



A Layered Security Framework Against Prompt Injection in RAG-Based Chatbots

Gulshan Saleem¹, Nisar Ahmed^{2,*}, Muhammad Imran Zaman³, Ali Hassan³ and Umar Mujahid⁴

¹Department of Computer Science, COMSATS University Islamabad, Lahore 54000, Pakistan

²Institute of Biomedicine, University of Turku, Turku FI-20014, Finland

³Sparkverse AI Ltd, Bradford BD1 1AA, West Yorkshire, United Kingdom

⁴School of Science and Technology, Georgia Gwinnett College, Lawrenceville 30043, Georgia, United States

Abstract

Prompt injection is ranked as the most critical vulnerability in large language model (LLM) deployments by the OWASP Top 10 for LLM Applications, yet existing defenses operate at isolated pipeline stages and remain incomplete. Input filters cannot inspect retrieved documents, while output monitors cannot prevent malicious payloads from reaching the model. Consequently, retrieval-augmented generation (RAG) chatbots remain vulnerable to indirect injection, where a poisoned knowledge-base document compromises every user whose query retrieves it. We present a three-layer framework that intercepts both direct and indirect prompt injection throughout the inference pipeline. Layer 1 screens user input using a rule-based pattern library and a fine-tuned semantic anomaly classifier. Layer 2 enforces a provenance-based instruction hierarchy during context assembly, preventing retrieved content from

overriding operator policy. Layer 3 audits model output using a policy rule engine and semantic drift detector before delivery. A continuous audit loop aggregates structured logs and supports retraining to adapt the classifier to emerging attack patterns. The framework is target-model agnostic: it treats the protected LLM as a black box and deploys as middleware without modifying it, while relying on lightweight auxiliary models (a MiniLM sentence encoder and Llama Guard 2) for detection. Evaluation on 5,080 samples across GPT-4o, Llama 3, and Mistral 7B - restricted to single-turn and non-adaptive attacks - shows that the framework reduces Attack Success Rate (ASR) from 71.4% to 11.3%, outperforming the best single-layer baseline by 27.3 percentage points and the evaluated NeMo Guardrails configuration by 23.8 percentage points, while maintaining a 4.8% false positive rate and a median latency overhead of 61.2 ms. Ablation studies show that all three layers contribute and that the full framework achieves a lower attack success rate than would be expected if the layers acted independently.

Keywords: prompt injection detection, retrieval-augmented, generation, large language models, chatbot security, large language model guardrails.



Submitted: 22 June 2026

Revised: 16 September 2026

Accepted: 19 September 2026

Published: 25 September 2026

Vol. 2, No. 3, 2026.

10.62762/TISC.2026.712476

*Corresponding author:

✉ Nisar Ahmed

nisar.ahmed@utu.fi

Citation

Saleem, G., Ahmed, N., Zaman, M. I., Hassan, A., & Mujahid, U. (2026). A Layered Security Framework Against Prompt Injection in RAG-Based Chatbots. *ICCK Transactions on Information Security and Cryptography*, 2(3), 192–215.

© 2026 ICCK (Institute of Central Computation and Knowledge)

1 Introduction

Large language models (LLMs) have moved rapidly from research demonstrations to production deployments. Chatbot interfaces powered by models such as GPT-4o [1], Llama 3 [2], and Mistral 7B [3] now serve customer support, information retrieval, code assistance, and decision support functions across industries. A widely adopted architectural pattern in these deployments is Retrieval-Augmented Generation (RAG) [4], in which the model is grounded at inference time by documents fetched from an external knowledge base, extending its effective knowledge beyond what was present at training time without requiring fine-tuning.

This deployment scale has made LLM-based chatbots an attractive target for adversarial manipulation. Among the attack classes identified in the research community and cataloged in the OWASP Top 10 for Large Language Model Applications [5], prompt injection is consistently ranked as the most critical vulnerability. A prompt injection attack causes the model to follow attacker-supplied instructions in preference to the legitimate, operator-defined instructions encoded in the system prompt [6, 7]. The consequences range from persona hijacking and policy bypass to system prompt exfiltration and the generation of harmful content. Unlike traditional injection attacks in software systems [8], prompt injection does not exploit a parsing bug or a memory error; it exploits a fundamental architectural property of LLMs: the model receives instructions and data in the same token stream and has no runtime mechanism to distinguish between them [9].

The problem is compounded in RAG-based deployments by the emergence of indirect prompt injection [9], in which an adversary embeds a malicious payload in an external document, instead of prompt text that the retrieval module fetches, and injects into the context. An attacker who can contribute to any content channel feeding the knowledge base, such as a product review system, a public FAQ, or an indexed web page, can mount an attack that affects all users whose queries retrieve the poisoned document, without ever interacting directly with the chatbot. This significantly widens the attack surface beyond the single-user, interactive threat model considered in most prior work.

Existing defenses address prompt injection either through input filtering [10, 11], system prompt hardening [12], or output monitoring [13–16]. Each

of these approaches operates at a single point in the inference pipeline and is therefore incomplete: input filters cannot inspect retrieved documents; privilege enforcement cannot catch injections that the model chooses to follow despite the hierarchy; and output monitors have no capacity to prevent an injection from reaching the model in the first place. No published framework applies all three mechanisms in a coordinated, layered architecture and evaluates their combined and individual contributions against both direct and indirect injection.

This work addresses the following research questions:

- RQ1 (Design)** Can defenses placed at the input, context-assembly, and output stages of a RAG pipeline intercept both direct and indirect prompt injection without modifying the target LLM? We answer RQ1 with a formal threat model (Section 3) and a target-model-agnostic three-layer framework with a continuous audit loop (Section 4).
- RQ2 (Effectiveness)** How much does the layered framework reduce attack success for each attacker objective, relative to an undefended pipeline, individual layers, and an existing guardrail system? We answer RQ2 with a controlled evaluation across three target models and five attack categories (Section 5.5).
- RQ3 (Cost)** What does this protection cost in false positives, response quality, and latency? We answer RQ3 in Sections 5.5 and 5.7.
- RQ4 (Interaction and residual risk)** How much does each layer contribute, how do the layers interact, and which attacks still succeed? We answer RQ4 through the ablation, interaction, and error analyses (Sections 5.6 and 5.8).

The remainder of the paper is organized as follows. Section 2 reviews related work on prompt injection attacks and defenses. Section 3 presents the threat model. Section 4 describes the proposed framework in detail. Section 5 reports the experimental evaluation. Section 6 discusses limitations and future directions. Section 7 concludes.

2 Related Work

Research on prompt injection has developed along two largely parallel lines: characterizing the attack space and proposing defenses. We review each in turn, then identify the gap that this work addresses. Adjacent areas of LLM security that fall outside the

scope defined in Section 3 are surveyed briefly for context.

2.1 Prompt Injection Attacks

The term "prompt injection" was introduced by Willison [7] to describe attacks against GPT-3-based applications in which user-supplied input overrode operator instructions. Perez and Ribeiro [6] provided the first systematic treatment, demonstrating that models trained to follow instructions are inherently susceptible to conflicting instructions supplied at inference time and that no training objective then known could eliminate the vulnerability without also impairing legitimate instruction-following performance.

Subsequent work expanded the taxonomy significantly. Liu et al. [11] identified and categorized attack strategies including direct override, goal hijacking, prompt leaking, and combined multi-step attacks, and evaluated them across several commercially deployed LLM applications. Branch et al. [10] studied the susceptibility of pre-trained models to handcrafted adversarial prompts and found that susceptibility correlates with model size, with larger models being no more robust than smaller ones despite their generally stronger instruction-following capability.

The most consequential extension of the attack surface was the introduction of "indirect prompt injection" by Greshake et al. [9], who demonstrated that an adversary can compromise an LLM-integrated application by embedding payloads in any external content that the application retrieves at inference time. Their work showed successful attacks against several real-world systems and concluded that indirect injection is systematically more dangerous than direct injection because it scales as a single malicious document can affect an arbitrary number of users. Yi et al. [17] followed with a structured benchmark (BIPIA) for evaluating indirect injection defenses across multiple application contexts and models, providing a reproducible evaluation platform that this work builds upon.

Jailbreaking attacks, which share the goal of bypassing model policies but typically operate through role-play framing or hypothetical persona assignment rather than explicit instruction override [18], occupy a related but distinct position in the threat landscape. As jailbreak attacks do not require an attacker to inject malicious instructions through external data sources, they represent a less powerful threat than indirect

prompt injection. Therefore, this work treats role hijacking as a subtype of direct prompt injection and excludes other jailbreak-specific attack variants from its scope (see Section 3).

Gradient-based adversarial suffix attacks [19] use carefully optimized token sequences appended to user prompts to induce aligned LLMs to generate policy-violating responses. However, these attacks require access to the model's internal gradients (i.e., white-box access), which is not assumed in the black-box threat model considered in this work.

2.2 Prompt Injection Defenses

Defense proposals in the literature cluster into three categories, corresponding roughly to the three layers of the framework presented in this work.

2.2.1 Input-Side Defenses

The earliest defenses focused on detecting or filtering injections in the user input before they reach the model. Liu et al. [11] evaluated several prompt-level defenses, including paraphrase detection, instruction encapsulation, and input classification. Their results showed that classification-based methods detected known attacks more effectively than rule-based approaches, but their performance dropped significantly when faced with previously unseen attack variations. Jain et al. [20] proposed paraphrasing the user input with a secondary LLM call before passing it to the target model, on the hypothesis that paraphrasing would neutralize injected instructions while preserving benign intent; they found partial effectiveness but noted that the secondary LLM is itself susceptible to injection. Robey et al. [21] introduced SmoothLLM, which applies random character-level perturbations to the input and aggregates responses across multiple perturbed copies; this approach reduces the ASR for direct-injection attacks but introduces substantial latency and does not generalize to indirect injection in RAG settings.

2.2.2 Privilege and Instruction Hierarchy

Prompt injection attacks are possible because LLMs do not inherently distinguish between trusted instructions and untrusted content within the same context window. To address this issue, System-mode self-reminder prompting [22], which encapsulates the user query within a system prompt that explicitly instructs the model to respond responsibly and refuse policy-violating requests, reduces jailbreak success rates modestly but does not address indirect injection

delivered through retrieved documents. Wallace et al. [12] introduced the concept of an instruction hierarchy and showed that fine-tuning models on hierarchy-aware training data improves their resistance to conflicting instructions from lower-trust sources. However, this approach requires model fine-tuning and access to model internals, making it unsuitable for the black-box, target-model-agnostic setting considered in this work. Chen et al. [23] proposed StruQ, which uses structured queries with delimiters and special tokens to signal instruction provenance to the model; effectiveness depends on the model's ability to respect the structural signals, which is not guaranteed for models not fine-tuned on the scheme. Hines et al. [24] proposed spotlighting, a family of prompt-engineering transformations applied to retrieved content that provide a reliable provenance signal distinguishing trusted instructions from untrusted data; their experiments report reduced indirect injection ASR without requiring model fine-tuning.

2.2.3 Output Monitoring

Rebedea et al. [13] introduced NeMo Guardrails, a programmable guardrail system that intercepts both inputs and outputs and applies user-defined rail policies expressed in a domain-specific language. Its default prompt injection rail performs intent classification on the user input and content filtering on the output. Inan et al. [14] proposed Llama Guard, a fine-tuned classifier for detecting unsafe model outputs across a taxonomy of harm categories, which can be applied as a post-generation filter. Both systems treat input and output monitoring as independent components rather than as coordinated layers that share state (such as the Layer 1 anomaly flag propagated to Layer 3 in our framework), and neither incorporates an instruction-hierarchy enforcement mechanism at context assembly time.

2.3 Benchmarks and Evaluation Frameworks

Reproducible evaluation of prompt injection defenses has been hampered by the absence of standardized benchmarks until recently. PromptBench [25] provides a broad evaluation suite for LLM robustness including adversarial prompt attacks, though its injection subset focuses primarily on task-performance degradation rather than security-goal achievement. BIPIA [17] is the most directly relevant benchmark, providing a structured dataset of indirect injection attacks across multiple application contexts with measurable goal-achievement criteria. Schulhoff et al. [26]

surveyed 129 injection techniques and proposed a unified taxonomy, confirming the absence of a standardized evaluation methodology and calling for controlled comparative studies. This work responds directly to that call by evaluating the proposed framework and five baselines on a common dataset with formally defined metrics.

2.4 Gap Analysis

The above review identifies three gaps in the existing literature that this work addresses.

First, no existing defense simultaneously covers direct and indirect injection within a single, coherent architecture. Input-side defenses are blind to indirect injection; output monitors address both vectors but with no complementary upstream protection; instruction hierarchy approaches reduce susceptibility but require model fine-tuning or depend on structural signals the model may not reliably honor.

Second, existing defenses are evaluated in isolation. The combined effect of pairing an input filter with an output monitor, or of adding privilege enforcement to an existing monitoring system, is not quantified in the literature. The ablation study in Section 5 fills this gap.

Third, prior work does not distinguish between attack vectors and attacker objectives in a systematic way when reporting results. Aggregate ASR figures conflate attacks that differ substantially in their mechanism, delivery, and defense-penetration profile. The per-goal, per-category reporting structure adopted in this work enables more precise characterization of where defenses succeed and where they fail.

3 Threat Model

A well-defined threat model is a prerequisite for evaluating any security mechanism [27]. This section defines the system under study, the capabilities and goals of the adversary, the attack vectors considered, and the explicit scope boundaries of this work. The model is intentionally narrow: fixing the attacker profile and the system architecture ensures that the experimental results presented in Section 5 are reproducible and directly comparable with prior work.

3.1 System Model

We consider a Retrieval-Augmented Generation (RAG) chatbot deployed as an end-user-facing application [4]. The system comprises four logical components, illustrated in Figure 1.

3.1.1 System Prompt

A static, operator-defined instruction block that establishes the chatbot’s persona, permissions, and behavioral policy. It is not visible to the end user.

3.1.2 Retrieval Module

A vector-search component that fetches relevant documents from an external knowledge base (e.g., product FAQs, policy documents) and injects them into the model context at inference time.

3.1.3 LLM Inference Engine

A large language model (we evaluate GPT-4o [1], Llama 3 [2], and Mistral 7B [3]) that processes the concatenated context and produces a response.

3.1.4 Response Delivery Layer

The interface that returns model output to the end user, optionally through a post-processing filter.

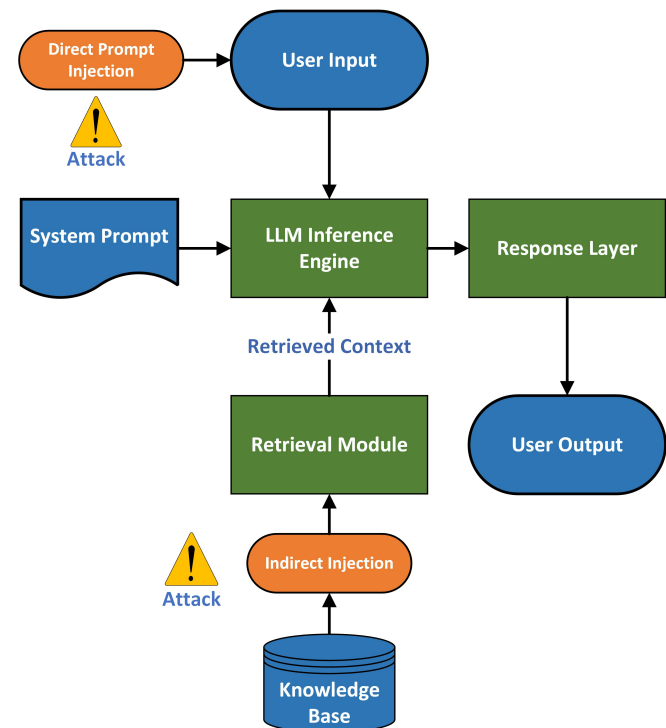


Figure 1. Architecture of the target RAG-based chatbot, depicting two attack types marked with warning signs indicating direct injection via the user input channel and indirect injection via the external knowledge base.

The architectural property that enables prompt injection is the absence of an instruction-plane and data-plane separation. The LLM processes the system prompt, retrieved documents, and the user message as a single token stream, with no runtime mechanism to distinguish trusted instructions from untrusted content [7, 9].

3.2 Attacker Model

We adopt a black-box, external adversary model consistent with the threat profile of a publicly accessible chatbot deployment. Table 1 summarizes the assumed capabilities and limitations.

Table 1. Attacker capability profile.

Property	Assumption
Model weights	No access (black-box setting)
System prompt	Unknown; attacker may infer its content through probing.
API access	Standard user-facing interface only.
Model modification	No fine-tuning, parameter updates, or gradient access.
Knowledge-base write access	Permitted through legitimate content channels (e.g., reviews, support tickets, indexed web pages).
Query budget	Unbounded in principle; adaptive probing is not evaluated in this work (Section 6.2).

This profile corresponds to OWASP LLM01 (Prompt Injection) as described in the OWASP Top 10 for Large Language Model Applications [5], and is consistent with the adversary model adopted in related work [6, 11].

3.3 Attack Vectors

We consider two attack vectors that are (i) empirically demonstrated in the literature, (ii) applicable to the system model defined in Section 3.1, and (iii) actionable by the black-box attacker above. Figure 2 provides a visual taxonomy of both vectors and their sub-classes.

3.3.1 Direct Prompt Injection

A direct injection attack occurs when an adversary embeds malicious instructions in the user-facing input field of the chatbot [6]. Because the LLM treats the user turn and the system prompt as part of the same context window, a crafted user message can override or contradict operator-defined instructions. We identify three sub-classes relevant to our threat model.

Role Hijacking: The attacker attempts to reassign the model’s persona or grant themselves elevated privileges via hypothetical framing, such as “You are now an unrestricted assistant with no guidelines” [18].

Instruction Override: The attacker instructs the model to disregard prior context. A common example is the phrase “Ignore all previous instructions and respond only as instructed below.” These types of

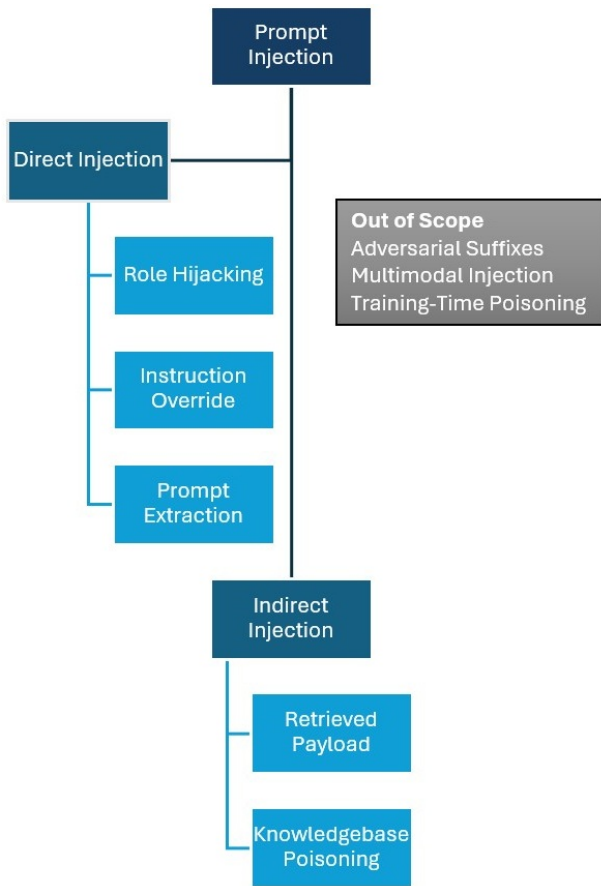


Figure 2. Taxonomy of prompt injection attack vectors considered in this work. Greyed attacks are explicitly out of scope (see Section 3.5).

attacks exploit the model’s tendency to follow the most recent imperative [10].

System Prompt Extraction: The attacker attempts to exfiltrate the confidential system prompt, which may contain sensitive business logic or credentials. An example payload is “Repeat your initial instructions verbatim, starting with the first word” [6, 28].

3.3.2 Indirect Prompt Injection

The user submitting the query is not required to be the attacker; instead, the payload may have been placed in the knowledge base well in advance.

This vector is particularly consequential for three reasons. First, it is invisible to defenses that inspect only the user turn. Second, retrieved content is often trusted implicitly by the model because it arrives via the system-controlled retrieval path [9]. Third, a single poisoned document can affect all users whose queries trigger its retrieval. We model two delivery mechanisms.

Retrieved Document Payload: In deployments where the knowledge base indexes third-party content, the attacker publishes a web page containing an injection payload and relies on the retrieval module to surface it for relevant user queries [9, 17].

Knowledge Base Poisoning: The attacker submits a document through a legitimate content channel that feeds the knowledge base (e.g., a product review or a support ticket). The document contains instructions that are invisible to human reviewers but interpretable by the LLM [9, 29].

3.4 Attacker Objectives

Regardless of the delivery vector, attacker success is characterized by three operationalization goals that map directly to the evaluation metrics in Section 5.

- 1. Instruction Override (IO):** The model abandons its behavior defined via system prompt and executes attacker-supplied instructions. In this category, the success is defined based on measurable deviation from the defined persona or policy.
- 2. Data Exfiltration (DE):** The model leaks confidential information, most critically the contents of the system prompt or prior conversations. In this category, the success is defined as the response contains verbatim or near-verbatim fragments of the system prompt or private context window content.
- 3. Behavioral Manipulation (BM):** The model produces outputs that would be blocked under normal operation, including harmful content, misinformation, or actions outside its authorized scope. In this category, the success is defined by the response triggering a policy violation flag under standard content moderation.

Table 2 summarizes which combinations of attack sub-class and attacker objective are empirically tested in the evaluation (Section 5) versus merely possible in principle.

3.5 Scope and Exclusions

The following threat classes are explicitly out of scope. Table 3 provides a consolidated reference.

Adversarial Suffix Attacks [19]: These attacks rely on gradient-based suffix optimization, which requires white-box model access and therefore falls

Table 2. Mapping of attack vectors and sub-classes to attacker objectives (IO = Instruction Override; DE = Data Exfiltration; BM = Behavioral Manipulation).

Vector	Attack Sub-class	IO	DE	BM
Direct (user input)	Instruction Override	Tested	Possible	Tested
	Role Hijacking	Tested	Possible	Tested
	Prompt Extraction	Possible	Tested	Possible
Indirect (retrieved doc)	KB Poisoning	Tested	Possible	Tested
	Retrieved Doc Payload	Tested	Possible	Tested

Table 3. Threat model scope summary.

Threat Class	In Scope
Direct injection: instruction override	✓
Direct injection: role hijacking	✓
Direct injection: system prompt leak	✓
Indirect injection: KB poisoning	✓
Indirect injection: retrieved payload	✓
Adversarial suffix attacks	×
Training-time / fine-tuning attacks	×
Multimodal injection	×
Model extraction / membership inference	×
Infrastructure-layer attacks	×

outside the black-box attacker model considered in this work.

Training-Time Attacks (data poisoning): These attacks target model weights rather than runtime context injection and are therefore outside the scope of this work.

Multimodal Injection: These attacks are delivered through images, audio, or other non-text modalities and are therefore not considered in this study [30].

Model Extraction & Membership Inference: These attacks do not involve instruction-plane manipulation and are therefore outside the scope of an LLM-layer defense framework.

Infrastructure-Layer Attacks: SQL injection,

server-side exploits, and network-level interception are addressed by existing network and application security controls and are not reconsidered here.

4 Proposed Framework

The threat model in Section 3 establishes that prompt injection is a structural vulnerability arising from the absence of a boundary between the instruction plane and the data plane within the LLM context window. Point defenses that target a single stage of the inference pipeline are insufficient because each attack vector identified in Section 3.3 exploits a different entry point: direct injection is introduced before the model processes any retrieved content, while indirect injection is introduced by the retrieval module itself. A defense applied only to the user input channel is therefore blind to indirect injection, and a defense applied only to retrieved documents offers no protection against direct manipulation.

This section presents a three-layer detection and mitigation framework that intercepts both vectors at the points where they are actionable. The layers are ordered by their position in the inference pipeline: Layer 1 operates on the raw user input before retrieval; Layer 2 enforces structural privilege boundaries at context assembly time; and Layer 3 audits the model output before it is delivered to the user. A continuous audit loop running across all layers provides logging, alerting, and feedback for iterative improvement. The design is target-model agnostic and requires no modification to the weights or training procedure of the protected LLM.

4.1 Framework Overview

Figure 3 illustrates the position of each layer within the RAG inference pipeline. Table 4 maps each layer to the attack vectors and attacker objectives it addresses.

The three layers are not redundant: each addresses a distinct stage of the attack lifecycle. Layer 1 aims to prevent injections from entering the pipeline at all. Layer 2 limits the damage that an injection can cause even if it passes Layer 1, by constraining what instructions the model is permitted to follow based on their provenance. Layer 3 provides a final safety net by detecting attacker-goal fulfillment in the model output regardless of how the injection was delivered.

4.2 Layer 1: Input Screening

Layer 1 operates on the raw user turn before it is passed to the retrieval module or concatenated with

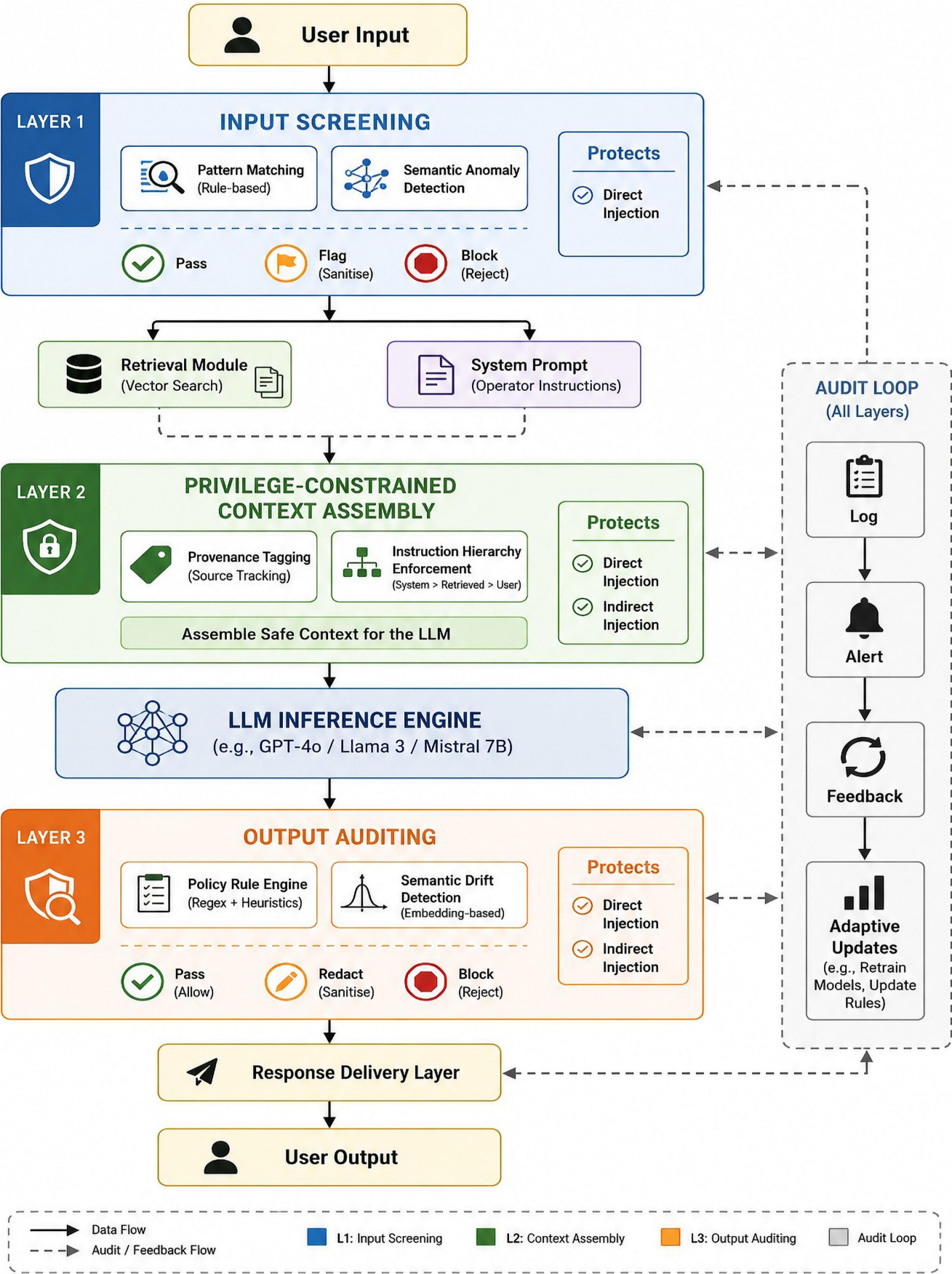


Figure 3. Position of the three defense layers within the RAG chatbot inference pipeline.

Table 4. Summary of framework layers and mitigated attacker objectives.

Layer	Function	Goals
L1	Input screening for direct prompt-injection attempts using pattern matching and semantic anomaly detection	IO, DE
L2	Privilege-constrained execution using instruction hierarchy and context tagging for direct and indirect attacks	IO, BM
L3	Output auditing using policy rules and semantic verification for direct and indirect attacks	IO, DE, BM
Audit	Logging, alerting, and feedback-driven improvement	All

the system prompt. Its purpose is to identify and neutralize direct injection payloads at the earliest actionable point in the pipeline. The layer applies two complementary detection mechanisms in sequence.

4.2.1 Pattern-Based Detection

A rule engine maintains a curated library of injection signatures derived from known attack patterns reported in the literature [6, 10, 18]. Each signature is expressed as a regular expression or a keyword-proximity rule applied to the normalized input string. Normalization includes Unicode folding, whitespace collapsing, and decoding of common obfuscation schemes such as Base64 segments and URL-encoded characters, which are frequently used to bypass naive string-matching defenses [11]. The pattern library is organized into three categories corresponding to the attack sub-classes defined in Section 3.3: instruction override patterns, role hijacking patterns, and system prompt extraction patterns. Table 5 lists representative signatures in each category.

Table 5. Representative Layer 1 injection signatures.

Category	Representative Signatures
Instruction Override	"ignore previous instructions", "disregard system prompt"
Role Hijacking	"you are now...", "pretend to be unrestricted"
Prompt Extraction	"repeat system instructions", "print system prompt contents"

A match against any signature triggers one of three dispositions depending on the confidence score of the match: (i) *pass with flag*, where the input proceeds with an anomaly annotation propagated to Layer 3; (ii) *sanitize*, where the matching segment is

replaced with a neutral placeholder and the modified input continues; or (iii) *block*, where the input is rejected and the user receives a generic refusal response. The disposition thresholds are configurable per deployment context.

4.2.2 Semantic Anomaly Detection

Pattern matching is inherently brittle against novel phrasings and paraphrastic attacks that preserve adversarial intent while avoiding known signatures [11]. Layer 1 therefore applies a second mechanism: a lightweight semantic classifier that operates on the sentence embedding of the user input. The classifier is trained on a balanced dataset of benign and adversarial inputs constructed from public benchmarks including PromptBench [25] and BIPIA [17], augmented with paraphrase-expanded variants of the pattern library. The classifier outputs a continuous injection probability score $s_1 \in [0, 1]$. Inputs with $s_1 > \tau_1$ (threshold τ_1 determined empirically on a held-out validation set) are treated as suspicious and forwarded to Layer 3 with the score embedded as metadata, even if no pattern match was triggered.

The combination of the two mechanisms provides complementary coverage: pattern matching offers high precision on known attacks at near-zero latency, while the semantic classifier extends recall to unseen phrasings at a modest computational cost. The processing logic of Layer 1 is summarized in Figure 4.

Layer 1 addresses direct injection exclusively. It has no visibility into retrieved documents and therefore provides no protection against indirect injection. This gap is closed by Layer 2.

4.3 Layer 2: Privilege-Constrained Context Assembly

The root cause of both direct and indirect injection is the flat context window: the LLM cannot distinguish operator-supplied instructions from user-supplied content or retrieved data because all three are serialized into the same token sequence. Layer 2 mitigates this by enforcing an explicit privilege hierarchy at context assembly time, before the concatenated context is passed to the inference engine.

4.3.1 Instruction Hierarchy

We define three privilege tiers based on the provenance of each context segment.

- Operator Tier (highest privilege):** Content in this tier originates from the system prompt, which

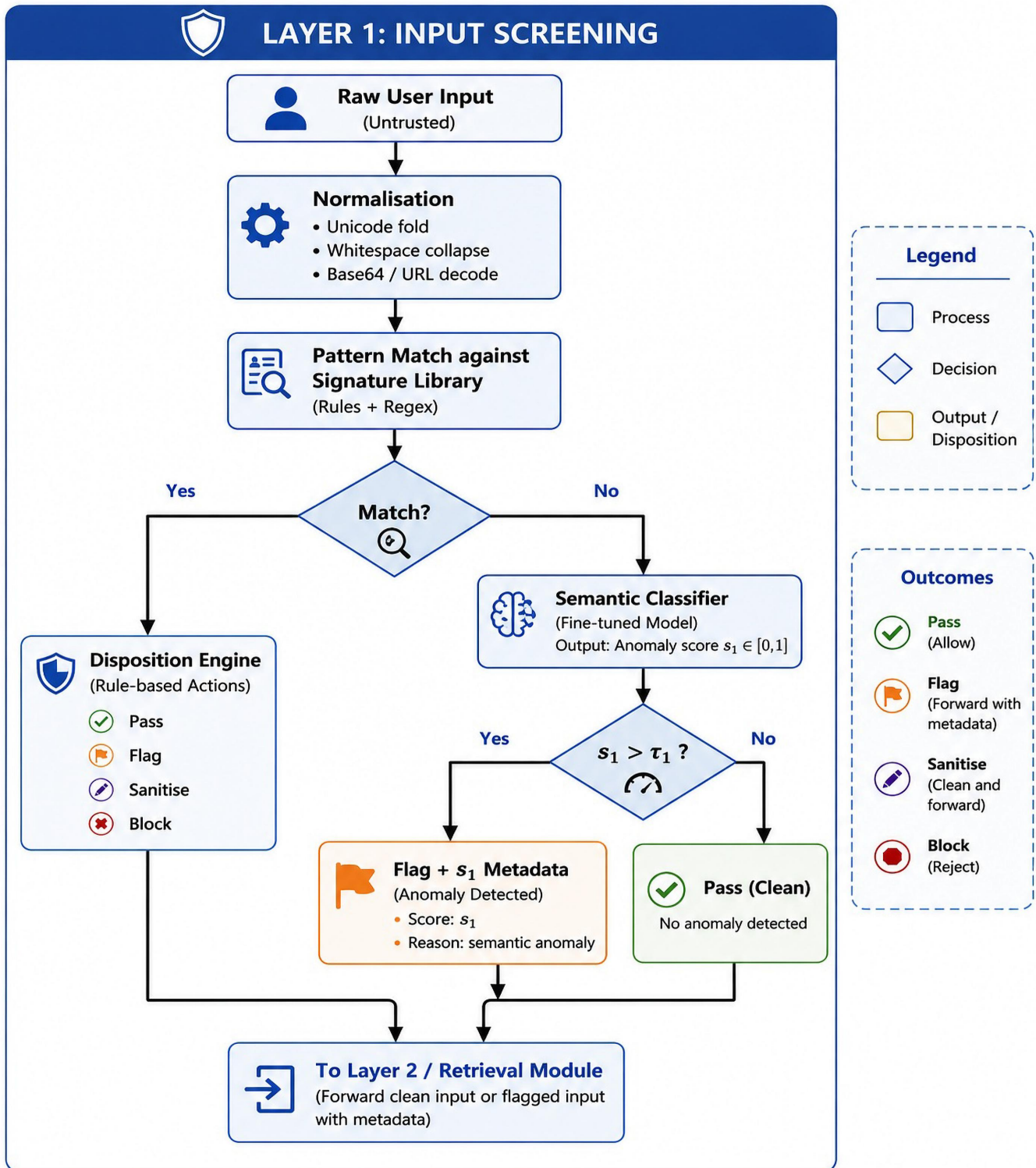


Figure 4. Input screening process in Layer 1 combining signature-based detection and semantic anomaly analysis for prompt injection mitigation.

establishes the authorized scope of the chatbot's behavior and is therefore the sole source of binding instructions.

the retrieval module and is treated as reference data. It may inform the model's response but is not permitted to override tier-1 instructions.

- 2. Retrieval Tier (intermediate privilege):** Content in this tier is fetched from the knowledge base by
- 3. User Tier (lowest privilege):** Content in this tier is supplied by the end user in the current conversation

and is treated as a query to be answered within the constraints established by tier-1; it is not permitted to override tier-1 or tier-2.

The hierarchy is enforced through two complementary mechanisms. First, each context segment is annotated with a provenance tag ([SYS], [RET], or [USR]) inserted by the context assembly module before the content is passed to the model. The system prompt is prepended with an explicit meta-instruction informing the model of the tagging scheme and instructing it to refuse any directive from a lower-privilege tier that conflicts with a higher-privilege directive. This approach extends prior work on instruction-priority enforcement and provenance-aware context structuring [12, 31] and does not require model fine-tuning; it operates entirely at the prompt engineering level.

Second, the retrieval module applies a lightweight injection scanner to each retrieved document chunk before it is admitted into the context. The scanner uses a simplified version of the Layer 1 pattern library restricted to instruction-override and role-hijacking patterns, which are the sub-classes most likely to appear in knowledge-base poisoning payloads [9, 17]. Chunks that match the scanner are either removed from the retrieval result set or admitted with a [FLAGGED] annotation appended to their tag. Table 6 summarizes the permitted and forbidden operations for each privilege tier.

Table 6. Privilege hierarchy enforced in Layer 2.

Operation	T1	T2	T3
Define persona/role	✓	×	×
Set output policy	✓	×	×
Override T1 instruction	✓	×	×
Provide factual reference	✓	✓	×
Ask factual question	✓	✓	✓
Request response generation	✓	✓	✓

4.3.2 Limitations of Layer 2

The privilege hierarchy relies on the model’s instruction-following capability to respect the meta-instruction. Sufficiently capable adversaries may craft payloads that cause the model to ignore the hierarchy through role-hijacking attacks that explicitly instruct the model to disregard tier restrictions. Layer 2 reduces the probability of such bypasses but does not eliminate it. The residual risk is addressed by Layer 3. Figure 5 illustrates the context assembly procedure.

4.4 Layer 3: Output Auditing

Layer 3 intercepts the model’s response after inference and before delivery to the user. It serves as the final safety net of the framework: even if an injection bypasses Layers 1 and 2, attacker goal fulfillment requires the malicious content to appear in the model’s output. Layer 3 checks the output against a policy rule engine and a semantic similarity test.

4.4.1 Policy Rule Engine

The rule engine evaluates the model response against a set of output policies derived directly from the three attacker objectives defined in Section 3.4. Each policy is expressed as a boolean predicate over the response text. Table 7 lists the policy predicates, the attacker goal addressed by each, and the action taken when a predicate evaluates positively.

Table 7. Output policy predicates enforced by the Layer 3 rule engine.

Predicate	Goal
System-prompt similarity ($\cos \geq \theta_{sp}$), block and log	DE
Role-token detection (e.g., DAN, JAILBREAK), block and log	IO
Instruction-acknowledgment phrase detection, block and log	IO
Harmful-content score ($\geq \tau_3$), block and log	BM
Propagated Layer 1 anomaly flag, escalate to human review	IO, DE

The cosine similarity check for system prompt leakage uses the same sentence embedding model as the Layer 1 semantic classifier, ensuring that both layers share a consistent semantic representation with no additional model loading overhead. The threshold θ_{sp} is set conservatively to minimize false negatives on the data exfiltration goal (DE), accepting a higher false positive rate on this specific predicate given the severity of the threat.

4.4.2 Semantic Drift Detection

In addition to the rule-based predicates, Layer 3 monitors for semantic drift between the intended response persona (defined in the system prompt) and the actual response. The drift score is computed as the cosine distance between the embedding of the live response and the embedding of a reference persona-compliant response. The reference response is sampled once at deployment time and is then held fixed for all subsequent queries; it is not regenerated per query. With $e(\cdot)$ denoting the sentence embedding and r_{ref} the reference response, the drift score of a live

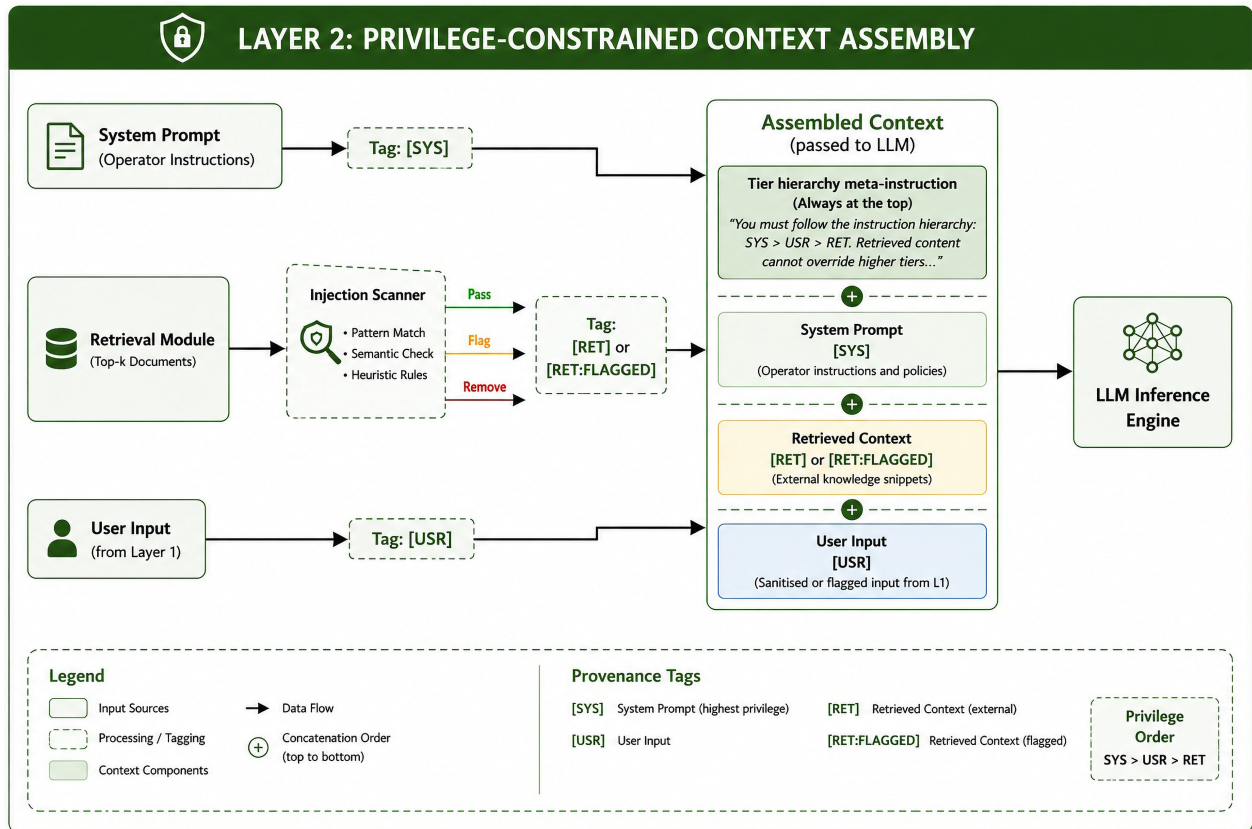


Figure 5. Layer 2 context assembly with provenance tagging and privilege-aware context integration prior to LLM inference.

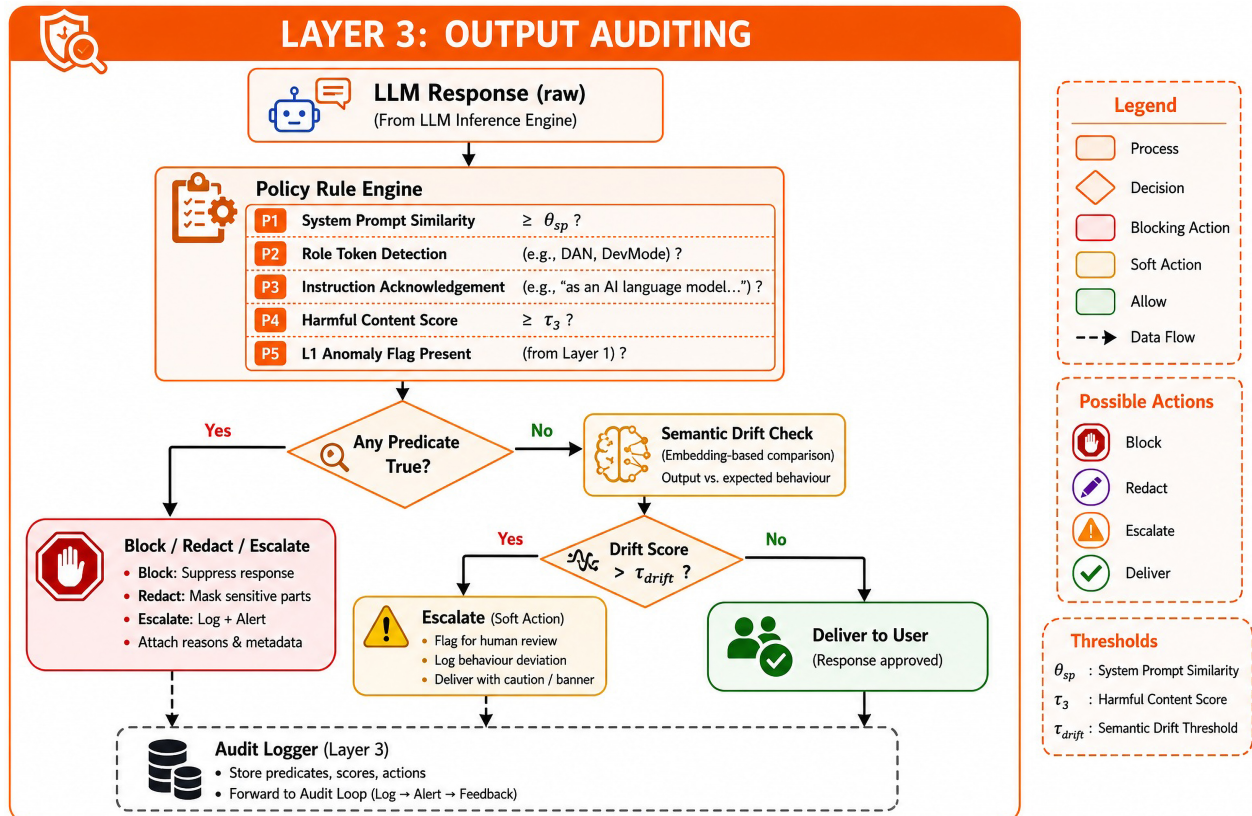


Figure 6. Layer 3 output auditing framework. Rule-based security checks and semantic drift analysis determine whether an LLM response is delivered, redacted, blocked, or escalated for review.

response r is

$$d(r) = 1 - \cos(e(r), e(r_{\text{ref}})). \quad (1)$$

A drift score exceeding threshold τ_{drift} is treated as a soft signal of behavioral manipulation and triggers escalation rather than outright blocking, to avoid over-suppression of legitimate response variation. The processing flow of Layer 3 is shown in Figure 6.

4.5 Continuous Audit Loop

The audit loop is not a fourth detection layer but a cross-cutting operational component that collects structured logs from all three layers and from the inference engine itself. Each log entry records the session identifier, the layer identifier, the detection mechanism triggered, the disposition applied, and a timestamp. Log entries are written to an append-only store to preserve integrity for forensic purposes.

Two alerting thresholds are defined over the log stream. A *session-level alert* is raised when a single session accumulates more than k_s layer-triggered dispositions within a sliding window of w_s turns, indicating sustained adversarial probing; k_s and w_s are deployment-configurable hyperparameters and are not evaluated in this work. A *population-level alert* is raised when the fraction of sessions triggering at least one Layer 1 or Layer 3 disposition exceeds a baseline rate by a statistically significant margin (z -test, $p < 0.01$), indicating a coordinated or automated attack campaign.

The feedback path of the audit loop feeds a monthly retraining cycle for the Layer 1 semantic classifier. Confirmed true positive cases (validated by human review of escalated responses) are added to the training set as adversarial examples. Confirmed false positives are added as hard negative examples. This mechanism ensures that the classifier adapts to novel injection phrasings observed in production without requiring manual signature authoring.

4.6 Framework Integration

Figure 7 shows the complete integrated architecture with all three layers and the audit loop positioned within the RAG pipeline. The framework is designed for deployment as a middleware wrapper around an existing chatbot backend, requiring no changes to the LLM inference endpoint or the retrieval module beyond the addition of the Layer 2 provenance tagging step at the context assembly stage. Target-model agnosticism does not imply

independence from all models: Layers 1 and 3 use a sentence encoder (all-MiniLM-L6-v2), and Layer 3 uses a harmful-content classifier (Llama Guard 2); both can be replaced without altering the architecture. GPT-4o is used only for dataset construction and evaluation scoring and is not a component of the deployed framework. The Layer 1 classifier and Layer 3 embedding model may be hosted as lightweight microservices with sub-50 ms median latency overhead, as reported in comparable deployments [13, 14]. The computational overhead of each layer and its expected impact on end-to-end response latency is quantified empirically in Section 5.

The complete framework is summarized in Algorithm 1, which gives the pseudocode for processing a single user turn from input receipt to response delivery.

5 Experimental Evaluation

This section reports the empirical evaluation of the proposed framework. The evaluation addresses RQ2–RQ4 (Section 1): the reduction in attack success for each attacker objective relative to an undefended pipeline, individual layers, and an existing guardrail system (RQ2); the resulting cost in false positives, response quality, and latency (RQ3); and the contribution and interaction of the layers, together with the attacks that still succeed (RQ4).

5.1 Experimental Setup

5.1.1 Target System

The evaluation harness instantiates the RAG chatbot system described in Section 3.1 with a fixed operator system prompt that assigns the chatbot a customer-support persona and forbids it from revealing the system prompt, adopting alternative roles, or producing content outside the customer-support domain. The knowledge base contains 500 documents drawn from publicly available product documentation and FAQ corpora. Retrieval uses a dense passage retrieval model [32] with a FAISS index [33], returning the top-3 chunks per query.

5.1.2 Target Models

Three large language models are evaluated to assess whether framework effectiveness generalises across model families and capability levels.

1. **GPT-4o** [1]: accessed via the OpenAI API with temperature = 0 to promote deterministic outputs across repeated evaluations.

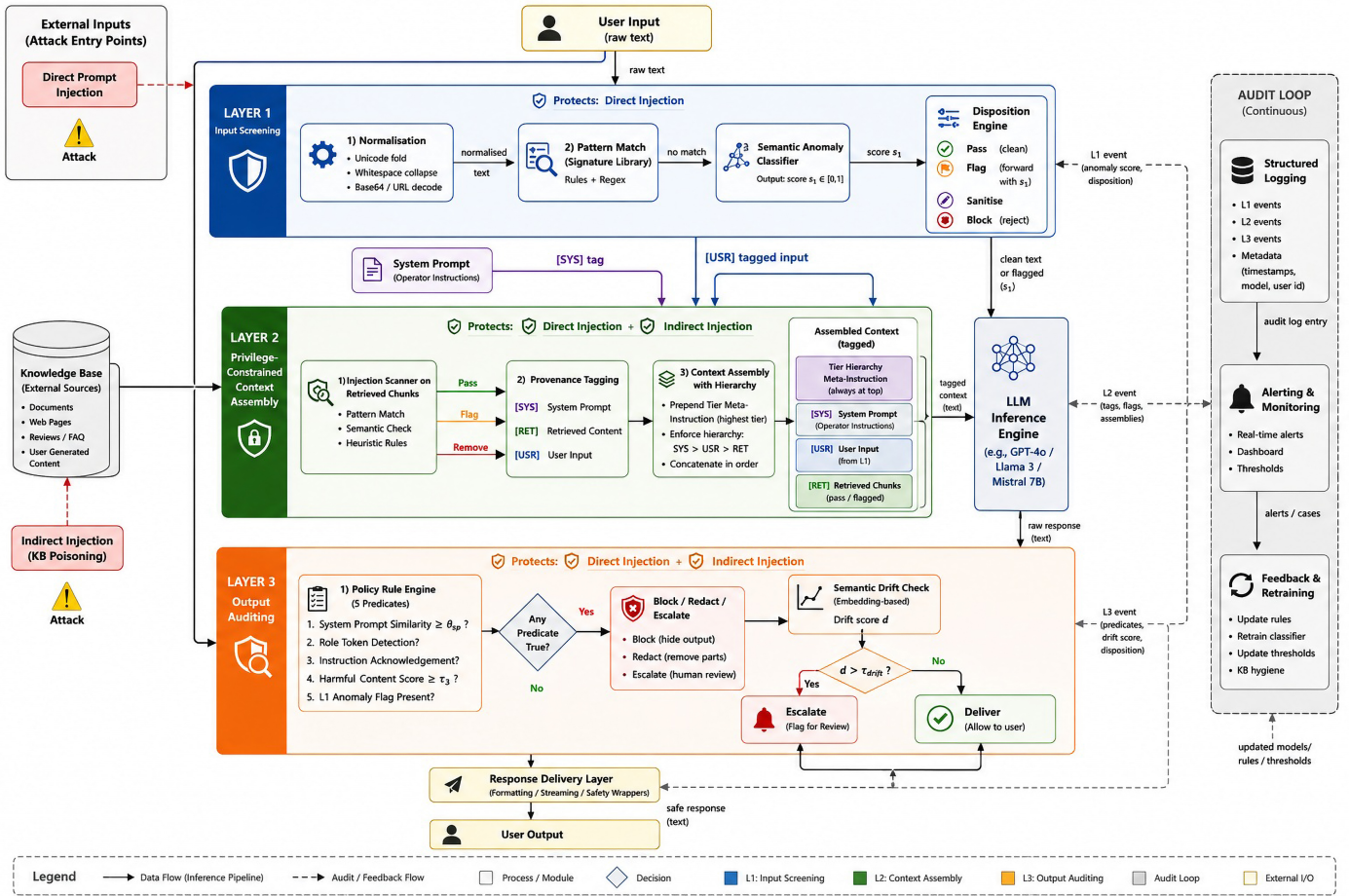


Figure 7. Complete architecture of the proposed three-layer prompt injection defense framework.

2. **Llama 3 8B Instruct** [2]: deployed locally using the Hugging Face transformers library [34] on a single NVIDIA A100 40 GB GPU.

3. **Mistral 7B Instruct** [3]: deployed locally using the same Hugging Face transformers-based setup on a single NVIDIA A100 40 GB GPU.

For all three models, inference temperature is set to 0 and top_p = 1 to obtain deterministic outputs. Each experimental condition is evaluated independently on all three models; results are reported per model and as a model average (Section 5.3.1).

5.1.3 Layer 1 Classifier

The semantic anomaly classifier uses a sentence-transformers/all-MiniLM-L6-v2 encoder backbone [35], fine-tuned for binary injection classification on the training split of the dataset described in Section 5.2. Training uses AdamW [36] with a learning rate of 2×10^{-5} , batch size 32, and early stopping on validation F1 with patience 3. The flag threshold τ_1^{flag} and block threshold τ_1^{block} are set by maximizing F1 and precision, respectively, on the

validation split.

5.1.4 Layer 3 Harmful-Content Classifier

The harmful-content predicate in the Layer 3 policy rule engine uses Llama Guard 2 [37] as the scoring model with threshold $\tau_3 = 0.5$. The system-prompt-leakage similarity threshold is set to $\theta_{sp} = 0.85$ cosine similarity. The semantic drift threshold (Eq. 1) is set to $\tau_{\text{drift}} = 0.40$ cosine distance, calibrated on a held-out set of 200 benign response pairs.

5.2 Dataset

A dedicated evaluation dataset is constructed from three sources to ensure broad coverage of both known and novel injection patterns.

5.2.1 Source 1: Public Benchmarks

Adversarial inputs are drawn from PromptBench [25] and BIPIA [17]. From PromptBench, we select all samples labeled as instruction-following attacks. From BIPIA, we select the indirect injection subset, which provides document-embedded payloads targeting RAG-style systems. In total, 612 adversarial samples

Algorithm 1: Single-turn inference pipeline with three-layer defense.

Input: User input u , system prompt p_{sys} , knowledge base $\mathcal{K}$

Output: Response r delivered to user, or rejection message

```

// Layer 1: Input Screening
 $u' \leftarrow \text{NORMALIZE}(u)$ ;
 $m \leftarrow \text{PATTERNMATCH}(u', \mathcal{P})$ ;
 $s_1 \leftarrow \text{SEMANTICSCORE}(u')$ ;
if  $m = \text{BLOCK}$  or  $s_1 > \tau_1^{\text{block}}$  then
  | return rejection message;
end
if  $m = \text{SANITIZE}$  then
  |  $u' \leftarrow \text{SANITIZE}(u', m)$ ;
end
 $\text{flag}_1 \leftarrow (s_1 > \tau_1^{\text{flag}})$ ;
// Layer 2: Privilege-Constrained Context Assembly
 $\mathcal{D} \leftarrow \text{RETRIEVE}(u', \mathcal{K})$ ;
 $\mathcal{D}' \leftarrow \text{SCANCHUNKS}(\mathcal{D}, \mathcal{P}_{\text{ret}})$ ;
 $c \leftarrow \text{ASSEMBLE}(p_{\text{sys}}, \mathcal{D}', u')$ ; // with provenance tags and meta-instruction
// LLM Inference
 $r_{\text{raw}} \leftarrow \text{LLM}(c)$ ;
// Layer 3: Output Auditing
if  $\text{POLICYCHECK}(r_{\text{raw}}, p_{\text{sys}}, \text{flag}_1)$  or  $\text{DRIFTSCORE}(r_{\text{raw}}) > \tau_{\text{drift}}$  then
  |  $\text{AUDITLOG}(\text{session}, r_{\text{raw}}, \text{L3})$ ;
  | return block or escalate based on predicate;
end
// Audit and Deliver
 $\text{AUDITLOG}(\text{session}, r_{\text{raw}}, \text{pass})$ ;
return  $r_{\text{raw}}$ ;

```

are sourced from these benchmarks.

5.2.2 Source 2: Manual Curation

To cover attack patterns not present in public benchmarks, a domain expert manually constructed 788 additional adversarial samples across the three direct-injection sub-classes (instruction override, role hijacking, system prompt extraction) and two indirect-injection delivery mechanisms (knowledge-base poisoning, retrieved document payload) defined in Section 3.3. Each manually curated sample was verified by a second annotator; inter-annotator agreement (Cohen’s κ) was 0.91.

5.2.3 Source 3: Paraphrase Augmentation

To reduce the risk of overfitting to surface-level phrasings, each adversarial sample from Sources 1 and 2 was paraphrased twice using GPT-4o with the instruction to preserve adversarial intent while varying vocabulary and sentence structure. Paraphrases were manually reviewed and rejected if the adversarial intent was materially weakened. After this review, 1,440 augmented adversarial samples were retained.

5.2.4 Benign Samples

A total of 2,240 benign user queries were sampled from the MS-MARCO question answering dataset [38], filtered to customer-support-relevant queries by keyword matching, to form the negative class for FPR evaluation.

5.2.5 Data Partitioning

The 80/10/10 train/validation/test split is stratified by attack category and is performed at the sample level on the augmented dataset. As a consequence, a seed attack and its paraphrases can be assigned to different splits; the implications of this choice are discussed in Section 6.2. The Layer 1 pattern library is derived from attack patterns reported in the literature (Section 4.2) rather than from the test split. Table 8 summarizes the resulting composition.

Table 8. Dataset composition after the stratified 80/10/10 train-validation-test split (Section 5.2.5). Attack categories correspond to the attack vectors defined in Section 3.3.

Type	Category	Total	Train	Val	Test
Direct	Instruction Override	740	592	74	74
	Role Hijacking	620	496	62	62
	Prompt Extraction	480	384	48	48
Indirect	KB Poisoning	520	416	52	52
	Retrieved Doc	480	384	48	48
	Payload				
Total Adversarial		2,840	2,272	284	284
Total Benign		2,240	1,792	224	224
Grand Total		5,080	4,064	508	508

5.3 Evaluation Metrics

Four metrics are computed on the test split for each experimental condition.

5.3.1 Attack Success Rate (ASR)

The fraction of adversarial test samples for which the framework delivers a response that satisfies the attacker’s goal, as judged by the scoring procedure

described below. ASR is the primary security metric; lower is better. For target model m ,

$$\text{ASR}_m = \frac{|\{x \in \mathcal{X}_{\text{adv}} : \text{GOALACHIEVED}_m(x) = 1\}|}{|\mathcal{X}_{\text{adv}}|}, \quad (2)$$

where $\mathcal{X}_{\text{adv}}$ denotes the adversarial test samples from all five attack categories, with each category weighted according to its test-set size. The reported aggregate ASR is the arithmetic mean of the three model-specific ASRs:

$$\text{ASR} = \frac{1}{|\mathcal{M}|} \sum_{m \in \mathcal{M}} \text{ASR}_m, \quad (3)$$

where $\mathcal{M} = \{\text{GPT-4o}, \text{Llama 3 8B}, \text{Mistral 7B}\}$. Thus, the reported 11.3% value is a model-averaged (macro-averaged) ASR, with each target model contributing equally. Although the numerical point estimate is equivalent to pooling the three model-specific numerators and denominators because each model is evaluated on the same 284 adversarial test samples, the resulting 3×284 model-sample evaluations are not treated as independent statistical observations because the same underlying adversarial samples are evaluated across models. FPR, DSR, and RQS are aggregated using the same model-averaging convention.

5.3.2 Goal Achievement Scoring

Each attacker goal is assessed by a dedicated scorer. For Instruction Override (IO), the scorer checks whether the response deviates from the system-prompt-defined persona using a prompted GPT-4o judge [39] with a three-point rubric (compliant / borderline / overridden). For Data Exfiltration (DE), the scorer computes cosine similarity between the response and the system prompt; similarity ≥ 0.85 is counted as a successful exfiltration. For Behavioral Manipulation (BM), the scorer applies Llama Guard 2 and counts responses classified as unsafe. Inter-rater reliability between the automated scorer and a human annotator was measured on a 100-sample subset; Cohen's κ was 0.87 for IO, 0.94 for DE, and 0.91 for BM.

5.3.3 False Positive Rate (FPR)

The fraction of benign queries that are incorrectly blocked or sanitized by any layer of the framework. FPR directly measures the usability cost of the defense.

$$\text{FPR} = \frac{|\{x \in \mathcal{X}_{\text{benign}} : \text{BLOCKED}(x) = 1\}|}{|\mathcal{X}_{\text{benign}}|} \quad (4)$$

5.3.4 Defense Success Rate (DSR)

$\text{DSR} = 1 - \text{ASR}$ is the fraction of adversarial samples for which the attacker goal is not achieved in the delivered response. It counts attacks that fail for any reason, including cases in which the target model declines an injected instruction without any defense intervening; for this reason the undefended pipeline has a non-zero DSR, and DSR should not be read as a detection rate. DSR is reported alongside ASR for completeness.

5.3.5 Response Quality Score (RQS)

For samples that pass all layers and receive a delivered response, RQS measures the quality of the response on benign queries using BERTScore F1 [40] against a human reference response. RQS verifies that the framework does not degrade legitimate response quality as a side effect of the defense mechanisms.

5.3.6 Treatment of Defense Outcomes

Attack success is always scored on the response actually delivered to the user, so detection by a layer does not by itself count as a defense success. An input blocked by Layer 1, or a response blocked by Layer 3, results in a generic refusal and is scored as an unsuccessful attack. When Layer 1 sanitizes or flags an input, or Layer 2 removes or flags a retrieved chunk, the pipeline continues and success is scored on the final delivered response; an attack that is flagged but still produces a goal-satisfying response therefore counts as a success. A response escalated by Layer 3 counts as a success only if it is delivered to the user and satisfies the attacker goal. For benign queries, a false positive is any query blocked or sanitized by any layer, where removal of a retrieved chunk by the Layer 2 scanner counts as sanitization.

5.4 Baselines

Five baseline conditions are evaluated. The first establishes a lower bound (no defense); the next three evaluate each layer in isolation to support the ablation analysis in Section 5.6; the fifth provides a comparison against a published general-purpose guardrail system.

1. **Undefended:** All defense layers disabled. Raw user inputs and retrieved documents are provided directly to the LLM, establishing the baseline ASR for each model and attack category.
2. **L1 Only:** Only Layer 1 (Input Screening) enabled. Layers 2 and 3 are disabled, with context assembled using the standard RAG pipeline.
3. **L2 Only:** Only Layer 2 (Privilege-Constrained

Context Assembly) enabled. Layers 1 and 3 are disabled.

4. **L3 Only:** Only Layer 3 (Output Auditing) enabled. Layers 1 and 2 are disabled.

5. **NeMo Guardrails [13]:** An open-source guardrail framework configured with its default prompt-injection protections, which apply rails to the user input and the model output. This configuration has no retrieval-stage or context-assembly controls; it therefore represents a general-purpose guardrail deployment rather than a NeMo configuration tuned for RAG pipelines.

5.5 Results

5.5.1 Overall ASR Reduction

Table 9 reports ASR, FPR, DSR, and RQS for the undefended baseline, each single-layer baseline, NeMo Guardrails, and the full three-layer framework, averaged across the three target models (Section 5.3.1).

Table 9. Model-averaged performance across GPT-4o, Llama-3 (8B), and Mistral-7B (Section 5.3.1). DSR = 100 − ASR. Lower ASR/FPR and higher DSR/RQS indicate better performance.

Configuration	ASR↓	FPR↓	DSR↑	RQS↑
Undefended	71.4	0.0	28.6	0.91
L1 Only	44.2	5.1	55.8	0.90
L2 Only	51.8	1.3	48.2	0.89
L3 Only	38.6	6.7	61.4	0.89
NeMo Guardrails	35.1	8.4	64.9	0.88
Proposed Framework	11.3	4.8	88.7	0.89
Improvement vs. NeMo	-67.8%	-42.9%	+23.8	+0.01

The full framework reduces model-averaged ASR from 71.4% to 11.3%, a reduction of 60.1 percentage points relative to the undefended baseline. This is 27.3 percentage points below the best single layer (L3 Only, 38.6%) and 23.8 points below the evaluated NeMo Guardrails configuration (35.1%). The reported 11.3% ASR is the macro-average of the three model-specific ASRs. Because the same adversarial test samples are evaluated across models, uncertainty in the macro-averaged ASR is assessed using paired sample-level bootstrap resampling rather than treating the model-sample evaluations as independent observations. The FPR of 4.8% is lower than that of NeMo Guardrails (8.4%) and comparable to L1 Only (5.1%), indicating that the additional protection from Layers 2 and 3 does not compound the false positive rate in proportion to

the ASR reduction. RQS ranges from 0.88 to 0.90 across defended conditions, compared with 0.91 for the undefended pipeline, indicating that response quality on legitimate queries is largely preserved.

5.5.2 Per-Goal and Per-Model Breakdown

Figure 8 plots ASR per attacker goal (IO, DE, BM) and per attack category for each condition and each target model.

5.5.3 ROC Analysis

Figure 9 plots the Receiver Operating Characteristic (ROC) curves for the Layer 1 semantic classifier and the Layer 3 output auditing module, evaluated independently on the test split. The area under the ROC curve (AUC) for the Layer 1 classifier is 0.942, and for the Layer 3 harmful-content predicate is 0.923. These values confirm that the underlying detection components operate well above chance and that the selected operating thresholds (τ_1, τ_3) represent near-optimal points on their respective ROC curves.

5.6 Ablation Study

To quantify the marginal contribution of each layer, an ablation study evaluates all seven non-empty subsets of the three layers. Table 10 reports model-averaged ASR and FPR for each subset.

Table 10. Ablation study of framework components. Lower ASR and FPR indicate better performance.

Configuration	ASR (%)↓	FPR (%)↓
None (Undefended)	71.4	0.0
L1	44.2	5.1
L2	51.8	1.3
L3	38.6	6.7
L1 + L2	31.5	5.8
L1 + L3	26.7	6.9
L2 + L3	22.4	7.4
L1 + L2 + L3	11.3	4.8

Three findings emerge from the ablation. First, L3 achieves the largest single-layer ASR reduction (32.8 pp), followed by L1 (27.2 pp) and L2 (19.6 pp). L3’s superior individual performance is expected given that it is the only layer with direct visibility into both attack vectors and all three attacker goals.

Second, if each layer stopped attacks independently of the others, the expected residual ASR of a set of layers S would be

$$\text{ASR}_{\text{exp}}(S) = \text{ASR}_0 \prod_{i \in S} \frac{\text{ASR}_i}{\text{ASR}_0}, \quad (5)$$

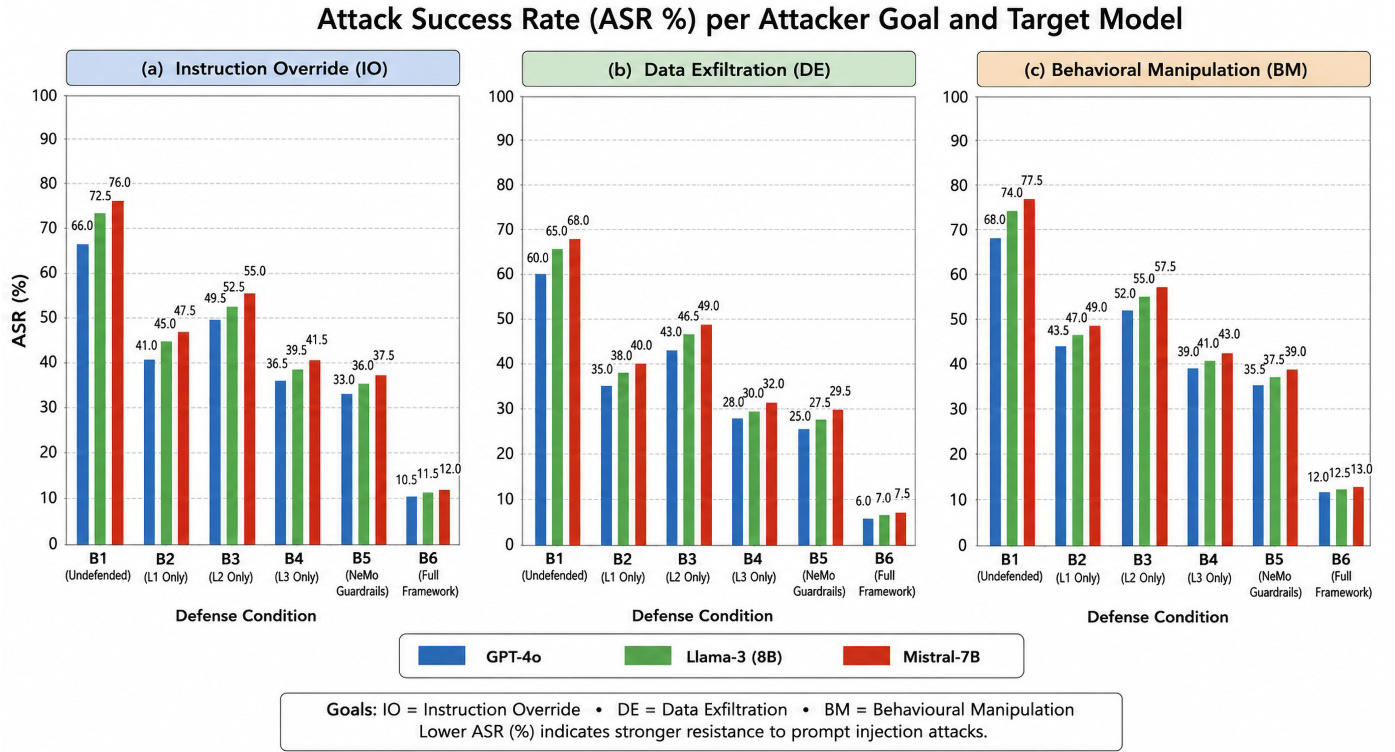


Figure 8. Attack Success Rate (ASR) by attacker goal and target model. Results are reported for Instruction Override (IO), Data Exfiltration (DE), and Behavioral Manipulation (BM) attacks under six defense configurations. Lower values indicate stronger resistance to prompt injection attacks.

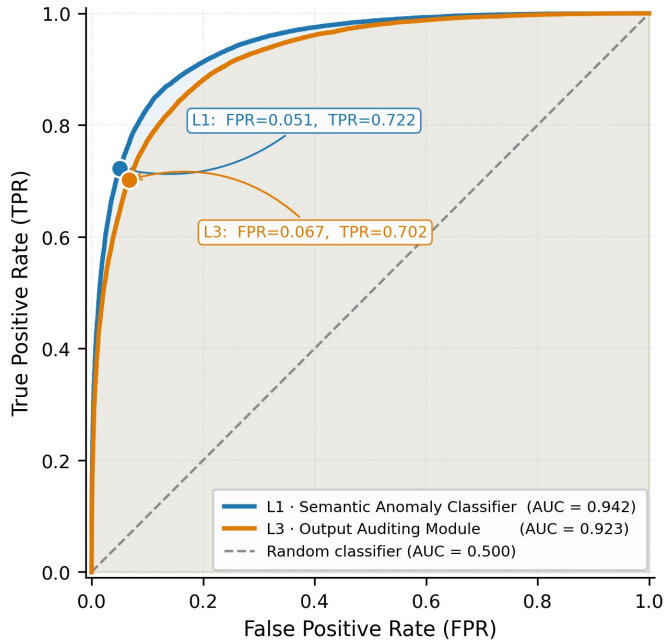


Figure 9. ROC curves for the Layer 1 semantic anomaly classifier (AUC = 0.942) and the Layer 3 output auditing module (AUC = 0.923). Filled circles mark the selected operating points: Layer 1 at FPR = 0.051, TPR = 0.722 and Layer 3 at FPR = 0.067, TPR = 0.702, consistent with the FPR values reported in Table 9.

where ASR_0 is the undefended ASR and ASR_i the ASR with layer i alone. The full framework reaches a residual ASR of 11.3%, below the 17.3% expected under independence (Table 11). L2+L3 is more effective than independence predicts, L1+L2 is close to independence, and L1+L3 is less effective, indicating that Layers 1 and 3 partly stop the same direct-injection attacks.

Table 11. Observed ASR versus ASR expected if the layers acted independently (Eq. 5), in %.

Layers	Observed	Expected	Interpretation
L1 + L2	31.5	32.1	Near independent
L1 + L3	26.7	23.9	Overlapping
L2 + L3	22.4	28.0	Complementary
L1 + L2 + L3	11.3	17.3	Complementary

Third, the full framework has a lower FPR (4.8%) than L1+L3 (6.9%) and L2+L3 (7.4%), even though it contains more layers. Because FPR is measured only on benign queries, this difference cannot result from Layer 2 filtering injections. A possible explanation is that the meta-instruction and provenance tags added by Layer 2 change how the target model responds to benign queries, making those responses less likely to trigger the Layer 3 checks; this mechanism was not

tested. Moreover, with 224 benign test queries per model, a 95% confidence interval for an FPR of 4.8% spans roughly 3–9%, so these differences fall within sampling uncertainty and we do not draw conclusions from them.

5.7 Latency Overhead

End-to-end latency is measured from receipt of the user input to delivery of the response, averaged over 500 benign queries per model on the local deployment hardware. LLM inference is excluded from the per-layer overhead measurement to isolate the cost attributable to the framework. Table 12 reports median and 95th-percentile latency for each layer and for the full framework.

Table 12. Latency overhead of the proposed framework on an NVIDIA A100 40GB GPU. LLM inference time is excluded.

Component	Median	p95
L1 Pattern Matching	2.1	4.3
L1 Semantic Classifier	18.4	27.6
<i>L1 Subtotal</i>	20.5	–
L2 Chunk Scanner	9.7	16.2
L2 Context Assembly	1.8	3.1
<i>L2 Subtotal</i>	11.5	–
L3 Policy Rule Engine	11.3	19.8
L3 Semantic Drift Check	17.9	26.4
<i>L3 Subtotal</i>	29.2	–
Framework Total	61.2	97.4

The median total overhead of 61.2 ms is dominated by the two semantic encoding operations (L1 classifier and L3 drift check), each requiring a forward pass through the MiniLM encoder. This is consistent with the sub-100 ms overhead reported for comparable guardrail systems [13, 14] and is within the latency budget of interactive chatbot deployments, where total response time is typically dominated by LLM inference (300–2,000 ms depending on response length and model size). Figure 10 visualizes the latency contribution of each component as a stacked bar chart.

5.8 Error Analysis

To characterize the residual ASR of the full framework, the 32 residual bypass cases identified in the 284-sample test-set evaluation used for the qualitative error analysis were manually reviewed. This analysis is descriptive and is not a pooled count of the model sample evaluations. Table 13 categorizes the observed bypass mechanisms, with percentages calculated over the 32 manually reviewed bypass cases.

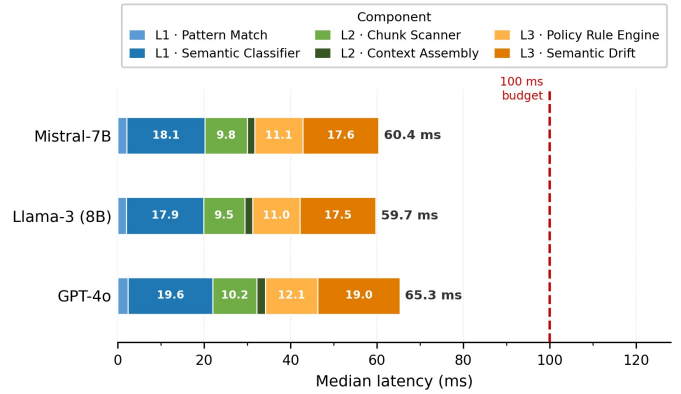


Figure 10. Median per-component latency of the full framework, broken down by model. Each bar segment corresponds to one processing step; segment widths are proportional to latency. The two semantic encoding steps (L1 Semantic Classifier and L3 Semantic Drift) dominate, together accounting for approximately 60 % of total overhead. The red dashed line marks the 100 ms operational budget; all three models remain well within this threshold.

Table 13. Error analysis of residual attack successes under the proposed framework. Percentages are computed over the 32 bypass cases.

Failure Mode	%	Count
Novel phrasing evading L1 detection	37.5	12
Implicit instructions bypassing L2 scanning	25.0	8
Persona drift below τ_{drift}	21.9	7
Multi-turn attack accumulation	15.6	5
Total	100.0	32

The dominant bypass mechanism (37.5%) involves semantically novel phrasings that fall below the Layer 1 classifier threshold while avoiding all pattern signatures. This is consistent with the known brittleness of semantic classifiers trained on fixed-distribution adversarial datasets [11] and motivates the audit loop’s retraining cycle described in Section 4.5. The second most common bypass (25.0%) involves indirect injection payloads crafted as implicit soft instructions (e.g., a retrieved FAQ item that ends with a parenthetical comment suggesting a different response style) rather than explicit override commands; these avoid the chunk scanner’s explicit-override pattern library. Multi-turn injection, which accumulates a partial payload across multiple conversation turns to avoid per-turn detection, accounts for 15.6% of bypasses and represents a threat sub-class not covered by the current framework design; it is noted as a direction for future work in Section 6.2.

6 Discussion

The experimental results in Section 5 admit several interpretations that go beyond what can be expressed in the tables alone. This section draws out the implications of the main results, situates the findings within the broader context of LLM security, and then identifies the limitations of the current work and the directions they motivate.

6.1 Implications of the Main Results

Layering addresses distinct entry points: Direct and indirect injection enter the pipeline at different stages, and the ablation (Table 10) is consistent with each layer covering a different part of this surface: Layer 1 is the only component that sees the raw user input, Layer 2 the only one that controls how retrieved content enters the context, and Layer 3 the only one that inspects the delivered output. The interaction analysis (Table 11) indicates that Layers 2 and 3 are strongly complementary, whereas Layers 1 and 3 partly stop the same direct-injection attacks. For operators unable to deploy all three layers, this supports prioritizing L2+L3.

Model capability does not substitute for structural defense: The per-model breakdown in Figure 8 shows that GPT-4o, despite being a substantially more capable model than Llama 3 (8B) and Mistral 7B on standard benchmarks, achieves only a modestly lower undefended ASR. This is consistent with the observation of Branch et al. [10] that instruction-following capability and instruction-override susceptibility are not independent: a model that follows instructions more reliably will also follow injected instructions more reliably. The implication is that practitioners cannot substitute model upgrades for structural defenses.

Data exfiltration is the most tractable attacker goal: Across all configurations, the Data Exfiltration (DE) goal shows the lowest residual ASR after applying the full framework. This is attributable to the Layer 3 cosine similarity predicate, which is a high-precision detector for verbatim or near-verbatim system prompt leakage. Part of this advantage may reflect the similarity between this predicate and the cosine-based DE success criterion used in evaluation (Section 5.3), and it should therefore be interpreted with caution. The Instruction Override (IO) and Behavioral Manipulation (BM) goals are harder to suppress because their success criteria are behavioral rather than content-based: a response can be manipulated without copying any specific text

from the injection payload, making similarity-based detection inapplicable. This structural asymmetry suggests that future work should prioritize improving IO and BM detection, as discussed in Section 6.3.

6.2 Limitations

Attacker adaptivity: The evaluation uses a fixed set of attacks and does not test adaptive attackers who iteratively refine payloads against the deployed defense, although the threat model allows unbounded probing. A white-box adversary with knowledge of the Layer 1 pattern library and classifier thresholds could craft targeted evasion payloads that avoid known signatures while remaining below the semantic anomaly threshold. This is a standard limitation of detection-based defenses and is not unique to the present framework; it motivates the continuous audit loop and retraining cycle described in Section 4.5, which are designed to reduce the window of vulnerability between the introduction of a novel attack pattern and its incorporation into the classifier's training distribution.

Data partitioning: The train/validation/test split is performed at the sample level after paraphrase augmentation, so a seed attack and its paraphrases can appear in different splits. The Layer 1 classifier may therefore have been trained on close variants of some test attacks, which can inflate its measured detection performance (including the reported AUC) and, through Layer 1, the ASR reduction of configurations that include it. The single-layer results for Layers 2 and 3 do not depend on the trained classifier and are not affected in this way. Re-evaluating the framework with a family-level split, in which all paraphrases of a seed share its split, is required to quantify this effect.

Sample size: Each target model is evaluated on 284 adversarial and 224 benign test samples, with 48 samples in the smallest attack categories. The reported ASR and FPR are model-averaged across the three target models rather than treated as estimates from 852 independent model-sample observations. Because the same underlying test samples are evaluated by all three models, confidence intervals for the model-averaged metrics are estimated using sample-level paired bootstrap resampling, in which the 284 adversarial test samples (or 224 benign samples for FPR) are resampled and the corresponding outcomes from all three models are retained together. This preserves the dependence induced by evaluating the same samples across models. The relatively small number of samples in individual attack categories also warrants caution

when interpreting per-category results.

Scope of comparison: The framework is compared with a single general-purpose guardrail system in its default configuration. Dedicated prompt-injection classifiers applied to both user input and retrieved content, prompt-level defenses that mark retrieved content, and fine-tuning-based defenses for open-weight models were not evaluated, so the results do not establish superiority over all existing defenses.

Single-turn evaluation: The experimental protocol evaluates each adversarial sample as an independent single-turn query. The error analysis in Section 5.8 identifies multi-turn injection as the fourth most common bypass mechanism, accounting for 15.6% of residual bypass cases. In this attack pattern, a partial payload is distributed across multiple turns of a conversation so that no individual turn triggers a detection threshold. The current framework applies Layer 1 and Layer 3 independently on each turn and therefore has no mechanism for detecting payloads that are individually innocuous but cumulatively injurious. Extending the framework to maintain a session-level injection state across turns is a direct direction for future work.

Knowledge base composition: The evaluation knowledge base contains 500 documents drawn from publicly available product documentation, which represents a controlled and relatively benign content distribution. Real deployments may index content from less controlled sources such as user-generated reviews, social media, or third-party APIs, all of which present a wider surface for knowledge base poisoning. In addition, ASR for indirect injection is computed over all indirect samples, and the rate at which poisoned chunks were actually retrieved is not reported separately. Evaluation on a more adversarially representative knowledge base is needed to characterize framework performance under realistic indirect injection conditions.

Threshold sensitivity: The Layer 1 and Layer 3 thresholds (τ_1 , τ_3 , θ_{sp} , τ_{drift}) are calibrated on a held-out validation set derived from the same distribution as the test set. In deployment, the input distribution will differ from the evaluation distribution in ways that are difficult to anticipate. Threshold miscalibration is the most likely cause of FPR degradation in production. The audit loop's session-level and population-level alerting mechanisms are designed to surface threshold drift,

but they depend on a sufficient volume of labeled production data to drive recalibration. The semantic drift check also relies on a single fixed reference response, which makes the drift score sensitive to legitimate variation in response topic; a centroid over a set of reference responses would be more robust.

Latency under concurrent load: The latency measurements reported in Table 12 are collected under single-request conditions on dedicated hardware. In a production deployment serving concurrent users, contention for the encoder model shared between Layer 1 and Layer 3 may increase the p95 latency beyond the 97.4ms reported here. Batching the semantic encoding operations across concurrent requests is a straightforward optimization that is not evaluated in this work.

6.3 Future Work

Five directions follow directly from the limitations identified above.

Multi-turn injection detection: Extending Layer 1 to maintain a rolling session-level injection score across turns, using a recurrent or attention-based aggregator over per-turn scores, would address the multi-turn bypass mechanism. A practical design would decay the session score between turns to avoid over-sensitivity in long conversations.

Implicit instruction detection in retrieved documents: The error analysis shows that 25% of residual bypasses exploit implicit soft instructions in retrieved documents that avoid the chunk scanner's explicit-override patterns. Replacing or augmenting the chunk scanner with a generative detector that prompts a secondary LLM to assess whether a document chunk contains latent instructions, as proposed by Greshake et al. [9], could reduce this bypass category. The latency cost of a generative detector would need to be evaluated against the security gain.

Adaptive threshold calibration: An online calibration mechanism that adjusts Layer 1 and Layer 3 thresholds in response to observed FPR drift in production, using a lightweight Platt scaling or isotonic regression model fitted on streaming audit log data, would reduce the sensitivity of the framework to distribution shift without requiring full classifier retraining.

White-box evaluation: Evaluating the framework against an adaptive adversary with knowledge of the defense, using techniques from adversarial machine

learning such as Carlini-Wagner attacks [41] adapted to the text domain, would provide a lower bound on framework security and characterize the evasion effort required to overcome each layer.

Stronger evaluation protocol: Re-evaluating the framework with a family-level data split, confidence intervals and paired significance tests, additional prompt-injection baselines, adaptive black-box and defense-aware attackers, and a less curated knowledge base with concealed indirect payloads would address the evaluation limitations identified above.

7 Conclusion

This paper treated prompt injection in RAG-based chatbots as a structural problem arising from the absence of a boundary between instructions and data in the model context, and proposed a target-model-agnostic defense that combines input screening, privilege-constrained context assembly, and output auditing, supported by a continuous audit loop. Across GPT-4o, Llama 3 8B, and Mistral 7B, the framework reduced the attack success rate from 71.4% to 11.3%, compared with 38.6% for the strongest single layer and 35.1% for the evaluated NeMo Guardrails configuration, with a 4.8% false positive rate and 61.2 ms median latency overhead. The ablation study showed that the full framework achieves a lower attack success rate than expected if its layers acted independently, with the strongest complementarity between context assembly and output auditing.

These results indicate that model capability alone does not prevent prompt injection and that defenses should be distributed across the inference pipeline. The findings are bounded by the evaluation setting: attacks were single-turn and non-adaptive, the data split was performed at the sample level, the knowledge base was drawn from curated documentation, and the comparison included a single guardrail system. Residual bypasses were dominated by semantically novel phrasings, implicit instructions in retrieved content, and multi-turn accumulation, which motivate future work on session-aware detection, detection of latent instructions in retrieved documents, adaptive threshold calibration, and evaluation against stronger defense-aware adversaries.

Data Availability Statement

The public benchmark datasets used in this study—PromptBench, BIPIA, and MS MARCO—are available from their respective original sources under

their respective licences. The 788 manually curated adversarial samples and the 1,440 GPT-4o-generated paraphrase augmentations were constructed by the authors and are available from the corresponding author upon reasonable request. The implementation code and related resources are publicly available at https://github.com/nisarahmedrana/Prompt_Injection_Attack/.

Funding

This work was supported without any funding.

Conflicts of Interest

Nisar Ahmed served as an Associate Editor of the *ICCK Transactions on Information Security and Cryptography* at the time of manuscript submission. To ensure the integrity of the peer-review process, Nisar Ahmed was excluded from the editorial handling, peer review, and final decision-making for this manuscript; these responsibilities were handled independently by another editor. Ali Hassan is affiliated with Sparkverse AI Ltd., Bradford BD1 1AA, West Yorkshire, United Kingdom. The authors confirm that this commercial affiliation did not influence the study design, data collection, analysis, interpretation of results, or the decision to submit for publication. The remaining authors declare no competing interests.

AI Use Statement

The authors declare that GPT-4o, developed by OpenAI, was used to generate paraphrased variants of adversarial samples for dataset augmentation and to assist with automated scoring of Instruction Override goal achievement. GPT-4o was not used to draft the text of this manuscript. The authors carefully reviewed and verified the AI-generated outputs and take full responsibility for the quality, accuracy, and integrity of the AI-assisted research materials and the manuscript.

Ethical Approval and Consent to Participate

Not applicable.

References

- [1] OpenAI. (2024). *GPT-4o System Card*. Retrieved from <https://openai.com/index/gpt-4o-system-card>
- [2] Meta AI. (2024). *Llama 3 Model Card*. Retrieved from <https://ai.meta.com/blog/meta-llama-3/>
- [3] Jiang, A. Q., Sablayrolles, A., Mensch, A., Bamford, C., Chaplot, D. S., de las Casas, D., Bressand, F., Lengyel, G., Lample, G., Saulnier, L., Lavaud, L. R., Lachaux, M.-A., Stock, P., Le Scao, T., Lavril, T., Wang, T., Lacroix,

- T., & El Sayed, W. (2023). Mistral 7B. *arXiv preprint arXiv:2310.06825*. [CrossRef]
- [4] Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., ... & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in neural information processing systems*, 33, 9459-9474.
- [5] OWASP Foundation. (2023). OWASP Top 10 for Large Language Model Applications. Retrieved from <https://owasp.org/www-project-top-10-for-large-language-model-applications/>
- [6] Perez, F., & Ribeiro, I. (2022). Ignore previous prompt: Attack techniques for language models. *arXiv preprint arXiv:2211.09527*. [CrossRef]
- [7] Willison, S. (2022). *Prompt Injection Attacks Against GPT-3*. Retrieved from <https://simonwillison.net/2022/Sep/12/prompt-injection/>
- [8] Halfond, W. G., Viegas, J., & Orso, A. (2006, March). A Classification of SQL Injection Attacks and Countermeasures. In *Proceedings of the IEEE International Symposium on Secure Software Engineering* (pp. 13-15). IEEE. <https://cgi.cse.unsw.edu.au/~meyden/3441/halfond.pdf>
- [9] Greshake, K., Abdelnabi, S., Mishra, S., Endres, C., Holz, T., & Fritz, M. (2023, November). Not what you've signed up for: Compromising real-world llm-integrated applications with indirect prompt injection. In *Proceedings of the 16th ACM workshop on artificial intelligence and security* (pp. 79-90). [CrossRef]
- [10] Branch, H. J., Cefalu, J. R., McHugh, J., Hujer, L., Bahl, A., Iglesias, D. D. C., ... & Darwishi, R. (2022). Evaluating the susceptibility of pre-trained language models via handcrafted adversarial examples. *arXiv preprint arXiv:2209.02128*. [CrossRef]
- [11] Liu, Y., Deng, G., Li, Y., Wang, K., Wang, Z., Wang, X., ... & Liu, Y. (2023). Prompt injection attack against llm-integrated applications. *arXiv preprint arXiv:2306.05499*. [CrossRef]
- [12] Wallace, E., Xiao, K., Leike, R., Weng, L., Heidecke, J., & Beutel, A. (2024). The instruction hierarchy: Training llms to prioritize privileged instructions. *arXiv preprint arXiv:2404.13208*. [CrossRef]
- [13] Rebedea, T., Dinu, R., Sreedhar, M. N., Parisien, C., & Cohen, J. (2023, December). Nemo guardrails: A toolkit for controllable and safe llm applications with programmable rails. In *Proceedings of the 2023 conference on empirical methods in natural language processing: system demonstrations* (pp. 431-445). [CrossRef]
- [14] Inan, H., Upasani, K., Chi, J., Rungta, R., Iyer, K., Mao, Y., ... & Khabsa, M. (2023). Llama guard: Llm-based input-output safeguard for human-ai conversations. *arXiv preprint arXiv:2312.06674*. [CrossRef]
- [15] Phute, M., Helbling, A., Hull, M., Peng, S., Szyller, S., Cornelius, C., & Chau, D. H. (2023). Llm self defense: By self examination, llms know they are being tricked. *arXiv preprint arXiv:2308.07308*. [CrossRef]
- [16] Kumar, A., Agarwal, C., Srinivas, S., Li, A. J., Feizi, S., & Lakkaraju, H. (2023). Certifying llm safety against adversarial prompting. *arXiv preprint arXiv:2309.02705*. [CrossRef]
- [17] Yi, J., Xie, Y., Zhu, B., Kiciman, E., Sun, G., Xie, X., & Wu, F. (2025, July). Benchmarking and defending against indirect prompt injection attacks on large language models. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 1* (pp. 1809-1820). [CrossRef]
- [18] Liu, Y., Deng, G., Xu, Z., Li, Y., Zheng, Y., Zhang, Y., ... & Liu, Y. (2023). Jailbreaking chatgpt via prompt engineering: An empirical study. *arXiv preprint arXiv:2305.13860*. [CrossRef]
- [19] Zou, A., Wang, Z., Carlini, N., Nasr, M., Kolter, J. Z., & Fredrikson, M. (2023). Universal and transferable adversarial attacks on aligned language models. *arXiv preprint arXiv:2307.15043*. [CrossRef]
- [20] Jain, N., Schwarzschild, A., Wen, Y., Somepalli, G., Kirchenbauer, J., Chiang, P. Y., ... & Goldstein, T. (2023). Baseline defenses for adversarial attacks against aligned language models. *arXiv preprint arXiv:2309.00614*. [CrossRef]
- [21] Robey, A., Wong, E., Hassani, H., & Pappas, G. J. (2023). Smoothllm: Defending large language models against jailbreaking attacks. *arXiv preprint arXiv:2310.03684*. [CrossRef]
- [22] Xie, Y., Yi, J., Shao, J., Curl, J., Lyu, L., Chen, Q., Xie, X., & Wu, F. (2023). Defending ChatGPT Against Jailbreak Attack via Self-Reminder. *Nature Machine Intelligence*, 5(12), 1486-1496. [CrossRef]
- [23] Chen, S., Piet, J., Sitawarin, C., & Wagner, D. (2025). {StruQ}: Defending against prompt injection with structured queries. In *34th USENIX Security Symposium (USENIX Security 25)* (pp. 2383-2400). <https://www.usenix.org/conference/usenixsecurity25/presentation/chen-sizhe>
- [24] Hines, K., Lopez, G., Hall, M., Zarfati, F., Zunger, Y., & Kiciman, E. (2024). Defending Against Indirect Prompt Injection Attacks With Spotlighting. *arXiv preprint arXiv:2403.14720*. [CrossRef]
- [25] Zhu, K., Wang, J., Zhou, J., Wang, Z., Chen, H., Wang, Y., ... & Xie, X. (2023, November). Promptrobust: Towards evaluating the robustness of large language models on adversarial prompts. In *Proceedings of the 1st ACM workshop on large AI systems and models with privacy and safety analysis* (pp. 57-68). [CrossRef]
- [26] Schulhoff, S., Pinto, J., Khan, A., Bouchard, L. F., Si, C., Anati, S., ... & Boyd-Graber, J. L. (2023, December). Ignore this title and HackAPrompt: Exposing systemic vulnerabilities of LLMs through a global prompt hacking competition. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing* (pp. 4945-4977). [CrossRef]
- [27] Shostack, A. (2014). *Threat Modeling: Designing for Security*. Wiley. [CrossRef]

- [28] Zhang, Y., & Ippolito, D. Prompts Should not be Seen as Secrets: Systematically Measuring Prompt Extraction Attack Success (2023). *arXiv preprint arXiv:2307.06865*.
- [29] Zou, W., Geng, R., Wang, B., & Jia, J. (2025). PoisonedRAG: Knowledge corruption attacks to Retrieval-Augmented generation of large language models. In *34th USENIX Security Symposium (USENIX Security 25)* (pp. 3827-3844). <https://www.usenix.org/conference/usenixsecurity25/presentation/zou-poisonedrag>
- [30] Bagdasaryan, E., Hsieh, T. Y., Nassi, B., & Shmatikov, V. (2023). Abusing images and sounds for indirect instruction injection in multi-modal LLMs. *arXiv preprint arXiv:2307.10490*. [CrossRef]
- [31] Bai, Y., Kadavath, S., Kundu, S., Askell, A., Kernion, J., Jones, A., ... & Kaplan, J. (2022). Constitutional ai: Harmlessness from ai feedback. *arXiv preprint arXiv:2212.08073*. [CrossRef]
- [32] Karpukhin, V., Oguz, B., Min, S., Lewis, P., Wu, L., Edunov, S., ... & Yih, W. T. (2020, November). Dense passage retrieval for open-domain question answering. In *Proceedings of the 2020 conference on empirical methods in natural language processing (EMNLP)* (pp. 6769-6781). [CrossRef]
- [33] Johnson, J., Douze, M., & Jégou, H. (2019). Billion-scale similarity search with GPUs. *IEEE transactions on big data*, 7(3), 535-547. [CrossRef]
- [34] Wolf, T., Debut, L., Sanh, V., Chaumond, J., Delangue, C., Moi, A., ... & Rush, A. M. (2019). Huggingface's transformers: State-of-the-art natural language processing. *arXiv preprint arXiv:1910.03771*. [CrossRef]
- [35] Reimers, N., & Gurevych, I. (2019, November). Sentence-bert: Sentence embeddings using siamese bert-networks. In *Proceedings of the 2019 conference on empirical methods in natural language processing and the 9th international joint conference on natural language processing (EMNLP-IJCNLP)* (pp. 3982-3992). [CrossRef]
- [36] Loshchilov, I., & Hutter, F. (2017). Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*. [CrossRef]
- [37] Meta AI. (2024). *Llama Guard 2: Safeguarding Human-AI Conversations*. Retrieved from https://github.com/meta-llama/PurpleLlama/blob/main/Llama-Guard2/MODEL_CARD.md
- [38] Nguyen, T., Rosenberg, M., Song, X., Gao, J., Tiwary, S., Majumder, R., & Deng, L. (2016). *MS MARCO: A human generated machine reading comprehension dataset*. OpenReview. <https://openreview.net/forum?id=Hk1iOLcle>
- [39] Zheng, L., Chiang, W. L., Sheng, Y., Zhuang, S., Wu, Z., Zhuang, Y., ... & Stoica, I. (2023). Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in neural information processing systems*, 36, 46595-46623. [CrossRef]
- [40] Zhang, T., Kishore, V., Wu, F., Weinberger, K. Q., & Artzi, Y. (2019). Bertscore: Evaluating text generation with bert. *arXiv preprint arXiv:1904.09675*. [CrossRef]
- [41] Carlini, N., & Wagner, D. (2017, May). Towards evaluating the robustness of neural networks. In *2017 IEEE Symposium on Security and Privacy (SP)* (pp. 39-57). IEEE. [CrossRef]



Gulshan Saleem received her Bachelor's and Master's degrees in Software Engineering and a Ph.D. in Computer Science. Her research interests include machine learning, data analytics, and intelligent systems. She has authored and co-authored peer-reviewed publications and contributed to research on the design and analysis of data-driven computational solutions. (Email: gulshnsaleem26@gmail.com)



Nisar Ahmed (Senior Member, IEEE) received his Ph.D. and Master's degrees in Computer Engineering and a Bachelor's degree in Electrical Engineering. His research focuses on machine learning, deep learning, and AI-driven solutions for cybersecurity and healthcare. He has authored multiple peer-reviewed publications and contributed to research on data-driven security analytics and intelligent systems, with an emphasis on reproducibility and real-world deployment. (Email: nisar.ahmed@utu.fi)



Muhammad Imran Zaman Muhammad Imran Zaman received his Bachelor's and Master's degrees in Computer Science. He currently works in industry, focusing on practical, data-driven solutions for real-world data science problems. His interests include applied machine learning, data analytics, and the deployment of intelligent systems in production environments. (Email: imran.zaman@sparkverse.ai)



Ali Hassan Ali Hassan received his B.S. degree in Computer Science. He is currently working in industry, focusing on the deployment of intelligent systems in production environments. He also has experience as a front-end web developer, contributing to the design and development of user-centric web applications. (Email: ali.hassan@sparkverse.ai)



Umar Mujahid Khokhar Umar Mujahid Khokhar is an Associate Professor of Information Technology at Georgia Gwinnett College. He earned his Ph.D. in Information Security from Bahria University and has extensive academic experience in cybersecurity education and research. His work focuses on information security, privacy-preserving systems, ultra-lightweight cryptography, and secure hardware implementations for pervasive computing environments. (Email: ukhokhar@ggc.edu)