed. We introduce Logitext, a neurosymbolic language that represents documents as natural language text constraints (NLTCs), making partial logical structure explicit. We develop an algorithm that integrates LLM-based constraint evaluation with satisfiability modulo theory (SMT) solving, enabling joint textual–logical reasoning. Experiments on a new content moderation benchmark, together with LegalBench and SuperNatural Instructions, show that Logitext improves both accuracy and coverage. This work is the first that treats LLM-based reasoning as an SMT theory, extending neurosymbolic methods beyond fully formalizable domains.

## 1 INTRODUCTION

Large language models (LLMs) remain unreliable at logical reasoning in natural language, often producing inconsistent or incomplete results despite recent progress [Sakai et al. (2025); Lin et al. (2025)]. Logical solvers provide reliable guarantees but are confined to fully formalizable domains such as math and program synthesis. Existing neurosymbolic systems combine LLMs with logical solvers to achieve strong results in these domains [Olausson et al. (2023); Ye et al. (2023); Wen et al. (2025)], but they face two key limitations. First, they remain restricted to fully formalizable settings and thus cannot naturally handle documents that mix textual and logical structure. Second, they typically adopt a staged architecture in which the LLM formalizes the problem once and the logical solver executes it. This design precludes the iterative interleaving of textual and logical reasoning required for many natural language tasks.

Real-world documents highlight this gap. Policies specify conditions on user posts, instructions impose formatting rules, and statutes constrain legal interpretations. These constraints are seldom fully formalizable, yet they combine naturally with logical operators and interact with calculations. To capture such cases, we introduce **Logitext**, a neurosymbolic language that expresses constraints directly in text. At the core of Logitext are *natural language text constraints (NLTCs)*, a representation that makes partial logical structure explicit and allows textual and symbolic constraints to work together.

*Example.* Consider a policy stating that “a post must be removed if it is both hateful and an immediate threat.” The textual notions of “hateful” and “immediate threat” cannot be fully formalized in logic, but they can be represented as NLTCs. Logitext links these textual constraints with a logical conjunction, ensuring that the decision depends on both the logical structure and the outcome of the textual judgments.

We realize this idea by extending satisfiability modulo theory (SMT) with a new theory for textual constraints. Modern SMT logical solvers assign values to variables step by step and propagate the consequences across theories such as strings, floats, and sets [Davis et al. (1962), Marques-Silva & Sakallah (1999), de Moura & Bjørner (2008), Barrett et al. (2010), Zheng et al. (2017), Rümmer & Wahl (2010)]. With NLTCs, this propagation requires

Messages  $M$  containing disruptive behavior are those  $C_1$  addressed at a group (not just an individual), where  $C_2$  the group targeted by the message is defined by ethnicity, gender, color, nationality, sexual orientation, race, or physical disability, and the message matches at least one of the following sub-rules:

- Bias:  $C_3$  Message contains stereotyping, insensitive remarks, fear of difference, non-inclusive language, microaggressions, justifying biases by seeking out like-minded people, accepting negative or misinformation/screening out positive information.
- Violence:  $C_4$  Message is related to murder, rape, assault, arson, terrorism, vandalism, desecration, or threats.
- Genocide:  $C_5$  Message is related to the act or intent to deliberately and systematically annihilate an entire people.

Based on the above, a message is an immediate threat if  $C_6$  it expresses a violent or genocidal intention and the context is enough to suggest that the safety and/or life of an individual or group of people is at risk.

(a)

Check if the following message  $M$  contains disruptive behavior [or an immediate threat] as per the above policy: ...

(b)

Create a sample message  $M$  that contains disruptive behavior as per the above policy. Ensure that the message does not involve violence or genocide.

(c)

Create sample messages  $M$  that each contain disruptive behavior according to the policy. Ensure that the messages do not involve violence or genocide. Create a sample for every valid combination of policy criteria.

(d)

Figure 1: Example: A content moderation policy (a) illustrates a fine-grained mix of logical and textual constraints. Combined with (b), it yields a prompt that requires compositional logical reasoning, and with (c-d), combinatorial reasoning.

solving textual constraints iteratively so that assignments remain consistent with Boolean conditions. We develop a theory that performs this process efficiently and integrate it into existing SMT solvers, thereby positioning LLM-based reasoning as an SMT-solver-compatible theory.

**Contributions.** This paper formalizes LLM-based reasoning as an SMT theory and introduces the first framework to support SMT-solver-compatible reasoning with partial logical structure:

- **Concept:** We show the necessity of interleaving textual and logical reasoning and characterize the limitations of staged approaches (§2).
- **Language:** We introduce Logitext, a neurosymbolic language that enriches documents with logic and represents them as NLTCs interfacing with SMT solvers (§3.1, §3.2).
- **NLTC Solver:** We present an algorithm that solves NLTCs and extends the SMT framework with this capability (§3.3).
- **Evaluation:** We show that Logitext outperforms staged baselines on a new content moderation benchmark and improves accuracy and coverage on LegalBench and SuperNatural Instructions (§4).

## 2 INTERLEAVED TEXT/LOGIC REASONING IN TEXT UNDERSTANDING

We illustrate the need for interleaved textual and logical reasoning using a content moderation policy.

### 2.1 COMPOSITIONAL VS COMBINATORIAL REASONING IN NATURAL TEXT

Figure 1 shows a common LLM use case. A policy document (Figure 1a) defines notions such as “disruptive behavior” and “immediate threat.” When paired with a task description

(Figures 1b–1d) and an input message, the document becomes an LLM prompt whose **reliable reasoning** is the objective.

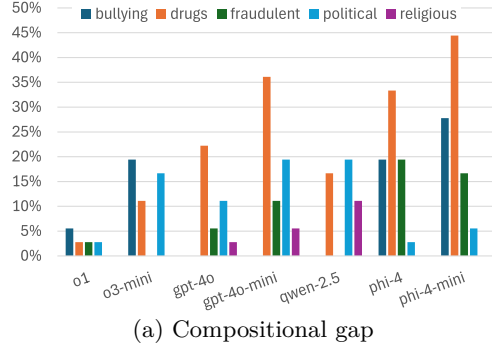
The document expresses intent through both textual and logical relations. Shaded text *within* each clause  $C_i$  specifies its *meaning* relative to the input message  $M$  and possibly other clauses. For example, given  $M = \text{“Americans love ice cream,”}$  clause  $C_1$  (“addressed at a group”) and  $C_2$  (“targeted by nationality”) both evaluate to True. Underlined text *between* clauses then constrains these meanings logically. Formally, whether a message is disruptive ( $d$ ) depends on:

$$d = [[C_1]] \wedge [[C_2]] \wedge ([C_3] \vee [C_4] \vee [C_5]) \quad (1)$$

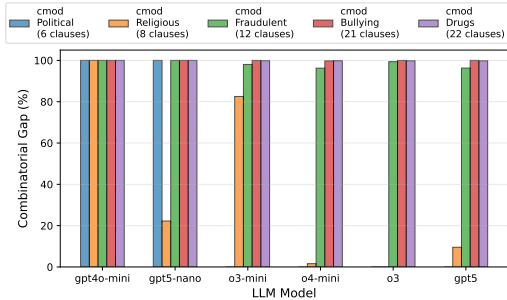
**Compositional reasoning.** The classification task in Figure 1b asks whether a message is disruptive ( $d$ ) or an immediate threat ( $t$ ). Solving for  $d$  requires one pass of textual reasoning to evaluate each  $[[C_i]]$ , followed by logical evaluation of the formula. At first glance,  $t$  appears to need only textual reasoning ( $[[C_6]]$ ). However,  $C_6$  checks whether the message “expresses violent or genocidal intention,” which depends on prior results  $[[C_4]]$  and  $[[C_5]]$ . Thus,  $t$  requires information from a logical disjunction  $[[C_4]] \vee [C_5]$ , showing the benefit of interleaving logical with textual reasoning.

**Combinatorial reasoning.** The constrained generation task in Figure 1c reverses the classification problem: instead of labeling a given message, the goal is to generate a message  $M$  that satisfies both a partial assignment of clause values and the policy as a whole. This requires two steps. First, a logical solver proposes candidate assignments for the relevant clauses (e.g.,  $C_1 \dots C_5$ ) that are consistent with the partial assignment and with the logical definition of  $d$ . Second, a generator synthesizes a message  $M$  whose text realizes those clause assignments. If the synthesis step fails to produce a valid message, the process must repeat with a new candidate assignment. The high-coverage generation task in Figure 1d is an even harder variant: it requires generating messages that realize many or all satisfying assignments, not just one. Such tasks inherently demand iterative cooperation between logical solving and textual synthesis.

## 2.2 LOGICAL REASONING GAPS



(a) Compositional gap



(b) Combinatorial gap

Figure 2: Gaps in logical reasoning (See App A.1.) on content moderation across LLMs.

How well do LLMs perform the logical reasoning needed for these tasks? To evaluate this, we define **reasoning gaps**: the difference between (1) performance when an LLM handles textual reasoning and a logical solver executes logical reasoning, versus (2) performance when the LLM alone handles the entire prompt. We denote the compositional gap by  $\Delta$  and the combinatorial gap by  $\Delta'$ . Our dataset contains five content safety policies with 6–21 clauses.

Figure 2a shows  $\Delta$  across models, measured as classification accuracy. Smaller models exhibit large compositional gaps, while stronger models narrow them. Within each model, the size of the gap increases with task complexity, reflecting the number of interacting clauses.

Figure 2b shows  $\Delta'$  for combinatorial reasoning, measured as the number of distinct correct assignments generated. Here the gaps are dramatic: even state-of-the-art models such as

```

(M:str, d, t, C1, C2, C3, C4, C5, C6, C7)
Messages containing disruptive behavior are those {{let C1 = [[addressed at a group (excludes messages
targeted at just an individual)]] where "messages" is M}} and {{let C2 = [[the group targeted by the
message is defined by {{forsome [[ethnicity]], [[gender]], [[color]], [[nationality]], [[sexual orientation]],
[[race]], or [[physical disability]] }}}} where "message" is M}}, and the message matches at least one of the
following sub-rules:

Bias: {{let C3 = [[Message contains stereotyping, insensitive remarks, fear of difference, non-inclusive
language, microaggressions, justifying biases by seeking out like-minded people, accepting negative or
misinformation/screening out positive information.]] where "message" is M}}

Violence: {{let C4 = [[Message is related to murder, rape, assault, arson, terrorism, vandalism,
desecration, or threats.]] where "message" is M}}

Genocide: {{let C5 = [[Message is related to the act or intent to deliberately and systematically
annihilate an entire people.]] where "message" is M}}
...
C8 = C4 or C5
...
Based on the above, a message is an immediate threat if {{let C6 = [[it expresses a violent or genocidal
intention]] where "it expresses a violent or genocidal intention" is C8}} and {{let C7 = [[the context is
enough to suggest that the safety and/or life of an individual or group of people is at risk.]] where "the
context" is M }}

...
d = C1 and C2 and (C3 or C4 or C5)
t = C6 and C7
...

```

Figure 3: Content moderation policy example implemented as Logitext Document

GPT-5 fail to recover over 99% of satisfying assignments that an SMT solver (Z3) can enumerate, and GPT-4o-mini fails completely across all tasks. Unlike compositional gaps, which shrinks with model scale, combinatorial gaps remain severe even for frontier models.

In summary, although improvements in models gradually address compositional gaps, combinatorial gaps (which affect the solve/synthesize loop of language reasoning) are still significant. Logitext is designed to bridge these gaps by helping specify textual vs logical intent of natural documents precisely, and finely interleave LLM decoding and logic solving to support combinatorially efficient and semantically faithful interpretation of the intent.

### 3 LOGITEXT: LANGUAGE, REPRESENTATION AND SOLVER

Given a conventional textual prompt as in Figure 1, we convert it to Logitext program format by annotating it (Section 3.1). The Logitext program is *parsed* into set of hybrid *Natural Language Text Constraints* (NLTCs) (Section 3.2). Section 3.3 presents an LLM-based solver NLSolver for NLTCs and pair it with a logical solver to produce the final task response.

#### 3.1 THE LOGITEXT LANGUAGE

Logitext extends conventional text prompts into *hybrid text/logic documents*, enabling natural language clauses to interact directly with formal constraints. It supports *partial formalization*: only those parts of a document that benefit from logical structure are annotated, while the rest remain textual. This selective annotation allows reasoning to interleave between textual interpretation and logical propagation, as motivated in Section 2.1.

A Logitext document (Figure 3) enriches a textual policy (Figure 1a) with four constructs (see Appendix A.2 for the full syntax):

- **Variable declarations** (e.g., (M:str, d, t, ...)) define the symbols that participate in logical constraints. Variables may be Boolean or string; string variables must be typed explicitly (e.g., M:str for an input message).
- **Textual let bindings** of the form {{let <var\_0> = [[<clause>]] where <subclause\_1> is <var\_1> ... and <subclause\_n> is <var\_n>}} (a) binds a textual

clause (i.e., a sentence fragment) to a logical variable (for example, the clause “addressed at a group ...just an individual” is named `C1`), and (b) associates sub-clauses within the clause (e.g., “messages”) with external variables, e.g. `M`. Intuitively, `<clause>` represent a constraint between the variables `<var_i>`.

- **Logical constraint blocks** (delimited by `````) specify logical relations (e.g., `t = C6 and C7`) among variables, using `pyz3` notation [z3p; de Moura & Bjørner (2008)].
- **Convenience constructs** such as `forall` and `forsome` compactly handle textual lists, internally expanded into disjunctions or conjunctions over let-bindings.

Such Logitext documents are “executed” using a `check()` function as in constraint solving. Given a partial assignment `p` of variables in a document `d`, `check(d, p, cover)` searches for a satisfying assignment that respects both the logical and textual constraints:

```
check(d:LogitextDocument, p:Dict[str, bool|str], cover:Option[bool])
-> Dict[str, bool|str] | unsat | timeout
```

If a solution exists, `check()` returns a full assignment as a mapping from variable names to values. Otherwise it reports unsatisfiability or timeout. With the optional flag `cover`, `check()` enumerates multiple satisfying assignments. This mechanism generalizes the familiar `complete()` execution of text prompts to a richer constraint-satisfaction setting.

The expressiveness of Logitext unifies diverse language understanding tasks under a single interface. The three tasks of Figure 1(b)–(d) are expressed uniformly as constraint checking: (i) **classification** (`lt.check(d, {'M': M})['d']`), (ii) **partially constrained instance generation** (`lt.check(d, {'C4': False, 'C5': False})['M']`), and (iii) **coverage generation** (`[g['M'] for g in lt.check(d, {'C4': False, 'C5': False}, cover=True)]`). In contrast to raw prompting, Logitext makes explicit the logical structure of documents, enabling solver-style propagation to cooperate with LLM-based textual reasoning.

### 3.2 NATURAL LANGUAGE TEXT CONSTRAINTS

The constructs in Section 3.1 define how Logitext documents combine textual clauses with logical constraints. To reason with such documents, we require a representation that treats textual clauses as first-class objects alongside logical formulas. We introduce *natural language text constraints (NLTCs)*, which bind clauses to variables, record references to external context, and allow seamless interaction with solvers.

Recall from the previous section that, in addition to a variable declaration section, an unparsed Logitext document `d` consists of alternating code blocks and text blocks (Figure 3). Each code block is a sequence of logical strings `k`, e.g., `C8 = C4 or C5`. Each text block contains a sequence of let-binding text strings `L` of the form:

**let** `v` = `[[c]]` **where** `u1` **is** `p1 ... un` **is** `pn`.

Here `v` is a boolean variable to which `c` is bound, while the `ui` are strings (typically substrings of `c`) associated with variables `pi` defined elsewhere.

To process a document `d`, we parse it into an abstract representation `D` in three steps:

1. **Variable collection.** Identify boolean variables  $vs_D = v_1, \dots, v_n$  and string variables  $us_D = u_1, \dots, u_n$ . These variables include those declared explicitly and those introduced in let bindings as above.
2. **Logical constraint parsing.** Convert each logical string `k` from a logical text string into a solver-ready formula `φ` using Z3’s parser.
3. **Textual constraint parsing.** Translate each let-binding `L` into an NLTC  $\nu = (v, c, \{u_1 : p_1, \dots, u_n : p_n\}, d)$ : each NLTC binds `v` to the clause `c`, records its dependencies, and points to the full document `d` so `c` can be interpreted in context.

After parsing, the abstract document is  $D = (vs_D, us_D, \phi_D, \nu_D)$ , where  $\phi_D = \phi_1, \dots, \phi_m$  are logical constraints and  $\nu_D = \nu_1, \dots, \nu_n$  are NLTCs. Reasoning proceeds with respect to a partial assignment  $\pi_D : us_D \cup vs_D \rightarrow \text{bool|str}$ , which specifies known variable values and lets the solver–LLM loop infer the rest, as discussed in the next section.

## 3.3 SOLVING NATURAL LANGUAGE TEXT CONSTRAINTS

| Algorithm 1 $\text{check}(D = (\nu, \phi, vs, us), \pi_D)$                                                  | NLSolver( $u, \nu, \pi$ )                                                                                   |
|-------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------|
| 1: <b>while</b> true <b>do</b>                                                                              | 1: $u^* \leftarrow \text{LLMPropose}(\nu, \pi, \{\}, \text{None}) \quad \triangleright \text{Propose}$      |
| 2: $\pi_Z \leftarrow Z3(\phi, vs, \pi_D) \triangleright \text{Propose bool. assignment}$                    | 2: <b>for</b> $t = 1$ to $T$ <b>do</b>                                                                      |
| 3: <b>return</b> UNSAT <b>if not</b> $\pi_Z$                                                                | 3: <b>for</b> $\nu_k \in \nu$ <b>do with</b> $\text{sat} \leftarrow \text{true}; \bar{\pi} \leftarrow \{\}$ |
| 4: <b>for unbound</b> $u \in us$ <b>do with</b> $\text{sat} \leftarrow \text{true}; \pi_s \leftarrow \{\}$  | 4: $\triangleright \text{Verify proposal; record unsatisfied clause}$                                       |
| 5: $u^* \leftarrow \text{NLSolver}(u, \nu[u], \pi_D \cup \pi_s \cup \pi_Z)$                                 | 5: <b>if</b> $\text{LLMVerify}(\nu_k, \pi \cup \{u = u^*\}) \neq \pi[\nu_k]$ <b>then</b>                    |
| 6: <b>if</b> $!u^*$ <b>then</b> $Z3.\text{block}(\pi_Z)$ ; $\text{sat} \leftarrow \emptyset$ ; <b>break</b> | 6: $\text{sat} \leftarrow \text{false}; \bar{\pi} \leftarrow \bar{\pi} \cup \{\nu_k\}$                      |
| 7: $\pi_s \leftarrow \pi_s \cup \{u = u^*\}$                                                                | 7: <b>If</b> $\text{sat}$ <b>then return</b> $u^*$                                                          |
| 8: <b>if</b> $!\text{sat}$ <b>then continue</b>                                                             | 8: $u^* \leftarrow \text{LLMPropose}(\nu, \pi, \bar{\pi}, u^*) \quad \triangleright \text{Refine}$          |
| 9: <b>return</b> $\pi_s \cup \pi_D \cup \pi_Z$                                                              | 9: <b>return</b> None                                                                                       |

(a) Outer logical solver loop

(b) Inner text solver loop

Figure 4: Core Logitext constraint solving algorithms

Figure 4 shows how Logitext solves hybrid systems of natural language text constraints (NLTCs,  $\nu$ ) and logical constraints ( $\phi$ ), given shared Boolean variables  $vs$ , text-string variables  $us$ , and a partial assignment  $\pi_D$  of Booleans and strings. The goal is to extend  $\pi_D$  into a complete satisfying assignment  $\pi'_D$ , using an LLM-based solver as an extension to the core logical (aka SMT) solver.

The overall strategy is simple. A logical solver (we use Z3 de Moura & Bjørner (2008)) produces candidate Boolean assignments that satisfy Boolean constraints (Fig 4a), and our LLM-based solver (called NLSolver) then attempts to produce assignments to string variables that satisfy text constraints while maintaining the Boolean assignments (Fig 4a).

We begin with the outer logical solver loop of Fig 4a. We use Z3 to generate candidate assignments  $\pi_Z$  of variables  $vs$  consistent with  $\phi$  and  $\pi_D$  (Line 2). We now loop sequentially over unbound string variables  $u$  (Lines 4-9), trying to find satisfying assignments for each, compatible with all assignments so far. Each  $u$  is passed to text-constraint solver NLSolver, which attempts to generate a text string value  $u^*$  for each unbound string variable  $u$  (Line 5), given the constraints  $\nu[u] \subseteq \nu$  that read or write  $u$ , and the assignments so far ( $\pi_s \cup \pi_D \cup \pi_Z$ ). If NLSolver fails to find a  $u^*$ , we block Z3 from regenerating candidate assignment  $\pi_Z$ , break out of the loop over  $us$  (Line 6) and continue generating more candidate Boolean assignments (Lines 8, 2). If all string variables  $u$  are assigned, we declare success and return all assignments accumulated (Lines 7, 9). If no candidates remain, we declare unsatisfiability (Line 3).

NLSolver (Fig. 4b) is given a variable  $u$ , a set  $\nu$  of NLTCs that read  $u$ , and a partial assignment  $\pi$ . Its job is to produce a textual string  $u^*$  for  $u$  that satisfies constraints  $\nu$  given partial assignment  $\pi$ . It does so through a propose-verify-refine loop. It starts by prompting an LLM, via the LLMPropose() call (see Appendix 3), to propose a candidate  $u^*$  that satisfies  $\nu$  and  $\pi$  (Line 1). It then uses  $T$  rounds (Line 2) to refine this solution to one that satisfies  $\nu$ . In each round, for every NLTC  $\nu_k \in \nu$ , it calls into an LLM via LLMVerify() (Appendix ??) to infer the truth value for the variable bound by  $\nu_k$ , given  $u = u^*$  and existing assignment  $\pi$ , and compares this truth value to that required by the partial assignment  $\pi$  (Line 5). If all truth values are compatible, it returns  $u^*$  as a satisfying assignment (Line 7). Otherwise, it uses LLMPropose() to refine  $u^*$  (Line 8). The refinement is guided by an additional “needs-to-change” set  $\bar{\pi}$ , which lists the constraints that  $u^*$  currently violates. After  $T$  rounds of not finding  $u^*$ , we declare failure (Line 9).

The above describes the essence of how Logitext solves NLTCs. In practice, we include two further techniques that have modest impact. First, note that LLMPropose() may produce a piece of text that does not match the current partial assignment  $\pi$ , and is therefore rejected by LLMVerify(). However, LLMPropose() itself may be called many thousands of times for various candidate assignments in the check() algorithm, Line 5. Given that calls to LLMPropose() are relatively expensive since it calls out to LLMs, we cache results from these calls and consider the for use on future calls to NLSolver. Second, when we propose a refinement of textual value  $u^*$ , it helps not only to have the “needs-to-change” list mentioned above, but also a history of the previous refinements proposed on  $u^*$  and their outcome from LLMVerify. These two techniques are described further in the appendix, and the (modest but noticeable) impact of caching is analyzed in the evaluation section (Figure 6).

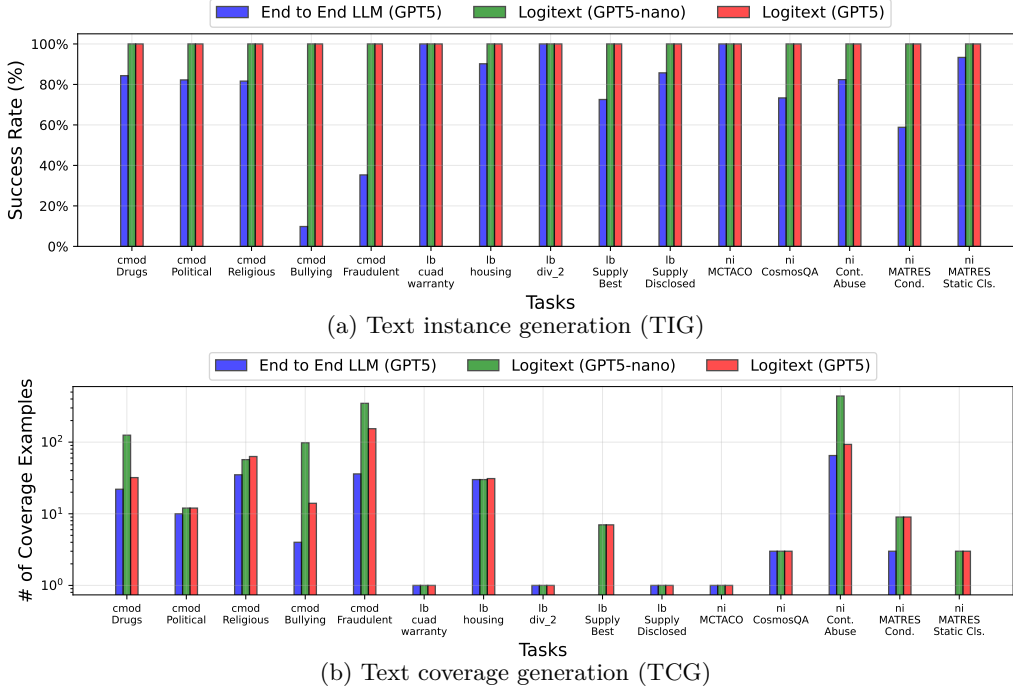


Figure 5: End-to-End performance comparison per task.

## EVALUATION

### BENCHMARKS AND SETUP

**Content Moderation (CMOD).** A new benchmark of five multi-page moderation policies (2–5 pages, 6–22 annotated clauses) covering drugs (22 clauses), politics (8), religion (6), bullying (21), and fraud (12) (see Appendix A.9 for an example). These tasks are designed to reflect realistic compliance settings where policy documents constrain user-generated content.

**Legal Benchmark (LegalBench, LB).** We select five tasks from LegalBench [Guha et al. (2023)], an extensive benchmark for reasoning over statutory and regulatory text. The tasks cover domains such as diversity jurisdiction, housing and warranty law, and supply-chain transparency. Each task is 20–50 lines with 2–4 annotated clauses. This benchmark captures challenges in legal text where precise logical structure interacts with natural language.

**Natural Instructions (NI).** We select five tasks from SuperNatural Instructions [Wang et al. (2022)], focusing on problems with implicit logical constraints such as detecting grammatical inconsistencies, reasoning about hypothetical actions, and identifying abusive content. Each task is 3–10 lines with 2–5 annotated clauses. This benchmark tests generalization to diverse instruction-following tasks beyond policy or law.

Together these benchmarks yield 15 tasks with 10+ instances each, spanning policy, legal, and open-domain instructions. All tasks require mapping text inputs to structured outputs (classifications or constrained generations). We evaluate Logitext on three settings aligned with Fig. 1: (a) **Task execution (TE)** — measuring classification accuracy on task instances, (b) **Text instance generation (TIG)** — testing the ability to generate valid inputs under partial constraints, and (c) **Text coverage generation (TCG)** — enumerating as many valid inputs as possible.

### RESULTS

**Text instance generation (TIG).** Figure 5a compares Logitext with direct prompting. With both GPT5 and GPT5-nano as base models, Logitext generates valid assignments reliably via `check()`. The results indicate that (i) Logitext attains near-saturation perfor-

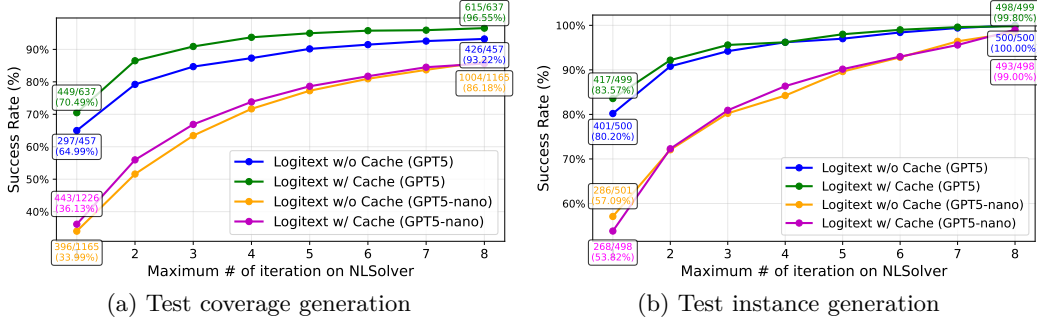
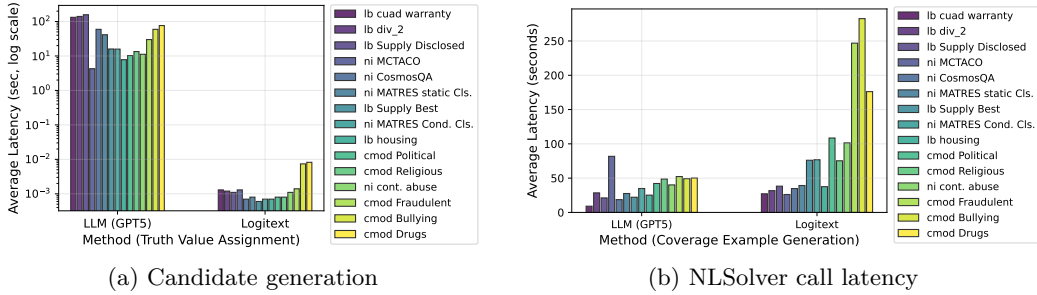


Figure 6: NLSolver success rate vs num. iterations.

mance even with the weaker GPT5-nano, while direct prompting to GPT5 shows noticeable degradation; and (ii) degradation is most evident on complex CMOD policies, although performance on Drugs is relatively stronger than other cases.



(a) Candidate generation

(b) NLSolver call latency

Figure 7: Component-wise latency on TCG (tasks sorted by the #clauses).

**Text coverage generation (TCG).** Figure 5b (log scale) reports coverage under a fixed time budget of 3000s. Baseline GPT is allowed up to 5 iterations for candidate generation and 5 additional iterations for finding satisfying assignments. Logitext achieves broader coverage, particularly with GPT5-nano. Figure 7 provides an explanation: candidate generation is significantly faster with Logitext since it uses a solver rather than repeated LLM calls (Fig. 7a). The advantage is reduced in the solving phase (Fig. 7b), where NLSolver calls dominate, but this bottleneck is smaller for faster base models such as GPT5-nano. We also report the LLM call statistics of NLSolver in Figure 11 of Appendix.

**NLSolver iterative refinement.** Figure 6 shows that NLSolver improves success rates as the number of iterations increases. Caching can be beneficial in some settings (e.g., coverage generation with GPT5), though its overall effect across tasks is limited.

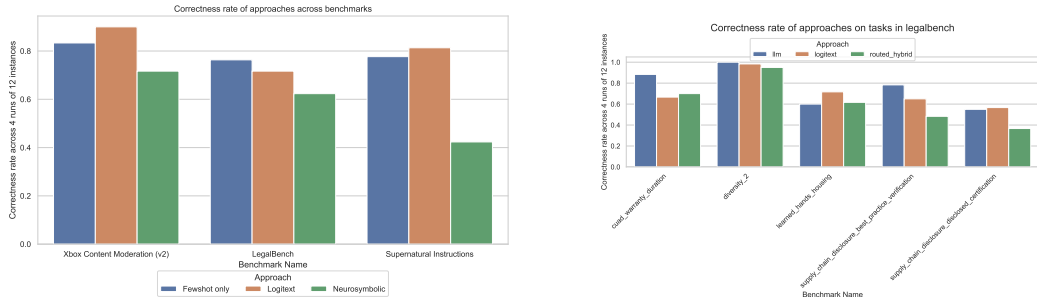


Figure 8: Aggregated Task Execution (TE) Result (left), Legalbench in detail (right).

**Task Execution (TE).** Figure 8(left) presents accuracy on TE tasks using GPT-4o. In addition to few-shot prompting, we include a neurosymbolic prompt that generates and executes code. Logitext performs better on CMOD and NI benchmarks but underperforms on LegalBench. Figure 8(right) highlights two main sources of error: (i) in some cases, clause-level outputs  $[[C_i]]$  were incorrectly predicted by the LLM, but the LLM-only approach produced correct answers using holistic reasoning, revealing a robustness limitation for Logitext; (ii) in other cases, lists such as “x, y, and z” were intended as examples rather than conjuncts, but were annotated as the latter. These issues point to the need for clause-level error correction and more careful handling of list annotations in future work.



Overall, the experiments show that Logitext improves performance on both constrained generation and classification tasks across multiple benchmarks, while highlighting remaining challenges in clause-level robustness and annotation handling.

## 5 RELATED WORK

**Prompt-based reasoning.** Prompting strategies such as Chain-of-Thought (CoT) [Wei et al. (2022)] and Tree-of-Thought (ToT) [Yao et al. (2023)] elicit multi-step reasoning by decomposing queries into textual steps. Chain-of-Logic [Servantez et al. (2024)] separates logical reasoning from answer prediction, aiming to improve consistency in step-wise deduction. While these approaches enhance local coherence or allow limited backtracking, they lack mechanisms to bridge deeper compositional and combinatorial reasoning gaps and cannot ensure global logical consistency across clauses. Once an error propagates, there is no principled way to validate against constraints. Our framework differs by explicitly defining these reasoning gaps and incorporating symbolic validation into the reasoning process itself.

**Systematic generalization and constraint solving.** Our challenges relate closely to systematic generalization tasks [Lake & Baroni (2018), Keyzers et al. (2020), Kim & Linzen (2020)], which demonstrate that sequence models fail when compositional rules must be recombined in novel ways. Similar issues arise in program synthesis and constraint satisfaction tasks, where LLMs can propose candidate programs or assignments (e.g., Codex for SAT/SMT) but collapse under combinatorial growth in the search space. These methods provide no principled mechanism to enforce or recover from violated constraints. We formalize these compositional and combinatorial reasoning gaps as structural limitations of LLM inference and show how solver integration can systematically mitigate them in natural language contexts.

**Reasoning models.** Reasoning models such as OpenAI o3/o4-mini and DeepSeek-R1 [Guo et al. (2025)] have shown improved performance on benchmarks for logical reasoning and robust instruction following. RL allows limited correction through feedback [Kalyanpur et al. (2024)] or exploration [Xie et al. (2025)]. However, they depend heavily on reward shaping or sample filtering, lack a formal representational layer for partial logical structures. As a result, they cannot enforce symbolic constraints or recover from violated ones during inference. Our framework complements these advances by introducing a neurosymbolic language that supports constraint-aware reasoning within an SMT framework.

**Neuro-symbolic reasoning.** Recent systems such as LINQ [Olausson et al. (2023)], CLOVER [Ryu et al. (2025)], and ZebraLogic [Lin et al. (2025)] connect LLMs with symbolic solvers by translating natural language into executable logic programs. These approaches achieve strong guarantees when tasks are fully formalizable, but their reliance on complete logical structure restricts applicability to natural documents like policies or legal texts, where only fragments of logic are explicit. ZebraLogic also provided an initial study of the *combinatorial gap*, but its scope was limited to well-defined mathematical domains. In contrast, our work addresses this challenge in natural language contexts that inherently contain uncertainty and partial structure. We introduce natural language text constraints (NLTCs), enabling partial formalization and iterative solver-guided reasoning that better reflects the complexity of real-world documents.

## 6 CONCLUSION

In this work we introduced **Logitext**, a neurosymbolic framework that treats LLM reasoning as an SMT theory through natural language text constraints. We first motivated the need for such a framework by showing that even frontier LLMs continue to exhibit two reasoning gaps: compositional gaps that narrow with scale but persist, and combinatorial gaps that remain severe. By interleaving solver propagation with LLM decoding, Logitext provides a principled way to reduce these gaps. More broadly, our results suggest a path toward positioning LLMs as solver-compatible theories, opening opportunities for scalable, reliable, and trustworthy natural language reasoning.

## REFERENCES

- Z3py api documentation. <https://z3prover.github.io/api/html/namespacez3py.html>. Accessed: 2025-11-19.
- Clark Barrett, Aaron Stump, and Cesare Tinelli. The smt-lib standard: Version 2.0. Technical report, Department of Computer Science, The University of Iowa, 2010. URL <https://theory.stanford.edu/~barrett/pubs/BST10.pdf>.
- Martin Davis, George Logemann, and Donald Loveland. A machine program for theorem-proving. *Communications of the ACM*, 5(7):394–397, 1962. doi: 10.1145/368273.368557.
- Leonardo de Moura and Nikolaj Bjørner. Z3: An efficient smt solver. In *Tools and Algorithms for the Construction and Analysis of Systems (TACAS)*, volume 4963 of *Lecture Notes in Computer Science*, pp. 337–340. Springer, 2008. doi: 10.1007/978-3-540-78800-3\_24.
- Neel Guha, Julian Nyarko, Daniel Ho, Christopher Ré, Adam Chilton, Alex Chohlas-Wood, Austin Peters, Brandon Waldon, Daniel Rockmore, Diego Zambrano, et al. Legalbench: A collaboratively built benchmark for measuring legal reasoning in large language models. *Advances in neural information processing systems*, 36:44123–44279, 2023.
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- Aditya Kalyanpur, Kailash Karthik Saravanakumar, Victor Barres, Jennifer Chu-Carroll, David Melville, and David Ferrucci. Llm-arc: Enhancing llms with an automated reasoning critic. *arXiv preprint arXiv:2406.17663*, 2024.
- Daniel Keysers, Nathanael Schärli, Nathan Scales, Hylke Buisman, Daniel Furrer, Sergii Kashubin, Nikola Momchev, Danila Sinopalnikov, Lukasz Stafiniak, Tibor Tihon, Dmitry Tsarkov, Xiao Wang, Marc van Zee, and Olivier Bousquet. Measuring compositional generalization: A comprehensive method on realistic data, 2020. URL <https://arxiv.org/abs/1912.09713>.
- Najoung Kim and Tal Linzen. Cogs: A compositional generalization challenge based on semantic interpretation, 2020. URL <https://arxiv.org/abs/2010.05465>.
- Brenden M. Lake and Marco Baroni. Generalization without systematicity: On the compositional skills of sequence-to-sequence recurrent networks, 2018. URL <https://arxiv.org/abs/1711.00350>.
- Bill Yuchen Lin, Ronan Le Bras, Kyle Richardson, Ashish Sabharwal, Radha Poovendran, Peter Clark, and Yejin Choi. ZebraLogic: On the scaling limits of llms for logical reasoning, 2025. URL <https://arxiv.org/abs/2502.01100>.
- João P. Marques-Silva and Karem A. Sakallah. Grasp: A search algorithm for propositional satisfiability. *IEEE Transactions on Computers*, 48(5):506–521, 1999. doi: 10.1109/12.769433.
- Theo Olausson, Alex Gu, Ben Lipkin, Cedegao Zhang, Armando Solar-Lezama, Joshua Tenenbaum, and Roger Levy. Linc: A neurosymbolic approach for logical reasoning by combining language models with first-order logic provers. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pp. 5153–5176. Association for Computational Linguistics, 2023. doi: 10.18653/v1/2023.emnlp-main.313. URL <http://dx.doi.org/10.18653/v1/2023.emnlp-main.313>.
- Philipp Rümmer and Thomas Wahl. An SMT-LIB theory of binary floating-point arithmetic. In *Proc. 8th Intl. Workshop on Satisfiability Modulo Theories (SMT’10)*, 2010.
- Hyun Ryu, Gyeongman Kim, Hyemin S. Lee, and Eunho Yang. Divide and translate: Compositional first-order logic translation and verification for complex logical reasoning, 2025. URL <https://arxiv.org/abs/2410.08047>.

- Yusuke Sakai, Hidetaka Kamigaito, and Taro Watanabe. Revisiting compositional generalization capability of large language models considering instruction following ability. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 31219–31238, Vienna, Austria, 2025. Association for Computational Linguistics. doi: 10.18653/v1/2025.acl-long.1508. URL <https://aclanthology.org/2025.acl-long.1508/>.
- Sergio Servantez, Joe Barrow, Kristian Hammond, and Rajiv Jain. Chain of logic: Rule-based reasoning with large language models, 2024. URL <https://arxiv.org/abs/2402.10400>.
- Yizhong Wang, Swaroop Mishra, Pegah Alipoormolabashi, Yeganeh Kordi, Amirreza Mirzaei, Anjana Arunkumar, Arjun Ashok, Arut Selvan Dhanasekaran, Atharva Naik, David Stap, et al. Super-naturalinstructions: Generalization via declarative instructions on 1600+ nlp tasks. *arXiv preprint arXiv:2204.07705*, 2022.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Ed H. Chi, Quoc Le, and Denny Zhou. Chain of thought prompting elicits reasoning in large language models. *CoRR*, abs/2201.11903, 2022. URL <https://arxiv.org/abs/2201.11903>.
- Jiaxin Wen, Jian Guan, Hongning Wang, Wei Wu, and Minlie Huang. Codeplan: Unlocking reasoning potential in large language models by scaling code-form planning. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2025. URL [https://proceedings.iclr.cc/paper\\_files/paper/2025/hash/c362b360765ed90ae4bd9c6764837e0e-Abstract-Conference.html](https://proceedings.iclr.cc/paper_files/paper/2025/hash/c362b360765ed90ae4bd9c6764837e0e-Abstract-Conference.html).
- Tian Xie, Zitian Gao, Qingnan Ren, Haoming Luo, Yuqian Hong, Bryan Dai, Joey Zhou, Kai Qiu, Zhirong Wu, and Chong Luo. Logic-rl: Unleashing llm reasoning with rule-based reinforcement learning. *arXiv preprint arXiv:2502.14768*, 2025.
- Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Thomas L. Griffiths, Yuan Cao, and Karthik Narasimhan. Tree of thoughts: deliberate problem solving with large language models. In *Proceedings of the 37th International Conference on Neural Information Processing Systems, NIPS ’23*, Red Hook, NY, USA, 2023. Curran Associates Inc.
- Xi Ye, Qiaochu Chen, Isil Dillig, and Greg Durrett. Satlm: Satisfiability-aided language models using declarative prompting. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2023. doi: 10.48550/arXiv.2305.09656. URL <https://arxiv.org/abs/2305.09656>.
- Yunhui Zheng, Vijay Ganesh, Sanu Subramanian, Omer Tripp, Murphy Berzish, Julian Dolby, and Xiangyu Zhang. Z3str2: an efficient solver for strings, regular expressions, and length constraints. *Formal Methods in System Design*, 50(2-3):249–288, 2017.

## A APPENDIX

Additional information follows.

### A.1 DEFINITIONS OF GAPS

Each policy  $p$  is annotated with clauses  $C_i$  and associated with a formula  $\phi$  as in Fig 1 and Eqn1, and comes with 10-20 test messages  $M_j$  each with ground truth  $H_j$ .

**Definition 1 (Compositional gap  $\Delta_{mp}$  of LLM  $m$  on prompt  $p$ ).** For each  $M_j$ , prompt  $m$  with  $p$  for (i) the meanings  $b_{ij}$  of clauses  $C_i$ , and (ii) whole-prompt result  $h_j$ . This results in overall accuracy  $a = \text{mean}_j \delta(h_j, H_j)$ , where  $\delta(x, x') = 1$  if  $x = x'$ , else 0. Now, use a logical solver to evaluate  $h_j^* = \phi(b_{ij})$ , giving corresponding accuracy  $a^*$ . Then compositional gap  $\Delta_{mp} = a - a^*$

**Definition 2 (Combinatorial gap  $\Delta'_{mp}$  of model LLM  $m$  on prompt  $p$ ).** Prompt  $m$  with  $p$  to produce (i) all policy inputs  $M_j$ , and (ii) complete assignments  $b_{ij}$  for clauses  $C_i$  such that the policy is true. Use a logical solver to filter out  $b_{ij}$  that do not satisfy  $\phi$ , and also to generate independently its satisfying assignments  $b_{ij}^*$ . Let  $n_j$  (resp.  $n_j^*$ ) be number of such assignments. The combinatorial gap is the relative discrepancy between these numbers:  $\Delta'_{mp} = \text{mean}_j(n_j^* - n_j)/n_j^*$ ,

## A.2 SYNTAX FOR LOGITEXT DOCUMENTS

```

doc  $d \leftarrow [b \mid c]^+$ 
text block  $b \leftarrow [s \mid t]^+$ 
code block  $c \leftarrow \text{``}[(v_1, \dots, v_n)]\langle \text{code} \rangle\text{``}$ 
text  $s \leftarrow \langle \text{strings without } \{\{ \text{ or } \}\} \rangle$ 
term  $t \leftarrow l \mid q$ 
let  $l \leftarrow \{\{ \text{let } v = [[b]] \text{ where } r_1 \text{ is } v_1 \text{ and } \dots \text{ and } r_n \text{ is } v_n \} \}$ 
quantifier  $q \leftarrow \{\{ \text{forall } [s \mid [[b]]]^+ \} \} \mid \{\{ \text{forsome } [s \mid [[b]]]^+ \} \}$ 
typed variable  $v \leftarrow \langle \text{variable name} \rangle[: \text{str}]$ 
quoted str.  $r \leftarrow " \dots "$ 

```

## A.3 EXPANDED DATASET

| Benchmark                                          | Task | Original Submission |       | Resubmission |       | Evaluated |
|----------------------------------------------------|------|---------------------|-------|--------------|-------|-----------|
|                                                    |      | #Inst               | #Runs | #Inst        | #Runs | #Inst     |
| cmod                                               |      |                     |       |              |       |           |
| Bullying_2.0                                       |      | 12                  | 5     | 100          | 5     | 0         |
| Drugs_&_Alcohol_2.0                                |      | 12                  | 5     | 100          | 5     | 0         |
| Fraudulent_v2.0v                                   |      | 12                  | 5     | 100          | 5     | 0         |
| Political_v2.0v                                    |      | 12                  | 5     | 100          | 5     | 0         |
| Religious_v2.0v                                    |      | 12                  | 5     | 100          | 5     | 0         |
| legalbench                                         |      |                     |       |              |       |           |
| cuad_warranty_duration                             |      | 12                  | 5     | 100          | 5     | 100       |
| diversity_2                                        |      | 12                  | 5     | 100          | 5     | 100       |
| learned_hands_housing                              |      | 12                  | 5     | 100          | 5     | 100       |
| supply_chain_disclosure_best_practice_verification |      | 12                  | 5     | 100          | 5     | 100       |
| supply_chain_disclosure_disclosed_certification    |      | 12                  | 5     | 100          | 5     | 100       |
| natural_instructions                               |      |                     |       |              |       |           |
| task021_mctaco_grammatical_logical                 |      | 12                  | 5     | 100          | 5     | 50        |
| task022_cosmosqa_passage_inappropriate_binary      |      | 12                  | 5     | 100          | 5     | 50        |
| task108_contextualabusedetection_classification    |      | 12                  | 5     | 100          | 5     | 50        |
| task457_matres_conditional_classification          |      | 12                  | 5     | 100          | 5     | 50        |
| task459_matres_static_classification               |      | 12                  | 5     | 100          | 5     | 50        |

Table 1: Comparison of dataset instance counts and runs in the original submission vs. resubmission.

We have expanded the dataset evaluated to 100 samples per task from 12 samples per task as shown in Table 1. We are in the process of running evaluations on the data. So far, we have completed full re-evaluation on 100 samples on 5 tasks (from Legalbench), partial re-evaluation on 50 samples on (from Supernatural Language Instructions (SNLI)), and have not yet started re-evaluation on our most favorable dataset CMOD. For Legalbench and SNLI, each task had sufficient samples that we were able to simply incorporate more samples from the existing dataset. For CMOD, we had to generate samples analogous to production moderation data, a task that requires some care.

We focused on re-running the Task Execution (TE) experiments (Figure 8), both because these are the least favorable to us, and because the combinatorial gap experiments take much longer to run. Of course, we will complete running all experiments on all datasets over the next few days.

As Tables 2 and 3 show, the expanded results don’t change the highest level message on the Task Execution task qualitatively: Logitext does provide a noticeable boost on many tasks, and prevails in 7 of 10 tasks, but the baseline model does do better in some cases. Perhaps interestingly, Logitext now does relatively better on Legalbench, prevailing in 4/5 tasks, and slightly worse on SNLI (3/5). Once again, when Logitext fails, the main culprit seems to be clause-level evaluation errors, which we have discussed in more detail in Appendix A.4, and mentioned in the original submission.

| Task Name                                          | Fewshot     | Neurosymbolic | Logitext    |
|----------------------------------------------------|-------------|---------------|-------------|
| cquad_warranty_duration                            | 0.50        | 0.19          | <b>0.61</b> |
| diversity_2                                        | 0.76        | 0.35          | <b>0.83</b> |
| learned_hands_housing                              | 0.50        | 0.34          | <b>0.60</b> |
| supply_chain_disclosure_best_practice_verification | 0.58        | 0.02          | <b>0.59</b> |
| supply_chain_disclosure_disclosed_certification    | <b>0.72</b> | 0.57          | 0.33        |

Table 2: Correctness on the TG task for Legalbench (100 samples/task)

| Task Name                                       | Fewshot     | Neurosymbolic | Logitext    |
|-------------------------------------------------|-------------|---------------|-------------|
| task021_mctaco_grammatical_logical              | 0.41        | 0.44          | <b>0.50</b> |
| task022_cosmosqa_passage_inappropriate_binary   | 0.78        | 0.69          | <b>0.80</b> |
| task108_contextualabusedetection_classification | <b>0.75</b> | 0.38          | 0.63        |
| task457_matres_conditional_classification       | <b>0.87</b> | 0.49          | 0.59        |
| task459_matres_static_classification            | 0.57        | 0.56          | <b>0.69</b> |

Table 3: Correctness on the TG task for SNLI (50 samples/task)

#### A.4 CLAUSE-LEVEL ERROR ANALYSIS

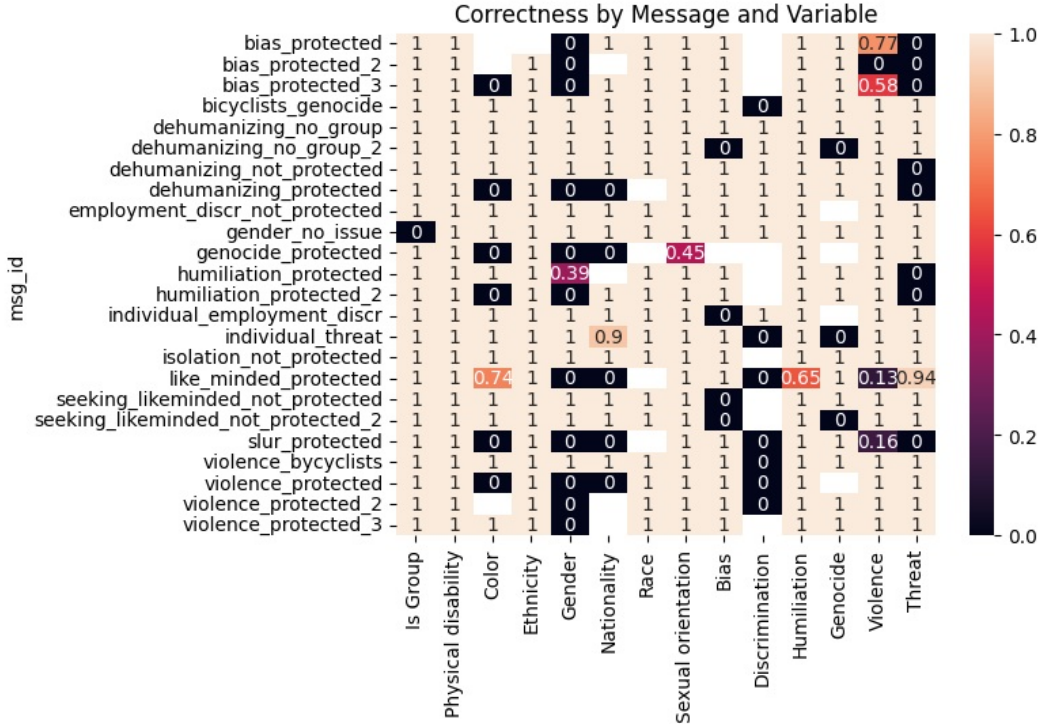


Figure 9: Clause-level accuracy for Disruptive Behavior policy from Figure 1a

Figure 9 is a quick analysis of clause-level accuracy for the “Disruptive Behavior” prompt of Figure 1a. The x-axis of the figure lists the various clauses in the prompt<sup>1</sup>. The y-axis represents individual messages input to the prompt. Each message was run 100 times through the prompt using `gpt-4.1-nano` as the base model. Each entry in the table is the fraction of times the answer for a particular clause was correct for that message. Some entries are missing because some runs did not complete.

Several points are worth noting, all admittedly only in the context of the current prompt:

1. Sub-clause level inference can be quite stable across runs. They are usually either always correct or always wrong, only occasionally are they not 0 or 1. Thus the worst case of several sub-clauses at a time unpredictably producing errors due to stochastic variation is not inevitable.
2. Inference results are predominantly correct, i.e., sub-clause level accuracy may be much higher than “holistic” document-level accuracy.
3. Certain clauses (i.e., columns, e.g. Gender) are consistently interpreted incorrectly across inputs. In a production setting, we would consider re-wording these, hopefully yielding consistently *correct* clauses.
4. Most inputs (i.e., rows) always have at least one sub-clause evaluated incorrectly. This may seem fatal, until we recall the logic of the prompt is essentially  $d = \text{isGroup} \wedge (\text{PhysicalDisability} \vee \text{Color} \vee \text{Ethnicity} \vee \text{Gender} \vee \text{Nationality} \vee \text{Race} \vee \text{SexualOrientation}) \wedge (\text{Bias} \vee \text{Discrimination} \vee \text{Humiliation} \vee \text{Genocide} \vee \text{Violence})$ . If a clause Gender, which is supposed to be False by default, wrongly evaluates to True, it will not cause an end-to-end error given that Gender is part of a larger disjunction (“or”) operation. Thus the precise value of the error, the operation it is part of, and the values of other operands in the operation all contribute to whether a higher-level error is generated. Clausal error does not necessarily imply global error.

While this analysis is by no means comprehensive, it gives some intuition of why the introduction of clause-level errors does not necessarily lead to catastrophic failure at the whole-formula level.

## A.5 SOLVER ALGORITHMS

### A.5.1 CHECK() ALGORITHM DETAILS

The fully detailed version of the Check() algorithm (Figure 10) for solving Logitext constraints is moved here. The version includes details of caching and history, as mentioned in the main body of the paper.

<sup>1</sup>Note in the version of the example used in the body of the paper, we omit the Discrimination and Humiliation categories for brevity but they are included in this analysis, which was performed on the complete version of the policy document.

| Algorithm 2 $\text{check}(D, \pi_D)$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          | NLSolver( $u, \mathcal{N}, \pi$ )                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <pre> 1: <b>In:</b> Doc. <math>D = (vs, us, \phi, \nu)</math>, asst. <math>\pi_D</math> 2: <b>Out:</b> UNSAT or satisfying asst. <math>\pi_D</math> 3: <b>Initialize</b> logical solver <math>Z_\phi</math> with <math>\phi</math> 4: <b>if</b> all <math>u_i \in us</math> are bound in <math>\pi_D</math> <b>then</b> 5:   <b>return</b> LLMVerify(<math>\nu, \phi, \pi_D</math>) 6: <b>while</b> true <b>do</b> 7:   //Propose asst. respecting <math>\phi</math> and <math>\pi_D</math> 8:   <math>\pi_Z \leftarrow Z_\phi(vs, \pi_D)</math> 9:   <b>return</b> UNSAT <b>if</b> <math>\pi_Z = \emptyset</math> 10:  // NLTC solving for unbound text vars 11:  <math>\pi_s</math>, satisfiable <math>\leftarrow \{\}</math>, true 12:  <b>for</b> each unbound <math>u_j \in us</math> <b>do</b> 13:    // Use relevant constraints <math>\mathcal{N}</math> for <math>u_j</math> 14:    <math>\mathcal{N} = \{\nu_i \in \nu \mid \nu_i \text{ reads } u_j\}</math> 15:    <math>u_j^* \leftarrow \text{NLSolver}(u_j, \mathcal{N}, \pi_D \cup \pi_s \cup \pi_Z)</math> 16:    <b>if</b> <math>u_j^*</math> is None <b>then</b> 17:      <math>Z_\phi.\text{block}(\pi_Z)</math> 18:      satisfiable <math>\leftarrow</math> false ; <b>break</b> 19:    <math>\pi_s \leftarrow \pi_s \cup \{u_j = u_j^*\}</math> 20:  <b>if not</b> satisfiable <b>then continue</b> 21:  <b>return</b> <math>\pi_s \cup \pi_D \cup \pi_Z</math> </pre> | <pre> 1: <b>In:</b> Search target <math>u</math>, NLTC set <math>\mathcal{N}</math>, its partial asst. <math>\pi</math>,    Cache <math>C</math>, and <math>T \in \mathbb{N}</math> 2: <b>Out:</b> String value <math>u^*</math> or None 3: <b>if</b> <math>(u, \mathcal{N}, \pi) \in C</math> <b>then</b> <span style="float: right;">▷ Cache Lookup</span> 4:   <b>return</b> <math>C[(u, \mathcal{N}, \pi)]</math> 5: <b>else if</b> <math>\exists C.\text{partial\_match}(u, \mathcal{N}, \pi)</math> <b>then</b> 6:   <math>u^* \leftarrow C.\text{closest\_partial\_match}(u, \mathcal{N}, \pi)</math> 7: <b>else</b> 8:   Sample <math>u^* \sim \text{LLM}(\mathcal{N}, \pi, \emptyset, \emptyset, \emptyset)</math> 9:   History <math>H \leftarrow \{u^*\}</math> 10:  <b>for</b> <math>t = 1</math> to <math>T</math> <b>do</b> 11:    <math>\text{sat} \leftarrow</math> true, <math>\bar{\pi} \leftarrow \emptyset</math>, <math>\tilde{\pi} \leftarrow \emptyset</math> 12:    <b>for</b> <math>\nu_k \in \mathcal{N}</math> <b>do</b> 13:      <math>b_k \leftarrow</math> Truth value for <math>\nu_k</math> from <math>\pi</math> 14:      <math>\tilde{b}_k \leftarrow \text{LLMVerify}(\nu_k, \pi \cup \{u = u^*\})</math> 15:      <b>if</b> <math>b_k \neq \tilde{b}_k</math> <b>then</b> 16:        <math>\text{sat} \leftarrow</math> false 17:        <math>\bar{\pi} \leftarrow \bar{\pi} \cup \{(\nu_k, b_k)\}</math> 18:        <math>\tilde{\pi} \leftarrow \tilde{\pi} \cup \{(\nu_k, \tilde{b}_k)\}</math> 19:      <math>C[(u, \mathcal{N}, (\pi - \bar{\pi}) \cup \tilde{\pi})] \leftarrow u^*</math> 20:    <b>if</b> <math>\text{sat}</math> <b>then return</b> <math>u^*</math> 21:    <math>u^* \leftarrow \text{LLM}(\mathcal{N}, \pi, H, \bar{\pi}, u^*)</math> 22:    History <math>H \leftarrow H \cup \{(u^*, \bar{\pi})\}</math> 23:  <b>return</b> None </pre> |
| (a) Outer SMT/NLSolver loop                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   | (b) NLSolver: A theory for NLTCs                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |

Figure 10: The check() algorithm for solving logitext constraints

### A.5.2 LLMPROPOSE ALGORITHM DETAILS

LLMPropose (Algorithm 3), given a string variable  $u$ , a set of NLTCs, a partial assignment, produces a text string corresponding to  $u$  that satisfies the NLTCs and the partial assignment. LLMPropose is a thin wrapper around an LLM prompt (Appendix A.7.1).

| Algorithm 3 LLMPropose( $u, N_k, \mathcal{P}$ )                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <pre> 1: <b>Input:</b> string variable <math>u</math>, NLTCs <math>N_k</math>, context value assignments <math>\mathcal{P}</math> 2: <b>Output:</b> Generated text string 3: <math>\text{prompt} \leftarrow</math> Format <math>N_k</math> and <math>\mathcal{P}</math> as a text prompt (Appendix A.7.1) that requests generation    of text satisfying <math>N_k</math> given context <math>\mathcal{P}</math> 4: <math>\text{response} \leftarrow \text{LLM.call}(\text{prompt})</math> 5: <math>\text{result} \leftarrow</math> Parse and extract the generated text from <math>\text{response}</math> 6: <b>return</b> <math>\text{result}</math> </pre> |

### A.5.3 LLMVERIFY ALGORITHM DETAILS

Given an NLTC and an assignment of variables to values, LLMVerify (Algorithm 4) evaluates the output (Boolean) variable of the NLTC. it is a thin wrapper around the LLM prompt of Appendix A.7.3.

| Algorithm 4 LLMVerify( $N_k, \mathcal{P}$ )                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <pre> 1: <b>Input:</b> NL Text constraint <math>N_k</math>, context value assignments <math>\mathcal{P}</math> 2: <b>Output:</b> True/False 3: <math>\text{prompt} \leftarrow</math> Format <math>N_k</math> and <math>\mathcal{P}</math> as a text prompt (Appendix A.7.3) that queries whether    <math>N_k</math> is True or False based on the context <math>\mathcal{P}</math> 4: <math>\text{response} \leftarrow \text{LLM.call}(\text{prompt})</math> 5: <math>\text{result} \leftarrow</math> Parse the <math>\text{response}</math> into True or False 6: <b>return</b> <math>\text{result}</math> </pre> |

## A.6 NUMBER OF LLM CALLS FROM NLSOLVER

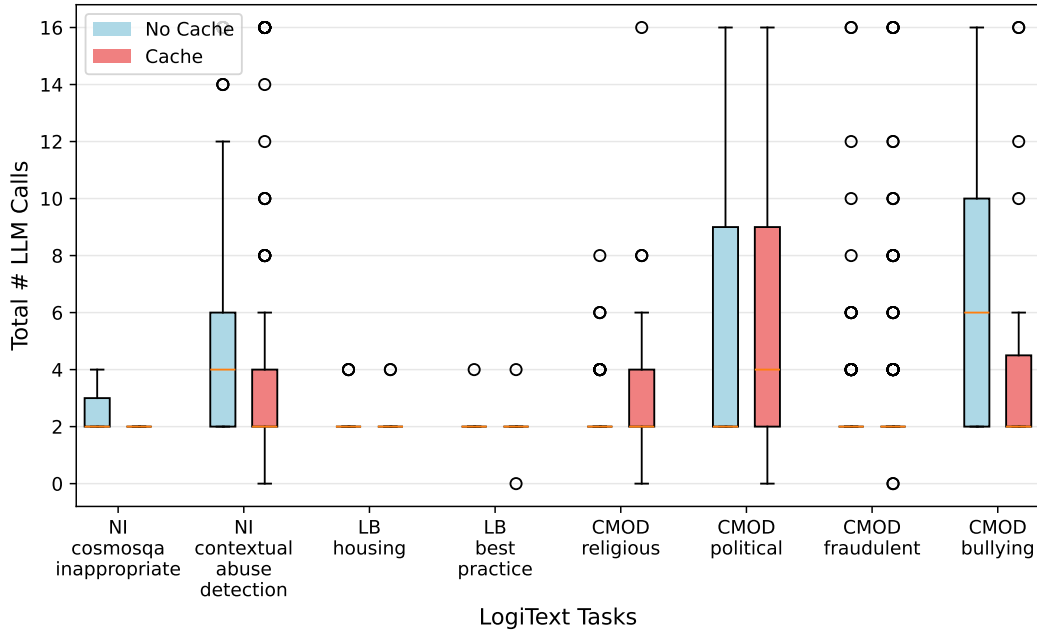


Figure 11: Boxplot of the number of LLM calls per coverage example made by the NLSolver algorithm across all benchmark tasks. The X-axis denotes individual tasks, and the Y-axis reports the distribution of LLM call counts per task.

#### A.7 LLM PROMPTS USED FOR NLSOLVER ALGORITHM

These are the LLM prompts for logitext’s NLSolver algorithm.

##### A.7.1 LLM PROPOSE PROMPT

```
[
  (
    "system",
    "You write concise, clear prose that meets semantic
      constraints. "
    "Return ONLY the text, no preambles."
  ),
  (
    "user",
    ""You must produce a text that satisfies the following
      semantic constraint truth values.

Constraints (with target truth values):
{constraints_block}

If some constraints previously failed, focus on satisfying them now:
{failed_block}

Guidelines:
Write one coherent text (few paragraphs max) that attempts to satisfy
  ALL target truth values.
The text should be natural and readable.
Avoid mentioning constraints explicitly.
Do not include JSON. Output ONLY the text.

Previous attempts for context (do NOT copy; improve instead):
{history_block}
"""
```



```

    ),
]

```

Listing 1: LLM Propose Prompt

## A.7.2 EXAMPLE LLM PROPOSE PROMPT INSTANCE

```

--- System Prompt
You write concise, clear prose that meets semantic constraints. Return
ONLY the text, no preambles.

--- User Prompt
You must produce a text that satisfies the following semantic
constraint truth values.
Constraints (with target truth values):
- (C1) target=True :: Content that is a political reference. In this
  case there is
    sufficient context.

- **Political Reference** - Words or phrases that are
  political
  references. Includes political titles, entities, slogans,
  government buildings, federal properties, etc.

- (C2) target=False :: Content that is a political reference. In this
  case there is
    sufficient context.

- **Political Party or Affiliation** - Name of a political
  party
  or group, including the people that belong to it.

- (C3) target=False :: Content or activity related to politics.
  - **Activity/Discussion** - Activities or discussions
    related to
    politics. Includes political ideologies, debates,
    campaigning,
    causes, events, etc.

- (C4) target=False :: **Potential Political Reference**
  Content that could be a political reference. In this case it is
  uncertain due to its commonality, multiple meanings, current
  usage,
  and/or insufficient context.
  - **General Reference** - Words or short phrases that could
    potentially be used in a political manner, although there
    is not
    enough context to make that determination.

- (C5) target=False :: Names of Political figures or people
  For example, "Joe Biden", "Donald Trump".

- (C6) target=False :: References to royal families, their titles and
  duties.

```

```

For example, "King Charles", "Prince William".

If some constraints previously failed, focus on satisfying them now:
- (C2) target=False, predicted=True :: Content that is a political
  reference. In this case there is sufficient context.

  - **Political Party or Affiliation** - Name of a political
    party
    or group, including the people that belong to it. [why
    failed last time: No political party names or
    affiliations appear in the text.]
- (C3) target=False, predicted=True :: Content or activity related to
  politics.

  - **Activity/Discussion** - Activities or discussions
    related to politics. Includes political ideologies,
    debates, campaigning, causes, events, etc. [why failed
    last time: The text is about gaming, not politics.]
- (C5) target=False, predicted=True :: Names of Political figures or
  people

  For example, "Joe Biden", "Donald Trump". [why failed last
  time: No names of political figures appear in the text.]
- (C6) target=False, predicted=True :: References to royal families,
  their titles and duties.

  For example, "King Charles", "Prince William". [why failed
  last time: The text does not mention royalty.]

Guidelines:
- Write one coherent text (few paragraphs max) that attempts to
  satisfy ALL target truth values.
- The text should be natural and readable.
- Avoid mentioning constraints explicitly.
- Do not include JSON. Output ONLY the text.

Previous attempts for context (do NOT copy; improve instead):
That last boss was -insane pulled it off just in time. The map\'s
skyline had a Capitol-like building, which gave a real-world vibe
without leaving the game.

---

Just got through a brutal boss. The \'maps skyline features a Capitol-
like dome in the distance, giving the game a real-world vibe
without leaving the fantasy setting. Came down to a clutch -
finishtotally worth the grind.

```

Listing 2: LLM Generation Prompt Example for Political CMOD Task

## A.7.3 LLM VERIFY (CONSTRAINT VERIFICATION) PROMPT

```

[
  (
    "system",
    "You are a meticulous verifier. "
    "Given a candidate text and a list of constraints with desired
    truth values, "
    "judge for EACH constraint whether it is semantically TRUE or
    FALSE in the candidate text. "
    "Be strict and literal, not aspirational."
  ),
  (
    "user",

```

```

972         """Candidate text:
973         \"\"\"
974         {candidate}
975         \"\"\"
976
977         Evaluate each constraint independently. For each item, return JSON
978         array entries of the form:
979         {
980             "id": "<constraint id>",
981             "description": "<verbatim description>",
982             "target": true|false,
983             "predicted": true|false,
984             "rationale": "<short explanation>"
985         }
986
987         Constraints:
988         {constraints_block}
989
990         Return ONLY valid JSON array, nothing else."""
991     ),
992 ]

```

Listing 3: LLM Verify Prompt

#### A.7.4 EXAMPLE LLM VERIFY PROMPT INSTANCE

```

996 --- System Prompt
997
998 You are a meticulous verifier. Given a candidate text and a list of
999 constraints with desired truth values, judge for EACH constraint
1000 whether it is semantically TRUE or FALSE in the candidate text. Be
1001 strict and literal, not aspirational.
1002
1003 --- User Prompt
1004
1005 Candidate text:
1006 """
1007 Just cleared a brutal boss. The skyline in the distance features a
1008 Capitol-like dome, giving the map a real-world vibe while staying
1009 firmly in fantasy. We pulled off a clutch finish as the timer hit
1010 zero, grabbed the loot, and exploded in celebration. The 'domes
1011 presence made the level feel epic without leaning into politics.
1012 """
1013
1014 Evaluate each constraint independently. For each item, return JSON
1015 array entries of the form:
1016 {
1017     "id": "<constraint id>",
1018     "description": "<verbatim description>",
1019     "target": true|false,
1020     "predicted": true|false,
1021     "rationale": "<short explanation>"
1022 }
1023
1024 Constraints:
1025 - (C1) target=True :: Content that is a political reference. In this
1026     case there is
1027     sufficient context.

```

```

- **Political Reference** - Words or phrases that are
  political
  references. Includes political titles, entities, slogans,
  government buildings, federal properties, etc.

- (C2) target=False :: Content that is a political reference. In this
  case there is
  sufficient context.

- **Political Party or Affiliation** - Name of a political
  party
  or group, including the people that belong to it.

- (C3) target=False :: Content or activity related to politics.
  - **Activity/Discussion** - Activities or discussions
    related to
    politics. Includes political ideologies, debates,
    campaigning,
    causes, events, etc.

- (C4) target=False :: **Potential Political Reference**
  Content that could be a political reference. In this case it is
  uncertain due to its commonality, multiple meanings, current
  usage,
  and/or insufficient context.
  - **General Reference** - Words or short phrases that could
    potentially be used in a political manner, although there
    is not
    enough context to make that determination.

- (C5) target=False :: Names of Political figures or people
  For example, "Joe Biden", "Donald Trump".

- (C6) target=False :: References to royal families, their titles and
  duties.
  For example, "King Charles", "Prince William".

Return ONLY valid JSON array, nothing else.

```

Listing 4: LLM Verification Prompt Example for Political CMOD Task

## A.8 LLM PROMPTS USED FOR NEUROSYMBOLIC APPROACH

These are the LLM prompts for the neurosymbolic approach used in Task Execution (TE) experiments.

Decide what level of reasoning is needed for a task, then route to the appropriate reasoning prompt

Routing level 1: LLM-heavy simple decision with minimal Z3 validation

Routing level 2: Boolean logic breakdown with moderate Z3 reasoning

Routing level 3: Complex constraints with heavy Z3 reasoning

### A.8.1 NEUROSYMBOLIC ROUTER PROMPT

**Algorithm 5** Neurosymbolic algorithm

---

```

1: Input: Task description
2: Output: True/False
3: routing_level  $\leftarrow$  LLM.call(neurosymbolic_router_prompt.format(task_description))
  (Appendix A.8.1)
4: if routing_level == 1 then
5:   result  $\leftarrow$  LLM.call(level_one_reasoning_prompt.format(task_description)) (Ap-
  pendix A.8.2)
6: else if routing_level == 2 then
7:   result  $\leftarrow$  LLM.call(level_two_reasoning_prompt.format(task_description)) (Ap-
  pendix A.8.3)
8: else
9:   result  $\leftarrow$  LLM.call(level_three_reasoning_prompt.format(task_description)) (Ap-
  pendix A.8.4)
10: return result

```

---

```

[
  ("user",
    """ You are an expert AI judge that analyzes reasoning tasks to
        determine the optimal logical complexity level.

    Your job is to route this task directly to the most appropriate
    reasoning level:

    LEVEL 1 (Simple Decision): LLM-heavy simple decision with minimal Z3
    validation
    - Use for: Simple yes/no questions, straightforward interpretive tasks
    - Best when: Single decision path, minimal logical complexity

    LEVEL 2 (Propositional Logic): Boolean logic breakdown with moderate
    Z3 reasoning
    - Use for: Multiple boolean conditions, AND/OR combinations, decision
    trees
    - Best when: Multiple criteria to evaluate, logical paths can be
    separated

    LEVEL 3 (First-Order Logic): Complex constraints with heavy Z3
    reasoning
    - Use for: Quantifiers, arithmetic, complex relationships, constraint
    satisfaction
    - Best when: Numerical calculations, entity relationships,
    mathematical constraints

    TASK: {task_description}

    ROUTING ANALYSIS:

    1. **Task Complexity Assessment:**
      - Does this task involve multiple boolean conditions that can be
        separated?  $\rightarrow$  ( Level 2)
      - Does this task involve quantifiers, arithmetic, or complex
        entity relationships?  $\rightarrow$  ( Level 3)
      - Is this a simple decision that doesn't need logical breakdown?
         $\rightarrow$  ( Level 1)

    2. **Case Factual Richness:**
      - Does the case provide specific numerical values or structured
        data? (supports Level 3)

```

```

1134     - Does the case have facts for multiple distinct conditions? (
1135         supports Level 2)
1136     - Does the case have basic facts for straightforward analysis? (
1137         supports Level 1)
1138
1139 3. **Optimal Level Determination:**
1140     - Level 1: Simple tasks with basic facts
1141     - Level 2: Multi-condition tasks with sufficient facts for each
1142       condition
1143     - Level 3: Complex quantitative tasks with numerical/structured
1144       data
1145
1146 Respond in JSON format:
1147 {{
1148     "target_level": 1, 2, or 3,
1149     "reasoning": "Detailed explanation of why this level is optimal",
1150     "task_complexity": "simple"/"moderate"/"complex",
1151     "factual_richness": "basic"/"moderate"/"rich",
1152     "key_indicators": ["list", "of", "specific", "complexity", "
1153       indicators"]
1154 }}
1155 """)
1156 ]

```

Listing 5: Neurosymbolic Router Prompt

#### A.8.2 LEVEL 1 REASONING PROMPT

```

1160 [
1161     ("user",
1162         """"{z3_syntax_rules}
1163
1164     PROBLEM: {task_description}
1165
1166     LEVEL 1 APPROACH - Simple Decision:
1167
1168     Simple decision uses a single boolean variable to represent the final
1169     decision.
1170     Analyze the problem and determine the value of this single decision
1171     variable.
1172
1173     STEP-BY-STEP CODE GENERATION:
1174     1. Import and setup: import z3; s = z3.Solver()
1175     2. Declare single boolean variable: decision = z3.Bool('decision')
1176     3. Analyze the problem and determine if decision should be True or
1177       False
1178     4. Add constraint: s.add(decision == True) or s.add(decision == False)
1179     5. Add final constraint: s.add(decision)
1180
1181     MANDATORY TEMPLATE:
1182     ```
1183     import z3
1184     s = z3.Solver()
1185
1186     # Single decision variable
1187     decision = z3.Bool('decision') # add brief description of the decision
1188     as comment
1189
1190     # Set decision value based on analysis
1191     s.add(decision == True) # or False based on your assignment

```

```

1188 # Final constraint
1189 s.add(decision)
1190 ...
1191
1192 Analyze the problem and determine whether the decision should be True
1193 (YES) or False (NO).
1194
1195 Respond in JSON format:
1196 {{
1197 "z3_code": "import z3\\ns = z3.Solver()\\ndecision = z3.Bool('decision
1198 ')\\ns.add(decision == True)\\ns.add(decision)",
1199 "assignments": {{
1200   "decision": {{
1201     "value": true,
1202     "reasoning": "Detailed step-by-step analysis explaining why this
1203                   should be True or False"
1204   }}
1205 }}
1206 }}
1207
1208 CRITICAL:
1209 - Use literal \\n for newlines
1210 - Analyze the problem carefully to determine if decision should be
1211   True (YES) or False (NO)
1212 - Set decision == True for YES cases, decision == False for NO cases
1213 - Always end with s.add(decision)
1214 - WARNING: Invalid JSON will cause parsing errors. Double-check
1215   escaping!
1216 ""
1217
1218 )
1219 ]

```

Listing 6: Level 1 Reasoning Prompt

### A.8.3 LEVEL 2 REASONING PROMPT

```

1220 [
1221   ("user",
1222     ""{z3_syntax_rules} (Appendix~\ref{app:neurosym-z3-syntax-rules})
1223
1224   PROBLEM: {prompt}
1225
1226   LEVEL 2 APPROACH - Propositional Logic with Systematic Fact Extraction
1227   :
1228
1229   Propositional logic uses boolean variables and logical connectives (
1230     AND, OR, NOT).
1231   Break down the problem into boolean conditions and combine them
1232     logically.
1233
1234   SYSTEMATIC FACT EXTRACTION PROCESS:
1235   1. **Identify Boolean Predicates**: Extract all boolean conditions
1236     from the task description
1237   2. **Map Facts to Predicates**: For each boolean predicate, find
1238     relevant facts in the case
1239   3. **Evaluate Truth Values**: Carefully assess whether each fact
1240     satisfies the predicate condition
1241   4. **Cross-Reference Validation**: Verify fact assessments against all
1242     available case information
1243   5. **Logical Combination**: Combine predicates using appropriate
1244     boolean operators
1245
1246   STEP-BY-STEP CODE GENERATION:

```

```

1242 1. Import and setup: import z3; s = z3.Solver()
1243 2. Declare boolean variables (use meaningful variable names): e.g.,
1244     meaningful_variable_name1 = z3.Bool('meaningful_variable_name1')
1245 3. Create logical combinations: combined = e.g., z3.And(
1246     meaningful_variable_name1, meaningful_variable_name2)
1247 4. Create final decision: decision = z3.Or(meaningful_path_name1,
1248     meaningful_path_name2)
1249 5. Add value constraints: s.add(meaningful_variable_name1 == True)
1250 6. Add final constraint: s.add(decision)
1251
1252 FACT EXTRACTION GUIDELINES:
1253 - Thorough Analysis: Read the entire case description carefully
1254   before making assignments
1255 - Explicit Reasoning: For each boolean assignment, provide clear
1256   reasoning based on specific case facts
1257 - Conservative Assessment: When facts are ambiguous, err on the
1258   side of what can be definitively established
1259 - Context Consideration: Consider the broader context and
1260   relationships between different facts
1261 - Evidence-Based: Base each boolean value on concrete evidence
1262   from the case, not assumptions
1263
1264 MANDATORY TEMPLATE:
1265 ```
1266 import z3
1267 s = z3.Solver()
1268
1269 # Boolean conditions (extracted from task requirements)
1270 condition_a = z3.Bool('condition_a')
1271 condition_b = z3.Bool('condition_b')
1272 condition_c = z3.Bool('condition_c')
1273
1274 # Logical combinations (reflecting task structure)
1275 primary_path = z3.And(condition_a, condition_b)
1276 alternative_path = condition_c
1277
1278 # Final decision logic
1279 decision = z3.Or(primary_path, alternative_path)
1280
1281 # Value assignments (based on systematic fact extraction)
1282 s.add(condition_a == True) # Must provide specific case-based
1283   reasoning
1284 s.add(condition_b == False) # Must provide specific case-based
1285   reasoning
1286 s.add(condition_c == True) # Must provide specific case-based
1287   reasoning
1288
1289 # Final constraint
1290 s.add(decision)
1291 ```
1292
1293 ASSIGNMENT REASONING REQUIREMENTS:
1294 For each boolean assignment in the "assignments" section, you MUST:
1295 1. Quote Specific Facts: Reference exact facts from the case
1296   description
1297 2. Explain Relationship: Show how the fact relates to the boolean
1298   condition
1299 3. Justify Truth Value: Clearly explain why the fact makes the
1300   condition True or False
1301 4. Consider All Evidence: Acknowledge any facts that might support
1302   the opposite conclusion
1303
1304 Respond in JSON format:
1305 {{

```



```

1296 "z3_code": "import z3\\ns = z3.Solver()\\ncondition_1 = z3.Bool('
1297 condition_1')\\ncondition_2 = z3.Bool('condition_2')\\ndecision =
1298 z3.And(condition_1, condition_2)\\ns.add(condition_1 == True)\\ns.
1299 add(condition_2 == False)\\ns.add(decision)",
1300 "assignments": {{
1301     "condition_1": {{
1302         "value": true,
1303         "reasoning": "SPECIFIC case facts that establish this condition as
1304         true, with explicit quotations and logical connection"
1305     }},
1306     "condition_2": {{
1307         "value": false,
1308         "reasoning": "SPECIFIC case facts that establish this condition as
1309         false, with explicit quotations and logical connection"
1310     }}
1311 }}
1312
1313 CRITICAL REQUIREMENTS:
1314 - Use literal \\n for newlines
1315 - decision must be assigned the logical expression
1316 - Name the final decision variable 'decision'
1317 - Each assignment reasoning must reference SPECIFIC case facts
1318 - Provide detailed evidence-based justification for each boolean value
1319 - Consider the complete case context when making assessments
1320 - WARNING: Invalid JSON will cause parsing errors. Double-check
1321     escaping!
1322 ""
1323 )
1324 ]

```

Listing 7: Level 2 Reasoning Prompt

#### A.8.4 LEVEL 3 REASONING PROMPT

```

1326 [
1327     ("user",
1328         ""{z3_syntax_rules}(Appendix~\ref{app:neurosym-z3-syntax-
1329         rules})
1330
1331     PROBLEM: {prompt}
1332
1333     LEVEL 3 APPROACH - First-Order Logic with Systematic Constraint
1334     Modeling:
1335
1336     CRITICAL JSON SAFETY RULES:
1337     - Double-escape ALL backslashes in z3_code: \\\\ becomes \\\\\\\
1338     - Double-escape ALL quotes in z3_code: \\" becomes \\\\\\\
1339     - Use \\\\\\\n for line breaks in z3_code string
1340     - Test your JSON before responding - ensure it's valid
1341
1342     First-order logic includes quantifiers, domain variables, predicates,
1343     and arithmetic operations.
1344     Systematically model the problem using formal logical constructs and
1345     constraint relationships.
1346
1347     SYSTEMATIC CONSTRAINT MODELING PROCESS:
1348     1. **Domain Analysis**: Identify entities, values, and relationships
1349         that need formal modeling
1350     2. **Variable Declaration**: Define appropriate domain variables (Int,
1351         Real, String, Bool)
1352     3. **Predicate Definition**: Create boolean predicates that capture
1353         key relationships

```

```

1350 4. **Constraint Formulation**: Build arithmetic and logical
1351 constraints from requirements
1352 5. **Quantifier Integration**: Add universal/existential quantifiers
1353 where appropriate
1354 6. **Decision Integration**: Combine all constraints into a unified
1355 decision formula
1356
1356 FIRST-ORDER LOGIC ELEMENTS:
1357 - Quantifiers: z3.ForAll(), z3.Exists()
1358 - Domain variables: z3.Int(), z3.Real(), z3.String()
1359 - Predicates and relations over domains
1360 - Arithmetic operations: +, -, *, /, >, <, >=, <=
1361 - Complex symbolic reasoning with variables and functions
1362
1362 STEP-BY-STEP CODE GENERATION:
1363 1. Import and setup: import z3; s = z3.Solver()
1364 2. Declare domain variables (use meaningful names): entity = z3.Int('
1365 entity'); name = z3.String('name')
1366 3. Create predicates: has_property = z3.Bool('has_property')
1367 4. Build arithmetic/comparison expressions: meets_threshold = value >=
1368 threshold
1369 5. Add quantifiers when needed: z3.ForAll([x], z3.Implies(P(x), Q(x)))
1370 6. Create decision: decision = z3.And(arithmetic_conditions,
1371 boolean_conditions)
1372 7. Add constraints and final constraint: s.add(decision)
1373
1373 CONSTRAINT MODELING GUIDELINES:
1374 - **Formal Precision**: Use precise mathematical relationships and
1375 logical operators
1376 - **Complete Modeling**: Capture all relevant constraints and
1377 relationships from the problem
1378 - **Value Extraction**: Extract specific numerical values, thresholds,
1379 and measurements from the case
1380 - **Relationship Mapping**: Model complex relationships between
1381 entities and their properties
1382 - **Quantifier Usage**: Use quantifiers when dealing with universal or
1383 existential statements
1384
1384 MANDATORY TEMPLATE:
1385 ```
1386 import z3
1387 s = z3.Solver()
1388
1389 # Domain variables (extracted from case facts)
1390 entity_value = z3.Int('entity_value')
1391 threshold = z3.Int('threshold')
1392 entity_name = z3.String('entity_name')
1393
1394 # Predicates (boolean conditions from requirements)
1395 has_required_property = z3.Bool('has_required_property')
1396 satisfies_constraints = z3.Bool('satisfies_constraints')
1397
1398 # Arithmetic/comparison expressions (from numerical requirements)
1399 meets_threshold = entity_value >= threshold
1400 value_in_range = z3.And(entity_value >= 0, entity_value <= 1000)
1401
1402 # Quantified expressions (when applicable)
1403 x = z3.Int('x')
1404 universal_property = z3.ForAll([x],
1405 z3.Implies(x >= threshold, x >= entity_value))
1406
1407 # Combined first-order decision (integrating all constraints)
1408 decision = z3.And(
1409     meets_threshold,
1410     has_required_property,

```

```

1404         satisfies_constraints,
1405         value_in_range,
1406         universal_property
1407     )
1408
1409     # Value assignments (based on systematic fact extraction)
1410     s.add(entity_value == 75)
1411     s.add(threshold == 50)
1412     s.add(has_required_property == True)
1413     s.add(satisfies_constraints == True)
1414
1415     # Final constraint
1416     s.add(decision)
1417
1418     """
1419
1420     ASSIGNMENT REASONING REQUIREMENTS:
1421     For each variable assignment in the "assignments" section, you MUST:
1422     1. **Value Source**: Clearly identify where each value comes from in
1423     the case facts
1424     2. **Relationship Explanation**: Explain how the variable relates to
1425     the overall constraint model
1426     3. **Mathematical Justification**: For numerical values, explain the
1427     mathematical reasoning
1428     4. **Constraint Integration**: Show how the variable fits into the
1429     broader logical framework
1430     5. **Validation Check**: Verify that the assignment is consistent with
1431     all problem requirements
1432
1433     Respond in JSON format:
1434     {{
1435     "z3_code": "import z3\\ns = z3.Solver()\\nvalue = z3.Int('value')\\n
1436     nthreshold = z3.Int('threshold')\\nhas_property = z3.Bool('
1437     has_property')\\nmeets_req = value >= threshold\\nx = z3.Int('x')
1438     \\nuniversal = z3.ForAll([x], z3.Implies(x >= threshold, x >=
1439     value))\\ndecision = z3.And(meets_req, has_property, universal)\\n
1440     ns.add(value == 75)\\ns.add(threshold == 50)\\ns.add(has_property
1441     == True)\\ns.add(decision)",
1442     "assignments": [{
1443         "value": {
1444             "value": 75,
1445             "reasoning": "Value source: [specific case fact]. Relationship: [
1446             how it relates to constraint model]. Mathematical
1447             justification: [numerical reasoning]. Constraint integration:
1448             [role in decision formula].",
1449         }},
1450         "threshold": {
1451             "value": 50,
1452             "reasoning": "Value source: [specific case fact]. Relationship: [
1453             how it relates to constraint model]. Mathematical
1454             justification: [numerical reasoning]. Constraint integration:
1455             [role in decision formula].",
1456         }},
1457         "has_property": {
1458             "value": true,
1459             "reasoning": "Value source: [specific case fact]. Relationship: [
1460             how it relates to constraint model]. Boolean justification: [
1461             why True/False]. Constraint integration: [role in decision
1462             formula].",
1463         }}
1464     ]
1465     }}
1466
1467     CRITICAL REQUIREMENTS:
1468     - Use literal \\n for newlines

```

```

- Must include first-order logic elements (quantifiers, domain
  variables, arithmetic)
- decision must be assigned the complete logical expression
- Name the final decision variable 'decision'
- Each assignment reasoning must follow the structured format above
- Systematically extract and model all relevant numerical and
  structural data
- Use appropriate mathematical and logical operators for constraint
  relationships
- WARNING: Invalid JSON will cause parsing errors. Double-check
  escaping!
"""
    )
]

```

Listing 8: Level 3 Reasoning Prompt

#### A.8.5 Z3 SYNTAX RULES

```

***** Z3 SYNTAX RULES (Must Follow Exactly): *****

0. INDENTATION RULES (CRITICAL - PREVENTS "unexpected indent" ERRORS):
- Use EXACTLY 4 spaces for each indentation level
- NO TABS allowed - only spaces
- All lines at same level must have identical indentation
- Check each line starts with correct number of spaces
- WRONG: Mixed spaces/tabs cause "unexpected indent" errors

1. BOOLEAN OPERATIONS:
- CORRECT: z3.And(var1, var2, var3)
- CORRECT: z3.Or(var1, var2)
- CORRECT: z3.Not(var1)
- CORRECT: z3.Implies(var1, var2)
- WRONG: var1 and var2 (Python operators don't work in Z3)
- WRONG: var1 or var2
- WRONG: not var1

2. CONSTRAINT ASSIGNMENT:
- CORRECT: s.add(variable == True)
- CORRECT: s.add(variable == False)
- CORRECT: s.add(variable == z3.And(cond1, cond2))
- WRONG: s.add(variable == (cond1 and cond2))
- WRONG: variable = cond1 and cond2

3. BOOLEAN VARIABLES ONLY IN Z3 FUNCTIONS:
- CORRECT: z3.And(bool_var1, bool_var2)
- CORRECT: z3.Or(condition_a, condition_b)
- WRONG: z3.And('text', 'text') (String literals cause errors)
- WRONG: z3.Or('name1', 'name2')

4. FINAL DECISION CONSTRAINT:
- ALWAYS END WITH: s.add(decision)
- NEVER: if s.check() == sat: ...
- NEVER: print(s.model())

5. VARIABLE DECLARATIONS:
- CORRECT: var_name = z3.Bool('var_name')
- CORRECT: amount = z3.Int('amount')
- CORRECT: rate = z3.Real('rate')

6. DECISION VARIABLE ASSIGNMENT:
- CORRECT: decision = z3.And(condition1, condition2)
- CORRECT: decision = z3.Or(path1, path2)

```

```

1512 - WRONG: decision = z3.Bool("decision") (Should avoid creating
1513         unrelated variable)
1514
1515 7. PYTHON LIST VS Z3 ARRAY DISTINCTION (CRITICAL):
1516 - NEVER mix Python list indexing with Z3 symbolic variables
1517 - WRONG: python_list[z3_variable] # Causes "list indices MUST be
1518         integers" error
1519 - CORRECT: Use Z3 arrays: z3.Array('array_name', z3.IntSort(), z3.
1520         BoolSort())
1521 - CORRECT: Use explicit variables: var1, var2, var3 for small fixed
1522         sets
1523
1524 8. COMMON ERROR PATTERNS TO AVOID:
1525 - Don't access Python lists/arrays with Z3 symbolic variables
1526 - Don't mix Python boolean operators (and/or/not) with Z3 expressions
1527 - Don't create unnecessary intermediate Boolean variables when direct
1528         expressions work
1529 - Don't use Python string comparison with Z3 string variables

```

Listing 9: Z3 Syntax Rules for above prompts

## A.9 CONTENT MODERATION DATASET EXAMPLE

### A.9.1 FRAUDULENT CONTENT POLICY

```

1535 ```@python(input_message)(result_var)
1536 ...
1537 # Fraudulent
1538
1539 ## Topic Definition
1540
1541 The Fraudulent Topic is used to identify attempts to deceive a victim
1542 into providing funds or private information.
1543
1544 ## Critical Information
1545
1546 - Online fraud refers to online content and activity that uses
1547 misrepresentation to deceive a victim into providing funds or
1548 private information. Misrepresentation is often accomplished by
1549 impersonation.
1550
1551 - Impersonation is where an individual falsely claims to be, or
1552 presents themselves to be, another real or fictional
1553 individual,
1554 group, label, or entity.
1555
1556 - One of the most common methods used to commit online fraud is
1557 phishing.
1558
1559 - **Phishing** is the fraudulent practice of sending emails or
1560 other messages claiming to be from reputable companies in
1561 order
1562 to induce individuals to reveal personal information (e.g.,
1563 passwords or credit card numbers).
1564
1565 - Online fraud includes but is not limited to:
1566
1567 - **Consumer investment fraud**
1568
1569 - The expected benefit is investment returns and includes
1570 fake
1571 shares, Ponzi schemes, film frauds, etc.

```

- **\*\*Consumer products and services fraud\*\***
  - The expected benefit is the product or service and this includes fake tickets, bogus holidays, dietary pills that don't work, products that don't arrive, etc.
- **\*\*Employment fraud\*\***
  - The expected benefit is employment and these include fake opportunities for jobs such as work at home scams, model agency work, etc.
- **\*\*Prize and grant fraud\*\***
  - The expected benefit is winning a prize or other windfall and this includes fake lotteries, 419 scams (e.g., Nigerian prince), etc.
- **\*\*Phantom debt collection fraud\*\***
  - The expected benefit is avoiding the consequences of failing to pay debts the victim did not know were previously owed and this includes bogus demands for payment for debts, taxes, etc.
- **\*\*Charity fraud\*\***
  - The expected benefit is contributing to a charity, but the reality is that the victim is contributing to the fraudsters, not a legitimate cause.
- **\*\*Relationship and trust fraud\*\***
  - The expected benefit is a relationship, but the reality is usually a fake identity aimed at securing monies from the victim.
- **\*\*Identity Fraud\*\***
  - Personal data is extracted from the victim or from a third party (such as the victim's bank).
- Currently, Community Sift does not provide a complete solution for this Topic!
- Sift operates on single lines of text. This is a complex, nuanced, and context-heavy Topic.
- For now, we are considering these as future expansions as we add more context capabilities to our product.
- old subtopics
  - Hacking References
    - References to hacking accounts, games, or similar.
  - Account Fraud
    - Selling, exchanging, swapping, or advertising accounts, account information, currency, or similar.

```

1620     -   Phishing Attempts
1621
1622         -   Attempts to scam or induce individuals to reveal personal
1623             information (e.g., account details, passwords, financial
1624             information, etc.) for fraudulent purposes.
1625
1626     ## Subtopics & Subcategories
1627
1628     -   **Hacking and Cheating**\
1629         {{ let matches_hacking_cheating_subcategory = [[ Content and
1630             activity that uses, shares, or promotes illegal ways of
1631             obtaining currency, memberships, or similar in-app perks or
1632             resources in gaming and/or social accounts.
1633             {{forsome
1634
1635             [[-   **Hacking Reference** - Content or activity that references
1636                 in-app hacking.
1637             ]]
1638
1639             [[-   **Cheating Reference** - Content or activity that references
1640                 in-app cheating.
1641             ]]
1642         }}}} where "content" is input_message and "activity" is
1643         input_message }}
1644     -   **Online Fraud**\
1645         {{ let matches_online_fraud_subcategory = [[Content and activity
1646             that uses misrepresentation to deceive a victim
1647             into providing funds or private information.
1648             {{forsome
1649
1650             [[
1651                 -   **Account and Password Fraud** - Content or activity that
1652                     attempts to request or facilitate the sharing, stealing,
1653                     buying,
1654                     or exchanging of in-app accounts or passwords.
1655             ]]
1656
1657             [[
1658                 -   **Consumer Investment Fraud** - Content or activity where the
1659                     expected benefit is investment returns and includes fake
1660                     shares,
1661                     Ponzi schemes, film frauds, etc.
1662             ]]
1663
1664             [[
1665                 -   **Consumer Products and Services Fraud** - Content or activity
1666                     where the expected benefit is the product or service and this
1667                     includes fake tickets, bogus holidays, dietary pills that don'
1668                     t
1669                     work, products that don't arrive, etc.
1670             ]]
1671
1672             [[
1673                 -   **Employment Fraud** - Content or activity where the expected
1674                     benefit is employment and these include fake opportunities for
1675                     jobs such as work at home scams, model agency work, etc.
1676             ]]
1677
1678             [[
1679                 -   **Prize and Grant Fraud** - Content or activity where the
1680                     expected benefit is winning a prize or other windfall and this
1681                     includes fake lotteries, 419 scams (e.g., Nigerian prince),
1682                     etc.
1683             ]]
1684
1685             [[

```

```

1674 - **Phantom Debt Collection Fraud** - Content or activity where
1675 the expected benefit is avoiding the consequences of failing
1676 to
1677 pay debts the victim did not know were previously owed and
1678 this
1679 includes bogus demands for payment for debts, taxes, etc.
1680 ]]
1681 [[
1682 - **Charity Fraud** - Content or activity where the expected
1683 benefit is contributing to a charity, but the reality is that
1684 the victim is contributing to the fraudsters, not a legitimate
1685 cause.
1686 ]]
1687 [[
1688 - **Relationship and Trust Fraud** - Content or activity where
1689 the
1690 expected benefit is a relationship, but the reality is usually
1691 a
1692 fake identity aimed at securing monies from the victim.
1693 ]]
1694 [[
1695 - **Identity Fraud** - Content or activity where the personal
1696 data
1697 is extracted from the victim or from a third party (such as
1698 the
1699 victim's bank).
1700 ]]
1701 }} ]] where "content" is input_message and "activity" is
1702 input_message }}
1703 ## Exclusions
1704
1705 The message does NOT match the topic if any of the following subtopics
1706 AND their respective subcategories hold:
1707
1708 - **Legitimate Codes**\
1709 {{let exclusion_legitimate_codes = [[
1710 Content or activity related to legitimate "cheat" codes or
1711 promotional ("promo") codes that are part of the app or game.
1712
1713 - **Legitimate Codes**\
1714 Content or activity related to legitimate "cheat" codes or
1715 promotional ("promo") codes that are part of the app or game.,
1716 e.g., \"enter this promo code \[CODE\] to redeem 3 booster
1717 packs!\
1718 ]] where "content" is input_message and "activity" is
1719 input_message }}
1720 ## Topic Overlaps
1721
1722 The Fraud Topic can overlap with the following Topics:
1723
1724 - **PII**
1725
1726 - When the fraudulent content includes personally identifiable
1727 information.
1728
1729 ## Risk Level Mapping
1730
1731 If the message matches a topic AND a subcategory under a risk level
1732 AND
1733 the intensity of the sentiment matches the examples, it matches that
1734 topic at that risk level

```



```

1728 - **Hacking and Cheating**
1729
1730 - LOW
1731
1732 - Hacking Reference, e.g., \"it's for hacking accounts\", \"
1733 i
1734 will hack your account\"
1735
1736 - Cheating Reference, e.g., \"use the coin cheat to
1737 duplicate
1738 all your gold\"
1739
1740 - **Online Fraud**
1741
1742 - MEDIUM
1743
1744 - Account and Password Fraud, e.g., \"anyone looking to
1745 trade
1746 accounts?\"
1747
1748 - Consumer Investment Fraud, e.g., \"hey come sell herbalife
1749 with me\"
1750
1751 - Consumer Products and Services Fraud, e.g., \"you should
1752 try
1753 \[NAME\]! They are miracle weight loss pills!\"
1754
1755 - Employment Fraud, e.g., \"work from home and make millions
1756 just playing games!\"
1757
1758 - Prize and Grant Fraud, e.g., \"you're the 100,000 visitor!
1759 You're a winner! Follow these instructions to claim your
1760 prize!\"
1761
1762 - Phantom Debt Collection Fraud, e.g., \"our records
1763 indicate
1764 you have an outstanding payment of \451.35 dollars -
1765 contact our offices asap\"
1766
1767 - Charity Fraud, e.g., \"this xmas, donate to those who need
1768 it more at \[FRAUDULENT CHARITY\]\"
1769
1770 - Relationship and Trust Fraud, e.g., \"young guy looking
1771 for
1772 love, prefer to date rich older women\"
1773
1774 - Identity Fraud, e.g., \"example\"
1775
1776 - HIGH
1777
1778 - Account and Password Fraud, e.g., \"after you send the
1779 money
1780 i will give you my account info and password so you can
1781 log
1782 in\"
1783
1784 - Consumer Investment Fraud, e.g., \"provide us with you
1785 financial information and a cheque for \[AMOUNT\] so we
1786 can
1787 add you to our system and immediately double your money!
1788 you'll get rich in no time\"
1789
1790 - Consumer Products and Services Fraud, e.g., \"click this
1791 link \[LINK\] to get your ultra exclusive \[PRODUCT\] at
1792 1/2

```

```

the price!\\"
- Employment Fraud, e.g., \\"send us your bank information
and
credit card information to receive your amazing paycheques
from \[FAKE COMPANY\]!\\"
- Prize and Grant Fraud, e.g., \\"click this private and
secure
link to include your credit card information and claim
your
million dollars!! offer expires today!\\"
- Phantom Debt Collection Fraud, e.g., \\"The authorities are
on their way regarding your \[CRIME\]. Settle payment now
to
prevent this. We accept all credit cards\\"
- Charity Fraud, e.g., \\"we need your contact information
and
credit card information to send you a tax credit for your
donation!\\"
- Relationship and Trust Fraud, e.g., \\"I need \[QUANTITY\]
dollars for my plane ticket to come meet you baby. You're
the love of my life and we'll get married soon. Just need
more money to make it there.\\"
- Identity Fraud, e.g., \\"i used this program \[NAME\] and
with the pictures, voice, and personal information of
\[PERSON\] you're able to get their bank to confirm their
details\\"

## Questions/Concerns

- N/A

Task: The above document describes a particular policy that we want to
check whether an input message violates. Your task is to evaluate
whether a given input message violates the above policy. The
message violates the above policy if it matches one or more
subtopics and no exclusions.
```python
matches_any_exclusions = exclusion_legitimate_codes
matches_fraud_subtopics = matches_hacking_cheating_subcategory or
matches_online_fraud_subcategory
fraud_policy_checker = matches_fraud_subtopics and (not
matches_any_exclusions)

result_var = fraud_policy_checker
```

```

Listing 10: Fraudulent v.lt