

SoK: Evaluating Jailbreak Guardrails for Large Language Models

Xunguang Wang*, Zhenlan Ji*, Wenxuan Wang†, Zongjie Li*, Daoyuan Wu‡, Shuai Wang*§

*The Hong Kong University of Science and Technology, †Renmin University of China, ‡Lingnan University
{xwanghm, zjiaei, zligo, shuaiw}@cse.ust.hk, wangwenxuan@ruc.edu.cn, daoyuanwu@ln.edu.hk

Abstract—Large Language Models (LLMs) have achieved remarkable progress, but their deployment has exposed critical vulnerabilities, particularly to jailbreak attacks that circumvent safety alignments. Guardrails—external defense mechanisms that monitor and control LLM interactions—have emerged as a promising solution. However, the current landscape of LLM guardrails is fragmented, lacking a unified taxonomy and comprehensive evaluation framework. In this Systematization of Knowledge (SoK) paper, we present the first holistic analysis of jailbreak guardrails for LLMs. We propose a novel, multi-dimensional taxonomy that categorizes guardrails along six key dimensions, and introduce a Security-Efficiency-Utility evaluation framework to assess their practical effectiveness. Through extensive analysis and experiments, we identify the strengths and limitations of existing guardrail approaches, provide insights into optimizing their defense mechanisms, and explore their universality across attack types. Our work offers a structured foundation for future research and development, aiming to guide the principled advancement and deployment of robust LLM guardrails.

1. Introduction

Large Language Models (LLMs) have demonstrated remarkable capabilities across a wide range of applications, revolutionizing fields from natural language understanding to content generation [1], [2], [3], [4], [5], [6]. However, their increasing sophistication and widespread adoption have also exposed significant vulnerabilities. A prominent concern is their susceptibility to *jailbreak attacks* [7], [8], where adversaries craft malicious inputs to bypass safety alignments and elicit harmful, biased, or unethical responses. The proliferation of such attacks underscores the urgent need for robust defense mechanisms. Among various defense strategies, *guardrails* [9], [10], [11] have emerged as a promising approach, aiming to monitor and control LLM interactions without altering the underlying model’s weights or core functionalities.

Guardrail-based defenses offer a distinct advantage over other defense methods (e.g., tuning-based approaches [12]) as they can effectively filter jailbreak attempts while preserving the integrity of the target LLM’s original output capabilities. Despite their potential, the current landscape of LLM guardrails is characterized by *siloed innovation*.

Numerous research teams and organizations have proposed various guardrail solutions, often tailored to specific scenarios, attack vectors (e.g., focusing primarily on single-turn attacks), or proprietary systems. This ad-hoc development has resulted in a fragmented ecosystem of defense mechanisms, lacking a unified understanding or a systematic classification framework to position and compare these disparate efforts.

The absence of a systematic perspective contributes directly to a critical limitation in existing guardrails: a general lack of *universality*. Many solutions are not readily adaptable across different LLMs, attack types, or deployment contexts. Furthermore, current evaluation practices for LLM guardrails often fall short of reflecting real-world operational constraints. Evaluations predominantly focus on raw defense efficacy against specific jailbreak benchmarks, frequently overlooking crucial factors such as *computational cost* (e.g., inference latency, GPU resource consumption) and *utility* (e.g., the rate of misclassifying benign prompts as malicious, thereby degrading user experience). This narrow evaluation scope hinders a comprehensive understanding of the practical trade-offs involved in deploying guardrails.

To address these critical gaps, this Systematization of Knowledge (SoK) paper provides the first comprehensive analysis and structuring of the rapidly evolving field of jailbreak guardrails for LLMs. We aim to consolidate the disparate research efforts, offering a clear and structured understanding of the current state-of-the-art. Our primary contributions are threefold: (1) we propose a novel, multi-dimensional taxonomy for classifying LLM guardrails, enabling a nuanced understanding of their design characteristics; (2) we introduce a holistic evaluation framework centered on the *Security-Efficiency-Utility* trifecta, promoting more practical and comprehensive assessments; and (3) we conduct extensive analysis based on our framework, yielding valuable insights into the performance of existing guardrails and identifying promising avenues for future research. Specifically, our contributions are as follows:

- **A Multi-Dimensional Guardrail Taxonomy:** We propose the first comprehensive taxonomy to categorize LLM guardrails along six critical dimensions:
 - *Intervention Stage:* Characterizing when the guardrail operates (Pre-processing, Intra-processing, or Post-processing of LLM interactions).
 - *Technical Paradigm:* Identifying the underlying mechanism (Rule-based, Model-based, or LLM-based).
 - *Safety Granularity:* Defining the scope of the guardrail

§Corresponding author.

detection (Token-level, Sequence-level, or Session-level).

- *Reactivity*: Distinguishing between static (pre-defined) and dynamic (adaptive) defense strategies.
- *Applicability*: Considering the guardrail’s requirements regarding LLM access (White-box vs. Black-box).
- *Interpretability*: Assessing the transparency of the guardrail’s decision-making process.
- **A Security-Efficiency-Utility (SEU) Evaluation Framework**: We introduce a novel framework for evaluating guardrails that balances three crucial aspects:
 - *Security*: Measuring the defense performance against a diverse range of jailbreak attacks.
 - *Efficiency*: Quantifying the operational overhead, including inference delay and GPU memory consumption.
 - *Utility*: Assessing the impact on legitimate user interactions, primarily through the false positive rate on benign queries.
- **Experimental Findings and Optimization Insights**: We leverage our taxonomy and evaluation framework to analyze existing guardrails and explore future directions:
 - We conduct a tri-objective (SEU) evaluation of mainstream guardrail methods to identify balanced solutions and those effective against diverse jailbreak categories.
 - We investigate specific hypotheses, such as the efficacy of session-level guardrails against multi-turn attacks, the influence of intervention stage on latency, the impact of technical paradigms on resource consumption, and the relationship between safety granularity and utility.
 - We explore the *universality* of guardrails by assessing their performance against other attack modalities, such as prompt injection attacks.

This SoK aims to provide researchers and practitioners with a clear roadmap for understanding, developing, and deploying LLM jailbreak guardrails. By systematizing existing knowledge and proposing a comprehensive evaluation methodology, we hope to foster more principled advancements in this critical area of LLM security. The code is available at <https://github.com/xunguangwang/SoK4JailbreakGuardrails>.

2. Jailbreak Attacks in LLMs

In this section, we first formally describe the jailbreak in LLMs and then introduce several typical jailbreak methods. **Jailbreak Formulation.** The jailbreak phenomenon indicates that specific malicious instructions can bypass the safety mechanisms of LLMs, leading to the generation of harmful or unethical outputs. This is particularly concerning as it highlights the potential for adversaries to exploit vulnerabilities in LLMs to produce toxic or harmful content. The jailbreak process can be viewed as a two-step procedure: (1) crafting an adversarial prompt P that elicits a harmful response from the LLM, and (2) evaluating the generated response R against a predefined harmful objective

G using a classifier JUDGE. JUDGE returns ‘True’ if the generated response R meets the harmful objective G , i.e., $\text{JUDGE} = \text{True}$, otherwise ‘False’. Let $\mathcal{T}$ represent the LLM’s vocabulary. Formally, we can define the classifier $\text{JUDGE} : \mathcal{T}^* \times \mathcal{T}^* \rightarrow \{\text{True}, \text{False}\}$. The adversary’s goal is to maximize the probability of generating responses classified as satisfying the harmful objective G . This can be expressed mathematically as:

$$\sup_{P \in \mathcal{T}^*} \Pr_{R \sim \text{LLM}(P)} [\text{JUDGE}(R, G) = \text{True}], \quad (1)$$

where $\Pr$ denotes the probability, which accounts for the inherent stochasticity of the LLM’s outputs when processing the input prompt P . The adversary iteratively refines prompts to identify those that maximize the likelihood of producing outputs deemed harmful by the classifier.

Existing Jailbreak Attacks. Jailbreak attacks can be broadly categorized into two types: single-turn and multi-turn jailbreaks. Single-turn jailbreaks involve crafting a single prompt to elicit harmful responses, while multi-turn jailbreaks exploit the interactive nature of LLMs by engaging in a dialogue with the model over multiple turns. Due to the maturity of research on single-turn attacks and the relative scarcity of multi-turn attack studies, a further breakdown of the single-turn attack is necessary. Building on prior works [11], [13] that categorize jailbreak attacks by their technical paradigms, we divide single-turn attacks into similar four types: manual, optimization-based, generation-based and implicit jailbreaks.

- **Manual Jailbreaks.** These attacks involve crafting prompts that exploit vulnerabilities in LLMs [14], [15], [16], [17], [18]. Wei et al. [15] identified two key weaknesses—out-of-distribution inputs and conflicts between safety objectives and model capabilities—to inform prompt design. Deng et al. [18] introduced AIM (Always Intelligent and Machiavellian), a proof-of-concept jailbreak prompt that served as a foundation for generating additional adversarial prompts. Shen et al. [17] proposed JailbreakHub, a crowdsourcing framework for collecting diverse jailbreak prompts.
- **Optimization-based Jailbreaks.** These methods iteratively refine adversarial prompts using techniques like gradient-based optimization or search strategies [7], [8], [19], [20], [21]. GCG [7] introduced a greedy coordinate gradient method to optimize adversarial suffixes, enabling transferable jailbreaks across models and prompts. Sitawarin et al. [19] extended this with GCG++, leveraging a proxy model to enhance optimization. Beyond gradient-based techniques, JSAA [20] employed random search for suffix optimization and BON [22] induces jailbreaks by repeatedly sampling augmentations of a harmful instruction until one succeeds, while AutoDAN [8] used a hierarchical genetic algorithm to create human-readable jailbreak prompts. RLbreaker [23] utilized reinforcement learning to efficiently search for adversarial prompts, outperforming stochastic methods like JSAA and AutoDAN.
- **Generation-based Jailbreaks.** These attacks use auxiliary LLMs to produce adversarial prompts [18], [24],

[25], [26], [27], [28]. PAIR [25] employs a feedback loop where the attacking LLM adjusts outputs based on the target LLM’s responses. Mehrotra et al. [26] enhanced this approach using tree-of-thought reasoning [29]. LLM-Fuzzer [30] automates adversarial prompt generation by mutating human-written templates. Additionally, Advprompter [27] trains a fine-tuned LLM to create both effective and human-readable adversarial prompts.

- **Implicit Jailbreaks.** This class of attacks is derived from the Linguistics-based and Encoding-based jailbreak categories identified in JBSHield [13]. Given that both strategies conceal malicious intent within the query text to circumvent LLM safety mechanisms, they are collectively categorized as implicit jailbreaks. For instance, Handa et al. [31] demonstrated word substitution as a simple evasion method. DrAttack [32] decomposes harmful prompts into smaller, less detectable sub-prompts. Puzzler [33] embeds clues within queries to guide the LLM toward producing harmful outputs indirectly. Another approach involves translating harmful prompts into languages where LLM safety mechanisms are weaker [15], [34], [35], [36], [37], [38]. Deng et al. [34] and Yong et al. [35] found that low-resource languages, such as Zulu, often exhibit less robust safety alignment. Obfuscation techniques [39], including encoding or encrypting harmful prompts, further reduce LLM sensitivity to malicious inputs [15], [38].
- **Multi-turn Jailbreaks.** One multi-turn attack strategy is the fine-grained task decomposition, which decomposes the original malicious query into several less harmful sub-questions [40], [41], [42]. While this decomposition strategy successfully circumvents current safety mechanisms, it may be easily mitigated by including these finer-grained harmful queries in safety training data. Alternatively, researchers propose to use human red teamers to expose vulnerabilities of LLMs against multi-turn attacks [43]. Moreover, Yang et al. [44] depend on the heuristics from [25] and its seed examples to implement their attacks. Crescendo [45] gradually steers benign initial queries towards more harmful topics. The implementation of Crescendo is based on the fixed and human-crafted seed instances, making it challenging to generate diverse and effective attacks. By contrast, ActorAttack [46] proposes to discover diverse attack clues inside the model’s prior knowledge. X-Teaming [47] achieves more effective and diverse multi-turn attacks by adaptive collaborative agents for planning, attack optimization, and verification.

3. Definition, Taxonomy & Evaluation

3.1. Jailbreak Defense

Jailbreak defense comprises methods to protect LLMs from jailbreak attacks intended to elicit harmful or policy-violating content. We classify these defenses based on their point of application into four categories: prompt-based, tuning-based, refining-based defenses, and guardrails, which respectively focus on modifying prompts, optimizing model

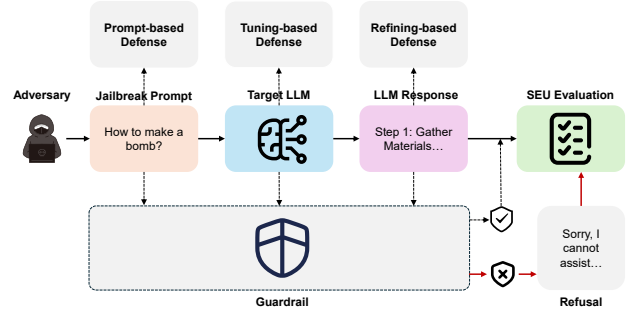


Figure 1. Illustration of a guardrail pipeline.

parameters, controlling model outputs, and external monitoring.

Prompt-based defenses alter the input prompt to enhance the LLM’s adherence to safety guidelines [16], [48], [49], [50]. For instance, Self-Reminder [48] prepends ethical reminders, while ICD [16] provides examples of harmful queries and safe responses. **Tuning-based defenses** optimize the LLM’s parameters to improve its inherent safety. This includes techniques like reinforcement learning from human feedback (RLHF) [51], [52], machine unlearning to erase harmful knowledge [53], model editing to patch vulnerabilities [54], supervised fine-tuning on curated safety datasets [55], [56], and adversarial training to bolster robustness against attacks [12]. **Refining-based defenses** steer the output generation process in real-time [57], [58], [59]. For example, RAIN [57] uses token-level evaluation to guide rewinds for safer generation, whereas Aligner [59] employs a lightweight model to correct initial responses towards alignment. **Guardrails** are external modules that monitor and control the interaction [9], [10], [60]. As illustrated in Figure 1, they can operate on inputs (pre-processing), internal model states (intra-processing), or outputs (post-processing). Distinct from the other categories, guardrails do not modify the target LLM’s inference; rather, they detect harmfulness and filter the interaction to prevent jailbreaks.

3.2. Jailbreak Guardrail Definition

A “Jailbreak Guardrail” refers to a specialized security mechanism designed for LLM systems, specifically to detect and prevent “jailbreak” attacks [9], [10], [11], [60]. Such guardrails typically function as a defensive layer, scrutinizing user inputs before they reach the LLM or vetting the model’s outputs before they are presented to the user. The primary objective is to ensure that the LLM does not generate harmful, unethical, or policy-violating content.

In the context of the jailbreak formulation introduced in Section 2, where an adversary crafts a prompt P aiming to elicit a response $R = \text{LLM}(P)$ such that $\text{JUDGE}(R, G) = \text{True}$ (indicating a harmful outcome based on objective G), a jailbreak guardrail introduces an additional checkpoint. Let $\mathcal{G}_R$ denote the guardrail system, and $\text{Assess}(\mathcal{G}_R, X)$ be its assessment function, which returns *allow* if content X (either P , R , or the internal feature F of the LLM) is deemed permissible, and *block* otherwise. When *block* is

executed, the final response R is replaced by a safe response R' , such as “Sorry, I cannot assist with that”, as shown in Figure 1.

A jailbreak attack is considered successful in the presence of such a guardrail if, and only if, the protective mechanisms of both the target LLM (i.e., its inherent safety alignment) and the guardrail are circumvented. This means the guardrail must deem the interaction (either the input prompt or the generated output) as acceptable, while the target LLM still produces content classified as harmful. More formally, if the guardrail inspects the input prompt P , a successful jailbreak occurs when:

$$\text{Assess}(\mathcal{G}_R, P) = \text{allow} \wedge \text{JUDGE}(R, G) = \text{True}. \quad (2)$$

Alternatively, if the guardrail inspects the model’s internal features F or output $R = \text{LLM}(P)$, a successful jailbreak is characterized respectively by:

$$\text{Assess}(\mathcal{G}_R, F) = \text{allow} \wedge \text{JUDGE}(R, G) = \text{True}, \quad (3)$$

$$\text{Assess}(\mathcal{G}_R, R) = \text{allow} \wedge \text{JUDGE}(R, G) = \text{True}. \quad (4)$$

This highlights that a successful adversary must not only craft a prompt that bypasses the LLM’s internal safety measures but also deceive the guardrail into permitting the harmful interaction or content.

As jailbreak techniques become increasingly sophisticated and diverse (as noted in Section 2), these guardrails face mounting challenges. They must evolve beyond detecting overtly malicious requests to identify subtle and nuanced jailbreak patterns and adversarial manipulations. The continuous enhancement of jailbreak guardrails is therefore critical for improving the safety, security, and regulatory compliance of AI applications.

3.3. Jailbreak Guardrail Taxonomy

This section categorizes existing guardrail approaches along several key dimensions. Our taxonomy considers:

- **Intervention Stages:** This dimension delineates *when* the guardrail operates within the LLM interaction pipeline—either at *pre-processing* (before the input reaches the LLM), *intra-processing* (during the LLM’s inference), or *post-processing* (after the LLM generates an output).
- **Technical Paradigms:** This refers to the underlying *methodology* employed by the guardrail. Approaches are classified as *rule-based* (relying on predefined rules or patterns), *model-based* (using statistical models or classifiers), or *LLM-based* (leveraging another LLM for analysis and decision-making).
- **Safety Granularity:** This specifies the *level of detail* at which the safety analysis is performed. It can be *token-level* (examining individual words or sub-word units), *sequence-level* (evaluating entire prompts or responses), or *session-level* (considering the context of the entire conversation history).
- **Reactiveness:** This dimension distinguishes how a guardrail responds to potentially harmful inputs. *Static*

defenses analyze inputs without modification, whereas *dynamic* defenses actively alter inputs—for example, through mutation or perturbation—to neutralize adversarial properties while aiming to preserve overall semantic meaning.

- **Applicability:** This criterion assesses the guardrail’s suitability for different LLM access models, with a particular emphasis on whether the mechanism can be effectively applied to *black-box* LLMs (i.e., closed-source models or those accessed via remote APIs where internal states are not accessible).
- **Explainability:** This focuses on whether the guardrail method provides interpretable insights into its safety judgments or offers clear rationales for the decisions it makes.

This multi-faceted classification provides a comprehensive framework for understanding and navigating the landscape of LLM guardrails. We have comprehensively compiled existing works on jailbreak guardrails by this taxonomy, as summarized in Table 1.

3.4. Guardrail Evaluation Framework

To enable a comprehensive and practical assessment of LLM guardrails, we propose the *Security-Efficiency-Utility* (SEU) Evaluation Framework. This framework is designed to capture the essential trade-offs involved in deploying guardrails in real-world LLM systems, moving beyond the narrow focus on raw defense efficacy. Below, we detail the three core dimensions of our framework and the specific metrics used for each.

Security: Defense Effectiveness. The primary objective of any guardrail is to enhance the security of LLM systems by mitigating jailbreak attacks. We evaluate defense effectiveness using two complementary metrics:

- **Attack Success Rate (ASR):** ASR measures the proportion of adversarial attempts that successfully bypass the guardrail and elicit harmful or unethical responses from the target LLM. Formally, it is defined as the percentage of attack queries for which the LLM system equipped with the guardrail fails to block or mitigate the attack. A lower ASR indicates a stronger defense.
- **Pass Guardrail Rate (PGR):** PGR measures the proportion of jailbreak attempts that successfully bypass the guardrail, indicating that the guardrail has classified the attempt as safe. For pre-processing and intra-processing guardrails, this refers to the proportion of malicious requests that the guardrail incorrectly identifies as benign. For post-processing guardrails, this refers to the proportion of instances where the guardrail fails to detect harmful content in the LLM’s response to a jailbreak attempt. A lower PGR signifies a more effective guardrail in blocking attacks.

Efficiency: Computational Overhead. In practical deployments, the operational efficiency of guardrails is a critical consideration, as excessive overhead can degrade user experience and increase infrastructure costs. We assess efficiency along two axes:

- **Extra Delay:** This metric captures the additional response latency introduced by the guardrail. It is computed as the difference between the end-to-end response time of the guardrail + LLM system and that of the standalone target LLM. Formally,

$$\text{Extra Delay} = T_{\text{guardrail} + \text{LLM}} - T_{\text{LLM}}, \quad (5)$$

where $T_{\text{guardrail} + \text{LLM}}$ and T_{LLM} denote the average response times with and without the guardrail, respectively.

- **GPU Memory Overhead:** This metric measures the increase in peak GPU memory consumption resulting from the integration of the guardrail. It is defined as the difference between the maximum GPU memory usage of the guardrail + LLM system and that of the target LLM alone:

$$\text{GPU Overhead} = M_{\text{guardrail} + \text{LLM}} - M_{\text{LLM}}, \quad (6)$$

where $M_{\text{guardrail} + \text{LLM}}$ and M_{LLM} represent the peak GPU memory usage with and without the guardrail, respectively.

Utility: Impact on Benign Queries. A robust guardrail should not only block malicious inputs but also preserve the utility of the LLM for legitimate users. We quantify utility loss using the following metric:

- **False Positive Rate (FPR):** FPR measures the proportion of benign (non-malicious) queries that are incorrectly flagged or blocked by the guardrail. It is defined as the percentage of normal user queries that are misclassified as attacks. A lower FPR indicates better utility preservation, as the guardrail minimally disrupts legitimate interactions with the LLM.

Discussion. By jointly considering Security, Efficiency, and Utility, the SEU Evaluation Framework provides a holistic basis for comparing and optimizing LLM guardrails. This tri-objective perspective enables the identification of solutions that achieve a balanced trade-off, rather than excelling in only one dimension at the expense of others. In our experimental analysis (§5 & §6), we employ this framework to systematically evaluate mainstream guardrail methods, offering actionable insights for both researchers and practitioners.

Alignment with Taxonomy Dimensions. The proposed taxonomy in §3.3 is intentionally structured to align with the empirical evaluation dimensions of the SEU framework, which form the core contribution of this work. Particularly, we would analyze the first three taxonomy dimensions (Intervention Stage, Technical Paradigm, and Safety Granularity) in depth, because they have a direct, measurable impact on Security, Efficiency, and Utility, forming the foundation of our quantitative analysis. In contrast, the latter three dimensions (i.e., Reactiveness, Applicability, and Explainability) are more qualitative in nature and represent critical future research directions for enhancing the adaptability, deployment flexibility, and transparency of guardrail systems. By focusing our empirical analysis on the dimensions most directly tied to operational performance, we provide a principled and actionable foundation for evaluating and optimizing jailbreak guardrails in practice.

4. Guardrail Analysis Based on Taxonomy

4.1. Intervention Stages

Guardrail mechanisms can be deployed at different stages of the LLM interaction pipeline, including pre-processing, intra-processing, and post-processing. Each stage serves a distinct purpose in identifying and mitigating jailbreak attempts:

Pre-processing Guardrails. These mechanisms operate on user inputs before they reach the target LLM, functioning as the first line of defense against jailbreak attempts. Pre-processing guardrails typically employ detection algorithms to identify potentially harmful prompts and then block them entirely. These guards are particularly valuable for their ability to prevent harmful prompts from ever reaching the model, thus conserving computational resources and reducing potential risks.

Early methods, such as Detecting Perplexity [64] and Perplexity Filter [65], compute the perplexity of input prompts to detect potential adversarial inputs. However, this approach is limited to GCG [7], [19], [28] attacks with unreadable adversarial suffixes amplifying the perplexity.

A more direct approach is to identify the semantic harmfulness of input sequences. Some methods focus on directly identifying toxic phrases or excerpts within the input text [11], [60], [88], while others assess the overall semantic harmfulness of the entire input [10], [60], [61], [62], [66], [73], [74], [79], [80], [82], [83], [84], [87], [89], [91], [100]. For instance, PromptGuard [80] and OpenAI Moderation [62] fine-tune pre-trained classifiers to assess the safety of input prompts. However, pre-processing guardrails may struggle with novel jailbreak techniques that do not exhibit clear patterns, e.g., implicit attack DrAttack [32] conceals malicious content within benign-looking prompts.

A more fundamental approach is to analyze the true intent of the query to filter out jailbreak requests, based on the premise that jailbreak attempts always involve malicious output targets. Leveraging the powerful language understanding capabilities of LLMs, we can directly utilize LLMs to identify the real intentions of requests to determine whether they are jailbreak attempts [11], [78], [95], [100]. For example, SelfDefend [11] with the intent prompt first summarizes the input intention and then assesses whether it constitutes a jailbreak request. Recently, some studies have employed LLM reasoning capabilities to analyze input intent [92], [93], [99] before the safety judgments. For instance, X-Guard [93] employs deep thinking to evaluate potential harms.

Summary 1: *Pre-processing guardrails are the first line of defense against jailbreak attempts, operating on user inputs before they reach the target LLM. They have evolved from simple perplexity detection to semantic harmfulness identification and, most recently, to LLM-based reasoning for analyzing input intent. This evolution is driven by the need to address increasingly sophisticated and covert attack methods.*

TABLE 1. WORKS ON GUARDRAILS CATEGORIZED BY 6 DIMENSIONS. THE BLACK CIRCLE INDICATES THE GUARDRAIL BELONGS TO THIS DIMENSION, AND THE WHITE CIRCLE OTHERWISE.

Paper	Venue	Intervention Stages			Technical Paradigms			Safety Granularity			Reactiveness		Applicability	Explainability
		Pre-processing	Intra-processing	Post-processing	Rule	Model	LLM	Token	Sequence	Session	Static	Dynamic		
Perspective API [61]	KDD'22 (2202.11176)	●	○	●	○	●	○	○	●	○	●	○	●	○
OpenAI Moderation [62]	AAAI'23 (2208.03274)	○	○	●	○	○	●	○	●	○	●	○	●	○
Self Defense [63]	arXiv:2308.07308	○	○	●	○	○	●	○	●	○	●	○	●	○
Detecting Perplexity [64]	arXiv:2308.14132	●	○	○	○	●	○	○	●	○	●	○	●	○
Perplexity Filter [65]	arXiv:2309.00614	●	○	○	○	●	○	○	●	○	●	○	●	○
erase-and-check [66]	COLM'24 (2309.02705)	●	○	○	○	●	○	○	●	○	●	○	●	○
SmoothLLM [67]	arXiv:2310.03684	○	○	●	●	○	○	●	○	○	○	●	●	○
NeMo Guardrails [9]	EMNLP'23 (2310.10501)	●	○	●	○	○	●	○	●	○	○	●	●	○
Llama Guard [10]	arXiv:2312.06674	●	○	●	○	○	●	○	●	○	○	●	●	○
GradSafe [68]	ACL'24 (2402.13494)	○	●	○	○	○	○	○	●	○	○	○	○	○
SemanticSmooth [69]	arXiv:2402.16192	○	○	●	○	○	●	○	●	○	○	●	●	○
LLMGuard [70]	arXiv:2403.00826	●	○	●	●	○	○	○	●	○	○	○	○	○
Gradient Cuff [71]	NeurIPS'24 (2403.00867)	○	●	○	○	○	○	○	●	○	○	○	○	○
AutoDefense [72]	arXiv:2403.04783	○	○	●	○	○	●	○	●	○	○	○	○	○
RigorLLM [73]	ICML'24 (2403.13031)	●	○	○	○	○	○	○	●	○	○	○	○	○
Aegis [74]	arXiv:2404.05993	●	○	●	○	○	○	○	○	○	○	○	○	○
LLMGuardrail [75]	CCS'24 (2405.04160)	○	○	○	○	○	○	○	○	○	○	○	○	○
RSAA [76]	CAMLS'24 (2406.03230)	○	○	○	○	○	○	○	○	○	○	○	○	○
Circuit Breaking [77]	NeurIPS'24 (2406.04313)	○	○	○	○	○	○	○	○	○	○	○	○	○
SelfDefend [11]	USENIX Security'25 (2406.05498)	●	○	○	○	○	○	○	○	○	○	○	○	○
GuardAgent [78]	ICML'25 (2406.09187)	○	○	○	○	○	○	○	○	○	○	○	○	○
WildGuard [60]	NeurIPS'24 (2406.18495)	●	○	○	○	○	○	○	○	○	○	○	○	○
R ² -Guard [79]	ICLR'25 (2407.05557)	●	○	○	○	○	○	○	○	○	○	○	○	○
Prompt Guard [80], [81]	Hugging Face (22 July 2024)	●	○	○	○	○	○	○	○	○	○	○	○	○
PrimeGuard [82]	arXiv:2407.16318	○	○	○	○	○	○	○	○	○	○	○	○	○
ShieldGemma [83]	arXiv:2407.21772	○	○	○	○	○	○	○	○	○	○	○	○	○
Adaptive Guardrail [84]	arXiv:2408.08959	○	○	○	○	○	○	○	○	○	○	○	○	○
EEG-Defender [85]	ICONIP'25 (2408.11308)	○	○	○	○	○	○	○	○	○	○	○	○	○
HSF [86]	WWW'25 (2409.03788)	○	○	○	○	○	○	○	○	○	○	○	○	○
MoJE [87]	AIES'24 (2409.17699)	○	○	○	○	○	○	○	○	○	○	○	○	○
Rapid Response [88]	arXiv:2411.07494	○	○	○	○	○	○	○	○	○	○	○	○	○
Pretrained Embeddings [89]	arXiv:2412.01547	○	○	○	○	○	○	○	○	○	○	○	○	○
Token Highlighter [90]	AAAI'25 (2412.18171)	○	○	○	○	○	○	○	○	○	○	○	○	○
Aegis2.0 [91]	arXiv:2501.09004	○	○	○	○	○	○	○	○	○	○	○	○	○
COT Fine-Tuning [92]	arXiv:2501.13080	○	○	○	○	○	○	○	○	○	○	○	○	○
GuardReasoner [93]	arXiv:2501.18492	○	○	○	○	○	○	○	○	○	○	○	○	○
Constitutional Classifiers [94]	arXiv:2501.18837	○	○	○	○	○	○	○	○	○	○	○	○	○
JBShield [13]	USENIX Security'25 (2502.07557)	○	○	○	○	○	○	○	○	○	○	○	○	○
EDDF [95]	ACL Findings'25 (2502.19041)	○	○	○	○	○	○	○	○	○	○	○	○	○
CURVALID [96]	arXiv:2503.03502	○	○	○	○	○	○	○	○	○	○	○	○	○
MirrorShield [97]	arXiv:2503.12931	○	○	○	○	○	○	○	○	○	○	○	○	○
JailGuard [98]	TOSEM'25 (19 March 2025)	○	○	○	○	○	○	○	○	○	○	○	○	○
X-Guard [99]	arXiv:2504.08848	○	○	○	○	○	○	○	○	○	○	○	○	○
Continuous Detector [100]	arXiv:2504.19440	○	○	○	○	○	○	○	○	○	○	○	○	○
Active Monitoring [100]	arXiv:2504.19440	○	○	○	○	○	○	○	○	○	○	○	○	○

Intra-processing Guardrails. These guardrails operate during the LLM’s inference process, analyzing internal model features or gradients to detect potential jailbreak attempts. Unlike pre-processing methods, intra-processing guardrails can observe how the model processes inputs internally, providing deeper insights into potential vulnerabilities.

On one hand, intra-processing guardrails rely on gradient information to identify potential jailbreak attempts. These methods analyze the gradients of the model’s inputs or parameters during inference to identify unusual patterns or anomalies that may indicate adversarial inputs. For example, GradSafe [68] computes the similarity between the input’s gradient w.r.t. the safety-critical parameters and the unsafe reference gradients. Gradient Cuff [71] compares the gradient norm of refusal loss w.r.t. the query prompt with a threshold. Token Highlighter [90] uses the gradient norm of the affirmation loss for each token in the user query to

locate the jailbreak-critical tokens.

On the other hand, intra-processing guardrails can analyze the model’s internal states for jailbreak detection [13], [75], [76], [77], [85], [86]. These methods leverage the model’s hidden states or other internal representations, to identify patterns indicative of jailbreak attempts. For example, Circuit Breaking [77] interrupts the LLM to output harmful content when harmful states are detected. JB-Shield [13] analyzes the differences of the LLM’s internal states between the jailbreak prompts and the benign queries. These approaches can provide more nuanced insights into the model’s behavior and vulnerabilities, enabling more effective detection of sophisticated jailbreak techniques. However, these approaches typically require white-box access to the target model, which limits their applicability to open-source LLMs or scenarios where model internals are accessible.

Summary 2: *Intra-processing guardrails operate during the LLM's inference process, analyzing internal model features or gradients to detect potential jailbreak attempts. They provide deeper insights into vulnerabilities but require white-box access to the model, limiting their applicability.*

Post-processing Guardrails. These mechanisms evaluate the LLM's generated outputs to identify and filter harmful content. As jailbreak attacks inherently aim to produce harmful outputs, post-processing guardrails serve as a crucial last line of defense. This ensures that even if malicious prompts circumvent earlier detection stages, their resultant outputs can still be intercepted.

Given that post-processing guardrails scrutinize the LLM's generated outputs, a primary strategy involves the direct detection of harmfulness within these responses. The most elementary of these methods employ keyword-based detectors to assess the safety of the LLM's outputs, primarily focusing on determining its jailbroken state [67], [70]. Building upon this foundational technique, a more sophisticated approach involves training dedicated classifiers to distinguish between harmful and harmless responses [61], [62], [70], [79]. For instance, initiatives like the Perspective API [61] and OpenAI Moderation [62] have developed transformer-based classifiers engineered to predict the probability of harmful content appearing in an LLM's response. Similarly, R^2 -Guard [79] embeds safety knowledge into probabilistic graphical models, enabling the computation of unsafe probabilities for any given LLM outputs. Elevating this classification paradigm further, albeit with increased computational demands, some techniques leverage the reasoning capabilities of other LLMs to assess response safety [9], [10], [60], [63], [69], [72], [74], [83], [91], [93], [94], [99], [100]. Self Defense [63], for example, filters harmful content by querying an LLM about the harmfulness of the initial response. Llama Guard [10] takes a more contextual approach by considering the input prompt in conjunction with the output to determine the risk category of the response. Progressing towards even more thorough analysis, GuardReasoner [93] and X-Guard [99] employ chain-of-thought reasoning before rendering a safety judgment on the LLM's output.

Beyond directly assessing the semantic harmfulness of the response, an alternative category of methods ingeniously leverages the discrepancies in outputs that arise from disrupting the adversarial characteristics of the initial jailbreak prompt [67], [69], [97], [98]. SmoothLLM [67], for instance, operates by randomly perturbing or permuting multiple copies of a given input prompt to generate a set of responses; the safety of the original request is then determined by a voting mechanism based on these perturbed responses. Advancing this concept, SemanticSmooth [69] and JailGuard [98] implement more complex mutations than the simple character-level alterations used by SmoothLLM, such as paraphrasing the prompt or translating it into other languages. In a similar vein, MirrorShield [97] generates "mirror" prompts that aim to preserve the syntactic structure of the input while ensuring its semantic safety. These

mutation-based methods capitalize on the inherent properties of jailbreak prompts to identify adversarial inputs, rendering them particularly effective against sophisticated attacks. Nevertheless, they may incur additional computational overhead due to the requirement for numerous response evaluations or intricate input transformations.

Although detecting the output of LLMs may appear more straightforward than deciphering ambiguous prompts and internal features, later methodologies will integrate the input prompts to thoroughly evaluate the safety of the query. However, post-processing safeguards may incur more latency than other paradigms due to the requirement of awaiting the LLM's response. Furthermore, mutation-based techniques that mandate multiple response evaluations are suspected to exacerbate this latency.

Summary 3: *Post-processing guardrails, by operating on the LLM's generated outputs to identify and filter harmful content, act as an essential safeguard. They intercept potentially harmful outputs that bypass earlier detection stages. However, over-reliance on the nature of responses may cause noticeable delay, particularly when employing mutation-based techniques that need multiple evaluations.*

Drawing upon a classification by intervention stages, our analysis reveals a critical gap in the current literature, which motivates the RQ below. As no prior work has systematically investigated this specific dimension, we undertake a thorough examination in this paper.

RQ 1: *Pre-processing guardrails can reject harmful inputs before they reach the target LLM, intra-processing mechanisms operate concurrently during LLM inference, and post-processing techniques must await LLM's outputs. This raises a pertinent question: To what extent does the specific intervention stage of a guardrail, be it pre-processing, intra-processing, or post-processing, influence overall response latency?*

4.2. Technical Paradigms

Guardrail mechanisms employ diverse technical approaches to detect and mitigate jailbreak attempts, including rule-based, model-based, and LLM-based approaches.

Rule-based Guardrails. These guardrails operate by employing predefined rules, patterns, or heuristics to detect potentially harmful inputs or outputs of LLMs. A typical rule-based approach includes utilizing keywords or regular expressions to identify specific patterns tied to harmful content. For instance, SmoothLLM, as referenced in [67], leverages keyword-based detectors to assess the safety of the LLM's outputs, primarily to determine its jailbroken state. Similarly, the PII Detector mentioned in [70] uses regular expressions to identify personally identifiable information, such as phone numbers and emails. This method mirrors the approach taken by the baseline Regex in [88], which also utilizes regular expressions to mitigate jailbreak attacks.

Transitioning from specific examples to an evaluation of their effectiveness, it is evident that while these methods benefit from straightforward and transparent pattern matching—attributes that contribute to their computational

efficiency and interpretability—their reliance on predefined patterns can be a significant drawback. Specifically, these rule-based systems may falter when encountering novel jailbreak techniques that deviate from known patterns which inherently limit their capability to combat more sophisticated attacks.

Summary 4: *Rule-based guardrails, while beneficial for their computational efficiency and ease of interpretation, face challenges when dealing with innovative jailbreak techniques that do not match existing predefined patterns.*

Model-based Guardrails. These guardrails adopt classifiers or statistical characteristics to distinguish between benign and harmful queries. Model-based approaches can capture more complex patterns than rule-based methods, enabling them to generalize better to novel jailbreak attempts. Learning a text-based classifier is a common approach for jailbreak detection. On one hand, we can use traditional machine learning models as the classifiers [73], [76], [79], [86], [87], [88], [89]. For instance, K-Nearest Neighbors (KNN) in RigorLLM [73], LightGBM in RSAA [76] and Random Forest in PretrainedEmbeddings [89]. On the other hand, neural networks are also widely applied for the safety classification [61], [62], [66], [70], [75], [80], [86], [88], [96]. For example, HSF [86] and CURVALID [96] use a simple Multilayer Perceptron (MLP) as the classifier. PromptGuard [80] and erase-and-check [66] fine-tune the pre-trained model (i.e., mDeBERTa and DistilBERT, respectively) to distinguish the safe and unsafe inputs. Besides, other methods used statistical characteristics to design their own algorithms on safety distinguish [13], [68], [71], [84], [85], [90], [95], [97], [98]. Detecting Perplexity [64] and Perplexity Filter [65] classify the input as a jailbreak request if the perplexity is higher than a threshold. Gradient discrepancies between safe prompts and adversarial prompts are employed in GradSafe [68], GradientCuff [71] and TokenHighlighter [90]. JailGuard [98] identifies the jailbroken state of responses by computing their KL-divergence. EEG-Defender [85], JBSHield [13] and MirrorShield [97] take the model’s internal feature similarities between the input and the jailbreak prompt as judgment basis.

The essence of model-based guardrails is to find a classification standard in distinguishing the benign and harmful requests, whether to learn a classifier or design a statistical algorithm. Compared with rule-based methods, model-based approaches can capture more complex patterns and generalize better to novel jailbreak attempts. They can also adapt to evolving threats by retraining or fine-tuning the classifiers. However, these methods typically require substantial training data and computational resources, especially when using deep learning models.

Summary 5: *Model-based guardrails, by adopting classifiers or statistical characteristics to distinguish between benign and harmful queries, can capture more complex patterns than rule-based methods, enabling them to generalize better to novel jailbreak attempts. However, these methods typically require substantial training data and computational resources, espe-*

cially when using deep learning models.

Observation: *Intra-processing guardrails are basically model-based guardrails, which use LLM’s internal features to detect potential jailbreak attempts. This is because model-based methods analyze the features and build classifiers instead of using simple character matching or a more complex LLM.*

LLM-based Guardrails. LLM-based guardrails represent a sophisticated approach to security, harnessing the inherent inferring capabilities of LLMs themselves to identify and counteract jailbreak attempts. Within this paradigm, research has progressed through distinct phases, each characterized by evolving methodologies.

Initially, methods tended to focus on directly determining the harmfulness of a request or providing a summary analysis after the judgment. For example, Self Defense [63] directly employs the target LLM to assess the safety of its own generated responses and subsequently furnish an explanation for its findings. In a similar vein, Llama Guard [10] operates by first identifying an unsafe text and then assigning it to a harmfulness category. Complementing these approaches, WildGuard [60] offers a multi-faceted assessment, simultaneously reporting the harmfulness status of the input prompt, the generated response, and whether the response was ultimately refused.

More recently, however, there has been a discernible shift towards methodologies that conduct a more detailed, upfront analysis before arriving at a safety judgment. Illustrating this trend, SelfDefend [11] first summarizes the underlying intention of the input and then assesses whether this intention constitutes a jailbreak request. Building upon this principle of preliminary in-depth analysis, both GuardReasoner [93] and X-Guard [99] employ chain-of-thought reasoning. This allows them to meticulously trace and analyze potential harms associated with a query, culminating in a final safety judgment.

Undeniably, the strength of these LLM-driven techniques lies in the excellent language understanding intrinsic to the models themselves. As a result, these approaches demonstrate considerable efficacy in detecting a diverse range of jailbreak attempts and notably improve the explainability of the safety judgments they provide. Nevertheless, a crucial trade-off exists. While effective and explainable, these advanced guardrails may introduce substantially more computational overhead when compared to rule-based and model-based techniques.

Summary 6: *Employing the reasoning capabilities of LLMs, LLM-based guardrails not only detect and mitigate jailbreak attempts effectively but also improve the explainability of safety judgments. Nonetheless, they introduce significantly greater computational overhead than traditional rule-based and model-based methods.*

We now present one RQ that focuses on the cost of LLM-based guardrails, which is a crucial aspect of their practical deployment. This RQ is particularly relevant given the increasing complexity and resource demands of LLM-

based approaches, especially in environments with limited computational resources.

RQ 2: *Given that rule-based, model-based, and LLM-based guardrails inherently possess different levels of computational complexity and resource requirements, a significant practical question emerges: To what extent does the choice of technical paradigm directly influence the GPU memory footprint of LLM guardrail mechanisms during their operational deployment?*

4.3. Safety Granularity

Guardrail mechanisms can operate at 3 different levels of detection granularity: token-level for individual words or tokens, sequence-level for an entire prompt or response, and session-level for entire conversation sessions.

Token-level Guardrails. These guardrails analyze individual tokens or small token groups to identify potentially harmful elements within inputs or outputs. Token-level approaches can pinpoint specific problematic components within a text, enabling more precise interventions. For instance, Token Highlighter [90] identifies specific tokens that contribute to harmful outputs. These fine-grained approaches enable targeted interventions but may miss harmful content that emerges from the broader context rather than specific tokens.

Sequence-level Guardrails. These guardrails evaluate entire prompts or responses as cohesive units, considering the overall semantic meaning rather than individual components. Sequence-level approaches can capture harmful content that emerges from the interaction between different parts of a text. For example, Llama Guard [10] and ShieldGemma [83] assess the holistic safety of the prompt sequence, while Constitutional Classifiers [94] evaluate outputs against predefined safety principles. These approaches can better capture contextual harms but may provide less granular insights into specific problematic elements.

Session-level Guardrails. These guardrails monitor entire conversation sessions, tracking the evolution of dialogue across multiple turns to identify potential jailbreak attempts that unfold gradually. Session-level approaches can detect sophisticated multi-turn attacks that might appear benign when individual messages are analyzed in isolation. For instance, Adaptive Guardrail [84] maintains conversation state to identify harmful patterns across turns. These comprehensive approaches are particularly valuable against advanced jailbreak techniques that exploit the sequential nature of conversations but typically require more complex implementation and greater computational resources. We now present two RQs that explore the impact of safety granularity on the effectiveness and utility of guardrail mechanisms.

RQ 3: *To what extent are current session-level guardrails truly effective in defending against sophisticated multi-turn jailbreak attacks?*

RQ 4: *How does the choice of safety granularity (i.e., token, sequence, or session-level) impact the utility of LLMs when implementing guardrail mechanisms?*

4.4. Reactiveness

Static Guardrails operate by analyzing inputs or outputs of LLMs without modifying them. For example, OpenAI Moderation [62] uses a fine-tuned model to classify input prompts into safety categories without altering the original query. In contrast, **Dynamic Guardrails** actively modify inputs or outputs to neutralize adversarial properties while preserving semantic meaning, which is beneficial for better identification of jailbreak attacks. The modification on the input prompt or the output response can be divided into character-level, token-level, word-level, and sentence-level changes.

Character-level changes involve randomly introducing and altering characters within the text input, such as typos, inserting new characters or character swaps, to disrupt adversarial prompts without changing the overall meaning [67], [69]. A prominent example is SmoothLLM [67], which perturbs multiple copies of the input prompt through character-level changes and uses a voting mechanism to determine safety. *Token-level* changes involve substituting, inserting, or deleting specific tokens to neutralize adversarial prompts. For instance, erase-and-check [66] removes potentially harmful tokens from the input prompt to construct multiple sanitized versions for safety voting like SmoothLLM. *Word-level* changes involve substituting words with synonyms or related terms to neutralize adversarial prompts [69], [98]. For instance, SemanticSmooth [69] converts all verbs into the past tense and substitutes both verbs & nouns with their semantic equivalents. *Sentence-level* changes involve modifying and rewriting the entire input query to expose the embedded attack intent [9], [69], [73], [88], [98], [99]. For example, NeMo Guardrails [9] encodes the input prompt into a structured format to clarify the request's intent. RigorLLM [73] paraphrases or summarizes the input to identify potential jailbreak attempts. X-Guard [99] translates the non-English input into English. Naturally, textual modifications can be applied simultaneously at multiple levels of granularity, including character, word, and sentence levels, such as in SemanticSmooth [69] and JailGuard [98]. From a broader perspective, dynamic guardrails can be viewed as a form of input/output augmentation that enhances the robustness of jailbreak detection by exposing hidden attack intents.

4.5. Applicability

The applicability of a guardrail is largely determined by its dependency on model internals. **White-box guardrails** require access to the target LLM's internal states, such as gradients or hidden activations, to detect jailbreak attempts. For instance, GradSafe [68] analyzes gradient patterns of safety-critical parameters, while JBShield [13] compares internal state divergences between benign and malicious queries. These methods can achieve high detection accuracy but are limited to scenarios where model internals are accessible. In contrast, **black-box guardrails** operate solely on input and output text, making them suitable for

TABLE 2. THE DETAILS OF OUR COLLECTED BENCHMARK DATASETS.

Dataset	# Prompts	Jailbreak Methods
JailbreakHub [17]	1000	IJP [17]
JailbreakBench [101]	100	GCG [7], AutoDAN [8]
		TAP [26], LLM-Fuzzer [30]
		DrAttack [32]
		X-Teaming [47]
MultiJail [34]	315	MultiJail
SafeMTData [46]	600	ActorAttack [46]
AlpacaEval [102]	805	Normal Prompts
OR-Bench [103]	1000	Normal Prompts

proprietary or API-based LLMs. Examples include Prompt Guard [80], which uses a lightweight classifier to filter malicious inputs, and WildGuard [60], which moderates both prompts and responses via a separate LLM. Black-box approaches offer broader deployment flexibility but may sacrifice granularity in detection. This dimension is critical for practical deployment: organizations using closed-source models like GPT-4 must rely on black-box guardrails, while those with custom LLMs can leverage white-box methods for finer control.

4.6. Explainability

Explainability refers to the ability of a guardrail to provide transparent, interpretable rationales for its safety judgments. **Opaque guardrails**, such as many model-based classifiers, output binary decisions without justification, which can hinder trust and debugging. For example, MoJE [87] employs simple linguistic statistical techniques to classify inputs as safe or unsafe without explaining its reasoning. While opaque guardrails can be efficient and effective, their lack of transparency may limit adoption in high-stakes applications where understanding failure modes is crucial. In contrast, **explainable guardrails** enhance usability by providing reasoning traces or confidence scores. GuardReasoner [93] employs chain-of-thought reasoning to step through potential harms before rendering a judgment, while SelfDefend [11] summarizes the intent of a query before assessing its risk. Such transparency helps developers understand failure modes, refine safety policies, and comply with regulatory requirements.

5. Benchmark & Leaderboard

5.1. Evaluation Setup

Datasets and Target Models. Based on the five categories of existing jailbreak attacks we surveyed in §2 — manual, optimization-based, generation-based, implicit, and multi-turn jailbreaks — we identify representative jailbreak attack methods in each category. We then collect six benchmark datasets, **JailbreakHub** [17], **JailbreakBench** [101], **SafeMTData** [46], **MultiJail** [34], **AlpacaEval** [102] and **OR-Bench** [103], from which we use their user prompts to build test queries for diverse guardrails. Table 2 lists the

details of our collected benchmark datasets. *JailbreakHub* is a framework that collects and categorizes wild jailbreak prompts designed to bypass safety restrictions in LLMs. To refine ASR/PGR to three decimal places, we randomly sample 1,000 in-the-wild prompts (IJP) from JailbreakHub as manual attacks. *JailbreakBench* is an open-source robustness benchmark specifically designed to evaluate and measure the vulnerability of LLMs to jailbreak attacks. We use all 100 harmful instructions from JailbreakBench to drive optimization-based jailbreaks (GCG [7] and AutoDAN [8]), generation-based jailbreaks (TAP [26] and LLM-Fuzzer [30]), implicit jailbreak attack DrAttack [32], and multi-turn jailbreak attack X-Teaming [47]. *SafeMTData* is a dataset contains the initial multi-turn jailbreak prompts produced by ActorAttack [46]. We choose all 600 queries from SafeMTData as multi-turn jailbreak attacks. *MultiJail* is the first manually-created multilingual jailbreak dataset, covering both high-resource and low-resource languages. We choose all 315 jailbreak prompts in Bengali from MultiJail as multilingual jailbreaks. *AlpacaEval* is an automatic evaluation framework designed to assess the performance of instruction-following LLMs. We use all 805 instructions from AlpacaEval as normal prompts. *OR-Bench* is the first large-scale benchmark measuring LLMs’ over-refusal of 80,000 seemingly toxic but benign prompts across multiple categories. We randomly select 1,000 prompts from OR-Bench as normal prompts to also calculate ASR/PGR to three decimal places. We evaluate the defense performance of guardrails on open-source LLMs (i.e., **Llama-3-8B-Instruct** [4] and **Vicuna-13b-v1.5** [2]) and a closed-source LLM (i.e., **GPT-4-0125-Preview** [104]), covering diverse model architectures and sizes. For evaluations on Vicuna-13b-v1.5 and GPT-4-0125-Preview, please refer to Appendix A and Appendix B, respectively.

Metrics. We measure the defense performance of guardrails by the attack success rate (ASR), i.e., the frequency with which jailbreak prompts in a benchmark dataset bypass the guardrail and the target LLM. The lower the ASR, the stronger the defense performance. Adopting the methodology from prior studies ([46], [47], [105]), we leverage GPT-4o’s robust evaluation capabilities for assessment. This involves supplying GPT-4o, acting as a judge, with harmful objectives and the corresponding model outputs. The judge assigns a score from 1 to 5, where higher scores denote more harmful and detailed model responses. An attack is deemed successful if it receives a score of 5 from the GPT-4o Judge. For a detailed explanation of the scoring rubric, please see [46], [105].

Attack Configuration. To assess the jailbreak defense performance of guardrails, we employ the most widely used jailbreak attacks, including a manual attack (**IJP** [17]), optimization-based attacks (**GCG** [7] and **AutoDAN** [8]), generation-based attacks (**TAP** [26] and **LLM-Fuzzer** [30]), implicit attacks (**DrAttack** [32] and **MultiJail** [34]), and multi-turn attacks (**ActorAttack** [46] and **X-Teaming** [47]). In the context of *IJP*, 1,000 adversarial queries were randomly sampled from the forbidden question set with jailbreak prompts [106], curated by JailbreakHub. Regarding

TABLE 3. THE ASR ($\downarrow$) / PGR ($\downarrow$) RESULTS FOR THE TARGET LLM (LLAMA-3-8B-INSTRUCT) WITH DIFFERENT GUARDRAILS AGAINST FIVE MAJOR CATEGORIES OF JAILBREAK ATTACKS, INCLUDING ROW AVERAGES. (PRE) AND (POST) DENOTE THE PRE-PROCESSING AND POST-PROCESSING VERSIONS OF THE GUARDRAILS, RESPECTIVELY. (DIRECT) AND (INTENT) DENOTE THE DIRECT PROMPT AND INTENT PROMPT BASED VERSIONS OF SELFDEFEND [11], RESPECTIVELY.

Guardrails	Manual	Optimization-based		Generation-based		Implicit		Multi-turn		Average
	IJP	GCG	AutoDAN	TAP	LLM-Fuzzer	DrAttack	MultiJail	ActorAttack	X-Teaming	
Llama-3-8B-Instruct	0.078/-	0.130/-	0.020/-	0.140/-	0.490/-	0.100/-	0.044/-	0.227/-	0.910/-	0.238/-
PerplexityFilter	0.078/1.000	0.100/0.620	0.020/1.000	0.140/1.000	0.480/1.000	0.100/1.000	0.044/1.000	0.227/1.000	0.960/1.000	0.239/0.958
SmoothLLM	0.115/0.261	0.020/0.020	0.030/0.110	0.140/0.170	0.500/0.810	0.150/0.660	0.032/0.575	0.893/0.893	0.850/0.910	0.303/0.490
Llama Guard (Pre)	0.062/0.563	0.100/0.390	0.020/0.480	0.140/0.460	0.450/0.680	0.100/0.840	0.044/0.952	0.220/0.967	0.910/1.000	0.227/0.704
Llama Guard (Post)	0.061/0.061	0.090/0.090	0.020/0.020	0.150/0.150	0.390/0.390	0.090/0.090	0.041/0.041	0.223/0.223	0.960/0.960	0.225/0.225
GradSafe	0.077/0.599	0.130/0.770	0.010/0.040	0.070/0.580	0.450/0.580	0.100/0.420	0.044/0.917	0.188/0.863	0.950/0.990	0.224/0.640
GradientCuff	0.016/0.058	0.080/0.140	0.000/0.020	0.070/0.090	0.360/0.480	0.030/0.130	0.016/0.149	0.118/0.648	0.640/0.710	0.148/0.269
SelfDefend (Direct)	0.020/0.267	0.030/0.080	0.010/0.120	0.080/0.180	0.090/0.140	0.070/0.590	0.038/0.752	0.133/0.702	0.970/0.990	0.160/0.425
SelfDefend (Intent)	0.022/0.285	0.030/0.070	0.010/0.130	0.120/0.180	0.030/0.130	0.010/0.130	0.032/0.584	0.152/0.767	0.940/0.970	0.150/0.361
WildGuard (Pre)	0.004/0.033	0.020/0.020	0.000/0.020	0.030/0.060	0.000/0.000	0.080/0.500	0.044/0.797	0.150/0.757	0.960/0.980	0.143/0.352
WildGuard (Post)	0.020/0.020	0.050/0.050	0.010/0.010	0.060/0.060	0.090/0.090	0.060/0.060	0.035/0.035	0.148/0.148	0.950/0.950	0.158/0.158
Prompt Guard	0.000/0.000	0.000/0.080	0.020/0.420	0.170/0.940	0.000/0.000	0.100/0.940	0.044/1.000	0.225/0.995	0.910/1.000	0.163/0.597
GuardReasoner (Pre)	0.000/0.009	0.000/0.000	0.000/0.010	0.030/0.070	0.010/0.020	0.080/0.360	0.029/0.349	0.143/0.740	0.920/0.960	0.135/0.280
GuardReasoner (Post)	0.023/0.023	0.040/0.040	0.000/0.000	0.050/0.050	0.050/0.050	0.030/0.030	0.022/0.022	0.107/0.107	0.950/0.950	0.141/0.141

GCG, its individual variant was selected, and the adversarial suffix was optimized against the target LLM employing a batch size of 512 and subjected to 500 optimization iterations. For the *AutoDAN* methodology, the hierarchically-guided genetic algorithm variant, specifically *AutoDAN-HGA*, was adopted. The genetic algorithm integral to *AutoDAN-HGA* operates with a crossover probability of 0.5, a mutation probability of 0.01, and undergoes 500 optimization iterations. For GCG and *AutoDAN*, we migrate jailbreak prompts optimized for Vicuna-13b to attack GPT-4. Concerning *TAP*, the Vicuna-13b-v1.5 model [2] was utilized as the attacking agent. The parameters for *TAP* were configured with a maximum depth of 5, a maximum width of 5, and a branching factor of 4. The designated target models for *TAP* included Llama-3-8B-Instruct [4], Vicuna-13b-v1.5 [2] or GPT-4-0125-Preview [104]. In the case of *LLM-Fuzzer*, GPT-3.5 served as the auxiliary model for generating mutational inputs, and the query limit directed at the target LLMs was established at 200. For *DrAttack*, jailbreak prompts were formulated using GPT-4o. With respect to *MultiJail*, the entirety of the 315 available queries in the Bengali language was selected. For the *ActorAttack* strategy, a corpus of 600 queries was sourced from the SafeMTData dataset [46] (specifically, the `SafeMTData/Attack_600.json` file available on Hugging Face). For *X-Teaming*, we set the attacking model as Qwen2.5-32B-Instruct [6] and use the TextGrad-based text optimization to refine jailbreak prompts. Regarding *AlpacaEval*, all 805 questions within the AlpacaEval dataset were utilized. For *OR-Bench*, a subset of 1,000 prompts was randomly selected from the OR-Bench dataset [103] (specifically, the `or-bench-80k.csv` file on Hugging Face).

It is pertinent to note that: *The prompts associated with IJP, MultiJail, ActorAttack, AlpacaEval, and OR-Bench are static in nature. Consequently, all guardrail mechanisms encounter identical input stimuli, irrespective of whether they are safeguarding Llama-3-8B, Vicuna-13b or GPT-4. In contrast, GCG, AutoDAN, and DrAttack are specifically*

tailored to either Llama-3-8B, Vicuna-13b or GPT-4. As such, guardrails receive uniform inputs when defending the same designated target LLM. Conversely, TAP, LLM-Fuzzer, and X-Teaming represent adaptive attack methodologies. This implies that guardrail systems are presented with varied inputs, even when applied to the identical target LLM.

Baselines. We evaluate our framework with popular jailbreak defense methods, including **Perplexity Filter** [65], **SmoothLLM** [67], **Llama Guard** [10], **GradSafe** [68], **GradientCuff** [71], **SelfDefend** [11], **WildGuard** [60], **Prompt Guard** [80], and **GuardReasoner** [93]. Specifically, *Perplexity Filter* leverages a Llama-2-7b model to calculate the perplexity of the input prompt. A jailbreak is considered to happen when the perplexity exceeds a threshold. We set this threshold at the maximum perplexity of any prompt in the JailbreakBench dataset of harmful behavior prompts. *SmoothLLM* perturbs the jailbreak prompts with character-level changes to enable the target LLM to perform defense. In this paper, we set *SmoothLLM* to conduct character swapping with a 10% perturbation percentage. *Llama Guard* is a fine-tuned Llama-2-7b model designed to detect the toxicity category of input prompts. *GradSafe* is a gradient-based detection method that identifies unsafe or jailbreak prompts in LLMs by analyzing the consistent gradient patterns of safety-critical parameters when paired with compliance responses. *GradientCuff* is a method for detecting jailbreak attacks on LLMs by analyzing the refusal loss landscape, leveraging gradient-based patterns to identify and block adversarial prompts while maintaining normal query performance. *SelfDefend* is a practical jailbreak defense framework for LLMs that uses a shadow LLM instance to concurrently detect harmful queries while the target LLM processes them, providing robust protection with minimal delay. *WildGuard* is an open, lightweight, multi-task moderation tool for LLMs that detects malicious user prompts, harmful model responses, and model refusal behavior. *Prompt Guard* is a security tool developed by Meta that detects and blocks malicious inputs (e.g., jailbreak attempts, prompt injections) in LLM applications,

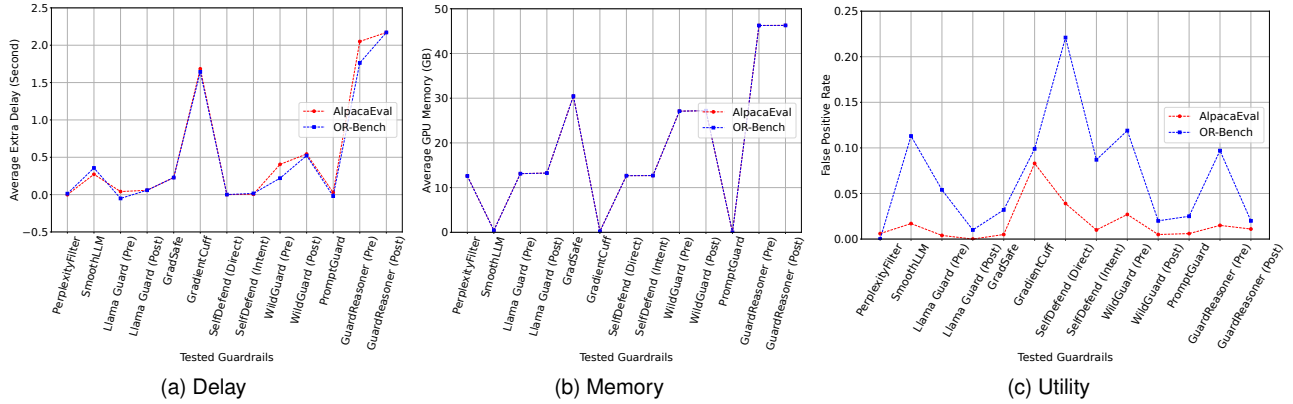


Figure 2. The delay, memory usage, and utility of guardrails.

using a lightweight classifier model Prompt-Guard-86M to filter harmful content in real time. *GuardReasoner* is a reasoning-based guard model designed to enhance the safety of LLMs by integrating explicit step-by-step reasoning into the moderation process. Our evaluations are implemented using PyTorch 2.6.0 and conducted on NVIDIA Hopper H800 GPUs.

5.2. Benchmark Evaluation

Defense Performance. We first analyze the defense performance of various guardrails. As delineated in Table 3, which presents the ASR, a lower value indicates superior defense capabilities. On average, *GuardReasoner* (Pre) demonstrates the most robust defense, achieving the lowest ASR of 0.135. Following closely is *GuardReasoner* (Post), underscoring the efficacy of the reasoning process prior to safety determination inherent in the *GuardReasoner* framework. Conversely, *SmoothLLM* exhibits the highest ASR of 0.303, rendering it the least effective in this cohort. This suboptimal performance may be attributed to its mechanism of token-level input perturbation, which appears to be primarily effective against jailbreak techniques characterized by adversarial suffixes, such as GCG, while offering limited protection against a broader spectrum of attacks.

Shifting focus to PGR, presented in Table 3, *GuardReasoner* (Post) achieves the best PGR of 0.141. Despite its superior precision in identifying malicious inputs, *GuardReasoner* (Post) does not attain state-of-the-art (SOTA) overall defense performance. A plausible explanation is its potential operational overlap with the target LLM’s intrinsic safety mechanisms. That is, there might be a significant number of instances where *GuardReasoner* (Post) identifies a response as safe, and concurrently, the target LLM also recognizes the harmful nature of the query and refuses to respond. Therefore, this overlap diminishes the unique contribution of *GuardReasoner* (Post) to the ASR reduction when compared to a guardrail like *GuardReasoner* (Pre) which operates on a different paradigm.

Efficiency. The efficiency of guardrails is a critical factor for practical deployment, which we assess in terms of latency and GPU memory consumption. Figure 2(a) illustrates the

extra delay introduced by different guardrail methodologies when processing normal inputs from the AlpacaEval and OR-Bench datasets. Perplexity Filter, Llama Guard, SelfDefend and PromptGuard stand out with negligible latency. In contrast, *GuardReasoner* and *GradientCuff* impose the most significant delays, with *GuardReasoner* (Post) being particularly notable. This suggests that the profound reasoning capabilities that afford *GuardReasoner* its enhanced defense performance come at the cost of increased processing time. The majority of other guardrails maintain an additional delay generally not exceeding 0.5 seconds.

From the perspective of GPU memory utilization, depicted in Figure 2(b), *GuardReasoner* again registers the highest memory footprint, consistent with its complex reasoning architecture. Conversely, *SmoothLLM*, *GradientCuff*, and *PromptGuard* are the most memory-efficient, with their consumption approaching negligible levels. This highlights a clear trade-off between the sophistication of the defense mechanism and its resource intensiveness.

Utility. Beyond security and efficiency, the utility of a guardrail, specifically its ability not to impede benign user interactions, is paramount. We measure the FPR on AlpacaEval and OR-Bench datasets, as shown in Figure 2(c). A higher FPR indicates a greater propensity to incorrectly flag legitimate prompts as malicious. *SelfDefend* (Direct) exhibits the highest FPR on OR-Bench, at 0.221. On AlpacaEval, *GradientCuff* records the highest FPR of 0.083. These figures suggest that these guardrails have a higher likelihood of intercepting normal user queries. Other guardrails with comparatively high FPRs include *SmoothLLM*, *SelfDefend* (Intent), *WildGuard* (Pre), and *GuardReasoner* (Pre). Although these three methods demonstrate strong defense performance (low ASR), their elevated FPRs underscore a critical trade-off between security and utility. Systems that are highly stringent in blocking threats may inadvertently penalize legitimate interactions, diminishing the overall user experience.

5.3. Leaderboard on SEU

To provide a holistic evaluation, we compare guardrails across five key metrics: ASR, PGR, Extra Delay, GPU

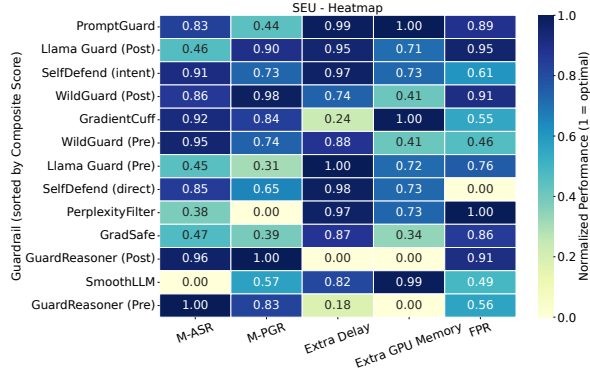


Figure 3. The heatmap of guardrails.

Memory, and FPR. We average ASR and PGR over the nine jailbreak attacks (cf. Table 3) to derive Mean-ASR (M-ASR) and Mean-PGR (M-PGR). The other metrics are measured using the OR-Bench dataset. For a unified ranking, we normalize each metric to a $[0, 1]$ range and invert the scores ($1 - \text{normalized value}$), ensuring higher values consistently indicate better performance. We then compute a Composite Score for each guardrail by averaging these five transformed scores. This score underpins the ranking visualized in the heatmap in Figure 3.

The analysis reveals inherent trade-offs, as no single guardrail excels across all dimensions. For instance, PromptGuard achieves the highest Composite Score but its low M-PGR suggests potential gaps in detection robustness. Conversely, GuardReasoner (Pre) ranks lower but provides superior defense (high M-ASR and M-PGR) at a significant cost to efficiency and utility. SelfDefend (Intent) offers a balanced profile, with its main weakness being a higher FPR. This leaderboard underscores that the optimal guardrail choice is context-dependent, contingent on the specific security requirements and operational constraints of a given deployment scenario. We believe this leaderboard will serve as a valuable resource for practitioners in selecting appropriate guardrails based on their unique needs.

6. Practical Insights & Implications

Answer to RQ3: Session-level Guardrails v.s. Multi-turn Jailbreaks. A critical question arises regarding session-level guardrails: given their reliance on LLM dialogue history (both input and output) for threat assessment, how effectively do they counter sophisticated multi-turn jailbreak attacks? Our analysis, focusing on three session-level guardrails—Llama Guard (Post), WildGuard (Post), and GuardReasoner (Post)—reveals nuanced performance. As indicated in Table 3, these guardrails maintain an ASR above 10% against the ActorAttack multi-turn jailbreak. Furthermore, when faced with the more adaptive X-Teaming attack, the ASR for most guardrails, including these session-level ones, exceeds 90%. GradientCuff is a partial exception with a 64% ASR, but this is still a high failure rate. These findings underscore a significant vulnerability of current

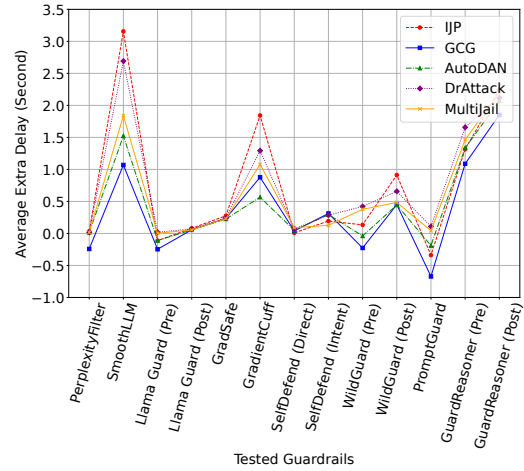


Figure 4. The delay of guardrails against different attack types.

session-level guardrails against advanced multi-turn attacks. The high ASR, particularly against adaptive attacks like X-Teaming, suggests that these defenses can be readily bypassed if the attack unfolds over several interactions. This highlights an urgent imperative to develop more robust guardrail methodologies specifically designed to address the evolving landscape of multi-turn jailbreaks.

Answer to RQ1: Intervention Stages on Delay. The intervention stage of a guardrail—whether it operates pre-processing (on user input), intra-processing (during LLM generation), or post-processing (on LLM output)—can significantly impact system latency. We investigate this relationship by examining the data presented in Figure 4. Observations indicate that, with the notable exception of GuardReasoner (Pre), pre-processing guardrails such as Perplexity Filter, Llama Guard (Pre), SelfDefend (Direct), SelfDefend (Intent), WildGuard (Pre), and Prompt Guard generally introduce negligible, or in some cases even negative, additional latency. The higher latency of GuardReasoner (Pre) is attributable to its more complex reasoning processes. In contrast, intra-processing and post-processing guardrails exhibit more varied latency profiles relative to each other. A key finding is that for identical detection models, post-processing variants consistently incur greater delay than their pre-processing counterparts (e.g., WildGuard (Post)’s delay is greater than that of WildGuard (Pre)). This phenomenon arises because post-processing methods inherently must await the completion of the target LLM’s generation phase before they can intervene. Conversely, pre-processing guardrails possess the advantage of potentially halting the LLM’s generation process immediately upon detecting a malicious input, thereby conserving computational time. Consequently, pre-processing guardrails, particularly those not reliant on extensive reasoning, generally offer a more latency-efficient solution for integrating safety measures.

Answer to RQ2: Technical Paradigms on GPU Memory Usage. The underlying technical paradigm of a guardrail—be it rule-based, traditional model-based, or LLM-based—is expected to influence its GPU memory foot-

print. We examine this correlation using data from Figure 2. The results show that the rule-based SmoothLLM incurs zero additional memory overhead, representing the most memory-efficient approach. Certain traditional model-based methods, specifically GradientCuff and PromptGuard, also demonstrate near-zero memory consumption, highlighting their lightweight nature. However, the landscape for model-based approaches is not uniform; GradSafe, another model-based technique, exhibits higher memory usage than several LLM-based methods, indicating significant variability in resource demands even within this category. As anticipated, LLM-based guardrails generally impose a greater memory burden. This is an intrinsic consequence of their design, which necessitates loading and executing a large language model for safety inference. This observation aligns with the expectation that leveraging large language models for safety assessment incurs a higher resource cost in terms of memory. While rule-based and optimized model-based solutions offer substantial memory efficiency, the choice of paradigm must be carefully weighed against the desired detection capabilities and specific deployment constraints.

Answer to RQ4: Safety Granularity on Utility. The granularity at which a guardrail performs its safety checks—whether at the token-level, sequence-level (assessing the entire input or output), or session-level (considering the dialogue history)—may significantly affect its utility, particularly its propensity to misclassify benign prompts, as measured by the FPR. This aspect is explored using data from Figure 2(c). Token-level guardrails, exemplified by SmoothLLM (which analyzes keywords in LLM responses) and SelfDefend (Direct) (which inspects harmful segments within queries), demonstrate relatively pronounced FPRs. Notably, SelfDefend (Direct) records the highest FPR on the OR-Bench dataset, exceeding 20%. This suggests that token-level mechanisms, while focused, may inadvertently penalize legitimate interactions due to a potential lack of broader contextual understanding. A comparative analysis further reveals that for the same underlying detection model, session-level guardrails (typically denoted by a “(Post)” suffix, leveraging both LLM input and output) consistently achieve markedly lower FPRs than their sequence-level counterparts (often denoted by a “(Pre)” suffix, relying solely on input). For instance, WildGuard (Pre) exhibits an FPR above 10% on OR-Bench, whereas the FPR for WildGuard (Post) remains below 5%. While sequence-level guardrails display a wider range of FPRs—some high, some low—session-level approaches generally maintain low FPR values across the board. These observations collectively suggest that session-level guardrails tend to offer superior utility by minimizing false positives. This improved performance is likely attributable to their comprehensive use of contextual information derived from the entire interaction history, enabling a more nuanced distinction between genuinely harmful prompts and benign ones that might share superficial characteristics with attacks.

Generalization: Cross-Attack Assessment. While LLM-based guardrails have demonstrated efficacy against jailbreak attacks, a critical and often overlooked considera-

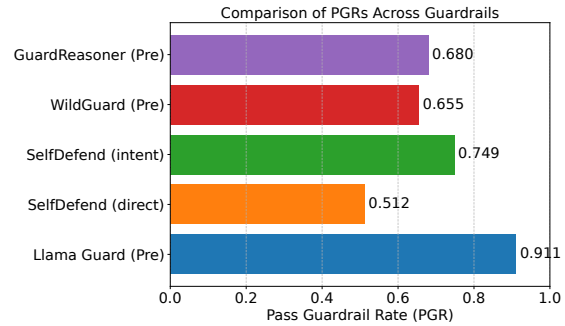


Figure 5. Cross-attack evaluation: injection attack on guardrails.

tion is their robustness against other adversarial manipulations, specifically prompt injection attacks. Given that these guardrails are themselves powered by LLMs, their susceptibility to injection attacks—which could potentially subvert their safety assessment capabilities—presents a significant security concern. To investigate this, we evaluated the performance of LLM-based guardrails against 203 distinct injection attack samples sourced from the “deepset/prompt-injections” dataset on Hugging Face.

Our primary finding is that these injection attacks did not compromise the fundamental operational integrity of the guardrails. That is, the guardrails were not coerced into abandoning their safety analysis function to produce arbitrary, irrelevant outputs (e.g., “hello world”). They continued to process the inputs for security threats as designed. However, their effectiveness in identifying and mitigating these injections was limited. We measured the Pass Guardrail Rate (PGR) for these attacks, with results presented in Figure 5. The data reveals that while LLM-based guardrails exhibit a non-trivial capacity to filter prompt injections, this capability is modest at best. This assessment underscores a crucial gap in the current state of guardrail technology: the need for broader cross-attack generalization. For a guardrail to be truly effective in practice, its defensive perimeter must extend beyond jailbreak attempts to also detect and neutralize other forms of attacks that could exploit its defense, such as prompt injections. This calls for the development of more versatile and robust guardrail mechanisms capable of addressing a wider spectrum of adversarial inputs.

7. Discussion

7.1. Limitations

While this work provides a comprehensive analysis of LLM jailbreak guardrails, several limitations should be acknowledged. First, our efficiency evaluations were conducted on specific hardware configurations (NVIDIA H800 GPUs), and the absolute latency and memory consumption metrics may vary across different hardware platforms and optimization techniques such as quantization. However, the relative performance trends and trade-offs among guardrails are expected to remain consistent.

Besides, our evaluation relies on LLM-as-a-judge (GPT-4o) for assessing jailbreak success, which, while efficient for large-scale evaluation, may have inherent limitations in reliability. To validate this approach, we hired three Ph.D. students to evaluate the jailbreak success of 100 randomly selected attack samples. Their average agreement with the LLM’s judgments was 0.9533, and their Fleiss’ Kappa was 0.9319, indicating strong alignment. Nevertheless, future work could benefit from more robust evaluation methodologies that mitigate potential biases in automated assessment.

7.2. Future Work

Based on our analysis, we identify several promising directions for future research in LLM jailbreak guardrails:

- **Cross-Attack Robust Guardrails:** Current guardrails demonstrate limited generalization across different attack modalities. Future work should develop unified defense mechanisms that can effectively handle diverse threats, including jailbreaks, prompt injections, and other adversarial manipulations simultaneously. This requires guardrails that can recognize attack patterns beyond jailbreaks.
- **Adaptive and Self-Evolving Defenses:** As jailbreak techniques continuously evolve, fixed guardrails face the risk of rapid obsolescence. Research is needed on self-adaptive guardrails that can learn from new attack patterns and update their detection capabilities without complete retraining. This could include online learning approaches, anomaly detection in deployment environments, and automated defense refinement.
- **Multimodal and Cross-Modal Guardrails:** With the rise of multimodal LLMs, new vulnerabilities emerge at the intersection of different modalities. Future research should address the unique challenges of multimodal jailbreaks [107] and develop guardrails that can analyze and protect across text, image, audio, and video inputs while maintaining efficiency.

8. Conclusion

This SoK paper comprehensively addresses the fragmented landscape of LLM jailbreak guardrails by introducing a novel multi-dimensional taxonomy and a SEU measurement framework. Our findings highlight the strengths, limitations, and interdependencies of existing defense mechanisms with a series of key insights. This work forms a structured foundation to guide the principled advancement and deployment of more robust LLM guardrails.

Acknowledgments

We thank the reviewers and the shepherd for their constructive comments. The HKUST authors are supported in part by an RGC CRF grant under contract C6015-23G and a research fund provided by HSBC. Daoyuan Wu was partially supported by Lingnan Grant SUG-002/2526.

References

- [1] H. Touvron, L. Martin, K. Stone, P. Albert, A. Almahairi, Y. Babaei, N. Bashlykov, S. Batra, P. Bhargava, S. Bhosale *et al.*, “Llama 2: Open foundation and fine-tuned chat models,” *arXiv preprint arXiv:2307.09288*, 2023.
- [2] W.-L. Chiang, Z. Li, Z. Lin, Y. Sheng, Z. Wu, H. Zhang, L. Zheng, S. Zhuang, Y. Zhuang, J. E. Gonzalez, I. Stoica, and E. P. Xing, “Vicuna: An open-source chatbot impressing GPT-4 with 90%* ChatGPT quality,” March 2023.
- [3] T. M. A. Team, “Mistral-7b-instruct-v0.2,” <https://huggingface.co/mistralai/Mistral-7B-Instruct-v0.2>, 2024.
- [4] A. Grattafiori, A. Dubey, A. Jauhri, A. Pandey, A. Kadian, A. Al-Dahle, A. Letman, A. Mathur, A. Schelten, A. Vaughan *et al.*, “The llama 3 herd of models,” *arXiv preprint arXiv:2407.21783*, 2024.
- [5] Anthropic, “Claude 3.5 sonnet,” <https://www.anthropic.com/news/claude-3-5-sonnet>, 2024.
- [6] Qwen Team, “Qwen2.5 technical report,” *arXiv preprint arXiv:2412.15115*, 2024.
- [7] A. Zou, Z. Wang, J. Z. Kolter, and M. Fredrikson, “Universal and transferable adversarial attacks on aligned language models,” *arXiv preprint arXiv:2307.15043*, 2023.
- [8] X. Liu, N. Xu, M. Chen, and C. Xiao, “AutoDAN: Generating stealthy jailbreak prompts on aligned large language models,” in *ICLR*, 2024.
- [9] T. Rebedea, R. Dinu, M. N. Sreedhar, C. Parisien, and J. Cohen, “NeMo Guardrails: A toolkit for controllable and safe LLM applications with programmable rails,” in *EMNLP: System Demonstrations*, 2023.
- [10] H. Inan, K. Upasani, J. Chi, R. Rungta, K. Iyer, Y. Mao, M. Tontchev, Q. Hu, B. Fuller, D. Testuggine *et al.*, “Llama Guard: LLM-based input-output safeguard for Human-AI conversations,” *arXiv preprint arXiv:2312.06674*, 2023.
- [11] X. Wang, D. Wu, Z. Ji, Z. Li, P. Ma, S. Wang, Y. Li, Y. Liu, N. Liu, and J. Rahmel, “SelfDefend: LLMs can defend themselves against jailbreaking in a practical manner,” in *USENIX Security*, 2025.
- [12] S. Khonneux, A. Sordoni, S. Günnemann, G. Gidel, and L. Schwinn, “Efficient adversarial training in LLMs with continuous attacks,” in *NeurIPS*, 2024.
- [13] S. Zhang, Y. Zhai, K. Guo, H. Hu, S. Guo, Z. Fang, L. Zhao, C. Shen, C. Wang, and Q. Wang, “JBSHield: Defending large language models from jailbreak attacks through activated concept analysis and manipulation,” in *USENIX Security*, 2025.
- [14] Y. Liu, G. Deng, Z. Xu, Y. Li, Y. Zheng, Y. Zhang, L. Zhao, T. Zhang, and Y. Liu, “Jailbreaking ChatGPT via prompt engineering: An empirical study,” *arXiv preprint arXiv:2305.13860*, 2023.
- [15] A. Wei, N. Haghtalab, and J. Steinhardt, “Jailbroken: How does LLM safety training fail?” in *NeurIPS*, 2023.
- [16] Z. Wei, Y. Wang, and Y. Wang, “Jailbreak and guard: Aligned language models with only few in-context demonstrations,” *arXiv preprint arXiv:2310.06387*, 2023.
- [17] X. Shen, Z. Chen, M. Backes, Y. Shen, and Y. Zhang, ““Do Anything Now”: Characterizing and evaluating in-the-wild jailbreak prompts on large language models,” in *CCS*, 2024.
- [18] G. Deng, Y. Liu, Y. Li, K. Wang, Y. Zhang, Z. Li, H. Wang, T. Zhang, and Y. Liu, “MASTERKEY: Automated jailbreaking of large language model chatbots,” in *NDSS*, 2024.
- [19] C. Sitawarin, N. Mu, D. Wagner, and A. Araujo, “PAL: Proxy-guided black-box attack on large language models,” *arXiv preprint arXiv:2402.09674*, 2024.
- [20] M. Andriushchenko, F. Croce, and N. Flammarion, “Jailbreaking leading safety-aligned LLMs with simple adaptive attacks,” in *ICLR*, 2025.

- [21] X. Jia, T. Pang, C. Du, Y. Huang, J. Gu, Y. Liu, X. Cao, and M. Lin, "Improved techniques for optimization-based jailbreaking on large language models," in *ICLR*, 2025.
- [22] J. Hughes, S. Price, A. Lynch, R. Schaeffer, F. Barez, S. Koyejo, H. Sleight, E. Jones, E. Perez, and M. Sharma, "Best-of-N jailbreaking," *arXiv preprint arXiv:2412.03556*, 2024.
- [23] X. Chen, Y. Nie, W. Guo, and X. Zhang, "When LLM meets DRL: Advancing jailbreaking efficiency via DRL-guided search," in *NeurIPS*, 2024.
- [24] E. Perez, S. Huang, H. F. Song, T. Cai, R. Ring, J. Aslanides, A. Glaese, N. McAleese, and G. Irving, "Red teaming language models with language models," in *EMNLP*, 2022.
- [25] P. Chao, A. Robey, E. Dobriban, H. Hassani, G. J. Pappas, and E. Wong, "Jailbreaking black box large language models in twenty queries," *arXiv preprint arXiv:2310.08419*, 2023.
- [26] A. Mehrotra, M. Zampetakis, P. Kassinik, B. Nelson, H. Anderson, Y. Singer, and A. Karbasi, "Tree of attacks: Jailbreaking black-box LLMs automatically," in *NeurIPS*, 2024.
- [27] A. Paulus, A. Zharmagambetov, C. Guo, B. Amos, and Y. Tian, "AdvPrompter: Fast adaptive adversarial prompting for LLMs," in *ICML*, 2025.
- [28] Z. Liao and H. Sun, "AmpleGCG: Learning a universal and transferable generative model of adversarial suffixes for jailbreaking both open and closed LLMs," in *COLM*, 2024.
- [29] S. Yao, D. Yu, J. Zhao, I. Shafran, T. Griffiths, Y. Cao, and K. Narasimhan, "Tree of thoughts: Deliberate problem solving with large language models," in *NeurIPS*, 2024.
- [30] J. Yu, X. Lin, Z. Yu, and X. Xing, "LLM-Fuzzer: Scaling assessment of large language model jailbreaks," in *USENIX Security*, 2024.
- [31] D. Handa, A. Chirmule, B. Gajera, and C. Baral, "Jailbreaking proprietary large language models using word substitution cipher," *arXiv preprint arXiv:2402.10601*, 2024.
- [32] X. Li, R. Wang, M. Cheng, T. Zhou, and C.-J. Hsieh, "DrAttack: Prompt decomposition and reconstruction makes powerful LLMs jailbreakers," in *EMNLP Findings*, 2024.
- [33] Z. Chang, M. Li, Y. Liu, J. Wang, Q. Wang, and Y. Liu, "Play guessing game with LLM: Indirect jailbreak attack with implicit clues," in *ACL*, 2024.
- [34] Y. Deng, W. Zhang, S. J. Pan, and L. Bing, "Multilingual jailbreak challenges in large language models," *ICLR*, 2024.
- [35] Z.-X. Yong, C. Menghini, and S. H. Bach, "Low-resource languages jailbreak GPT-4," *arXiv preprint arXiv:2310.02446*, 2023.
- [36] L. Shen, W. Tan, S. Chen, Y. Chen, J. Zhang, H. Xu, B. Zheng, P. Koehn, and D. Khashabi, "The language barrier: Dissecting safety challenges of LLMs in multilingual contexts," in *ACL*, 2024.
- [37] J. Li, Y. Liu, C. Liu, L. Shi, X. Ren, Y. Zheng, Y. Liu, and Y. Xue, "A cross-language investigation into jailbreak attacks in large language models," *arXiv preprint arXiv:2401.16765*, 2024.
- [38] Y. Yuan, W. Jiao, W. Wang, J. tse Huang, P. He, S. Shi, and Z. Tu, "GPT-4 is too smart to be safe: Stealthy chat with LLMs via cipher," in *ICLR*, 2024.
- [39] J. Chu, Y. Liu, Z. Yang, X. Shen, M. Backes, and Y. Zhang, "JailbreakRadar: Comprehensive assessment of jailbreak attacks against LLMs," in *ACL*, 2025.
- [40] Y. Meng, M. Xia, and D. Chen, "SimPO: Simple preference optimization with a reference-free reward," in *NeurIPS*, 2024.
- [41] Z. Zhou, J. Xiang, H. Chen, Q. Liu, Z. Li, and S. Su, "Speak out of turn: Safety vulnerability of large language models in multi-turn dialogue," *arXiv preprint arXiv:2402.17262*, 2024.
- [42] X. Liu, L. Li, T. Xiang, F. Ye, L. Wei, W. Li, and N. Garcia, "Imposter. AI: Adversarial attacks with hidden intentions towards aligned large language models," *arXiv preprint arXiv:2407.15399*, 2024.
- [43] N. Li, Z. Han, I. Steneker, W. Primack, R. Goodside, H. Zhang, Z. Wang, C. Menghini, and S. Yue, "LLM defenses are not robust to multi-turn human jailbreaks yet," *arXiv preprint arXiv:2408.15221*, 2024.
- [44] X. Yang, X. Tang, S. Hu, and J. Han, "Chain of attack: A semantic-driven contextual multi-turn attacker for LLM," *arXiv preprint arXiv:2405.05610*, 2024.
- [45] M. Russinovich, A. Salem, and R. Eldan, "Great, now write an article about that: The crescendo multi-turn LLM jailbreak attack," in *USENIX Security*, 2025.
- [46] Q. Ren, H. Li, D. Liu, Z. Xie, X. Lu, Y. Qiao, L. Sha, J. Yan, L. Ma, and J. Shao, "Derail yourself: Multi-turn LLM jailbreak attack through self-discovered clues," *arXiv preprint arXiv:2410.10700*, 2024.
- [47] S. Rahman, L. Jiang, J. Shiffer, G. Liu, S. Issaka, M. R. Parvez, H. Palangi, K.-W. Chang, Y. Choi, and S. Gabriel, "X-Teaming: Multi-turn jailbreaks and defenses with adaptive multi-agents," *arXiv preprint arXiv:2504.13203*, 2025.
- [48] Y. Xie, J. Yi, J. Shao, J. Curl, L. Lyu, Q. Chen, X. Xie, and F. Wu, "Defending ChatGPT against jailbreak attack via self-reminders," *Nature Machine Intelligence*, 2023.
- [49] A. Zhou, B. Li, and H. Wang, "Robust prompt optimization for defending language models against jailbreaking attacks," in *NeurIPS*, 2024.
- [50] Y. Mo, Y. Wang, Z. Wei, and Y. Wang, "Fight back against jailbreaking via prompt adversarial tuning," in *NeurIPS*, 2024.
- [51] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray, J. Schulman, J. Hilton, F. Kelton, L. Miller, M. Simens, A. Askell, P. Welinder, P. F. Christiano, J. Leike, and R. Lowe, "Training language models to follow instructions with human feedback," in *NeurIPS*, 2022, pp. 27 730–27 744.
- [52] Y. Bai, S. Kadavath, S. Kundu, A. Askell, J. Kernion, A. Jones, A. Chen, A. Goldie, A. Mirhoseini, C. McKinnon *et al.*, "Constitutional AI: Harmlessness from AI feedback," *arXiv preprint arXiv:2212.08073*, 2022.
- [53] W. Lu, Z. Zeng, J. Wang, Z. Lu, Z. Chen, H. Zhuang, and C. Chen, "Eraser: Jailbreaking defense in large language models via unlearning harmful knowledge," *arXiv preprint arXiv:2404.05880*, 2024.
- [54] W. Zhao, Z. Li, Y. Li, Y. Zhang, and J. Sun, "Defending large language models against jailbreak attacks via layer-specific editing," in *EMNLP*, 2024.
- [55] Z. Zhang, J. Yang, P. Ke, F. Mi, H. Wang, and M. Huang, "Defending large language models against jailbreaking attacks through goal prioritization," in *ACL*, 2024.
- [56] Y. Yuan, W. Jiao, W. Wang, J.-t. Huang, J. Xu, T. Liang, P. He, and Z. Tu, "Refuse whenever you feel unsafe: Improving safety in llms via decoupled refusal training," in *ACL*, 2025.
- [57] Y. Li, F. Wei, J. Zhao, C. Zhang, and H. Zhang, "RAIN: Your language models can align themselves without finetuning," in *ICLR*, 2024.
- [58] Z. Xu, F. Jiang, L. Niu, J. Jia, B. Y. Lin, and R. Poovendran, "SafeDecoding: Defending against jailbreak attacks via safety-aware decoding," in *ACL*, 2024.
- [59] J. Ji, B. Chen, H. Lou, D. Hong, B. Zhang, X. Pan, T. Qiu, J. Dai, and Y. Yang, "Aligner: Efficient alignment by learning to correct," in *NeurIPS*, 2024.
- [60] S. Han, K. Rao, A. Ettinger, L. Jiang, B. Y. Lin, N. Lambert, Y. Choi, and N. Dziri, "WildGuard: Open one-stop moderation tools for safety risks, jailbreaks, and refusals of LLMs," in *NeurIPS*, 2024.
- [61] A. Lees, V. Q. Tran, Y. Tay, J. Sorensen, J. Gupta, D. Metzler, and L. Vasserman, "A new generation of perspective API: Efficient multilingual character-level transformers," in *KDD*, 2022.

- [62] T. Markov, C. Zhang, S. Agarwal, F. E. Nekoul, T. Lee, S. Adler, A. Jiang, and L. Weng, "A holistic approach to undesired content detection in the real world," in *AAAI*, 2023.
- [63] M. Phute, A. Helbling, M. Hull, S. Peng, S. Szyller, C. Cornelius, and D. H. Chau, "LLM Self Defense: By self examination, LLMs know they are being tricked," *arXiv preprint arXiv:2308.07308*, 2023.
- [64] G. Alon and M. Kamfonas, "Detecting language model attacks with perplexity," *arXiv preprint arXiv:2308.14132*, 2023.
- [65] N. Jain, A. Schwarzschild, Y. Wen, G. Somepalli, J. Kirchenbauer, P.-y. Chiang, M. Goldblum, A. Saha, J. Geiping, and T. Goldstein, "Baseline defenses for adversarial attacks against aligned language models," *arXiv preprint arXiv:2309.00614*, 2023.
- [66] A. Kumar, C. Agarwal, S. Srinivas, A. J. Li, S. Feizi, and H. Lakkaraju, "Certifying LLM safety against adversarial prompting," in *COLM*, 2024.
- [67] A. Robey, E. Wong, H. Hassani, and G. J. Pappas, "SmoothLLM: Defending large language models against jailbreaking attacks," *arXiv preprint arXiv:2310.03684*, 2023.
- [68] Y. Xie, M. Fang, R. Pi, and N. Gong, "GradSafe: Detecting jailbreak prompts for LLMs via safety-critical gradient analysis," in *ACL*, 2024.
- [69] J. Ji, B. Hou, A. Robey, G. J. Pappas, H. Hassani, Y. Zhang, E. Wong, and S. Chang, "Defending large language models against jailbreak attacks via semantic smoothing," *arXiv preprint arXiv:2402.16192*, 2024.
- [70] S. Goyal, M. Hira, S. Mishra, S. Goyal, A. Goel, N. Dadu, K. DB, S. Mehta, and N. Madaan, "LLMGuard: Guarding against unsafe LLM behavior," in *AAAI*, 2024.
- [71] X. Hu, P.-Y. Chen, and T.-Y. Ho, "Gradient Cuff: Detecting jailbreak attacks on large language models by exploring refusal loss landscapes," in *NeurIPS*, 2024.
- [72] Y. Zeng, Y. Wu, X. Zhang, H. Wang, and Q. Wu, "AutoDefense: Multi-agent LLM defense against jailbreak attacks," *arXiv preprint arXiv:2403.04783*, 2024.
- [73] Z. Yuan, Z. Xiong, Y. Zeng, N. Yu, R. Jia, D. Song, and B. Li, "RigorLLM: Resilient guardrails for large language models against undesired content," in *ICML*, 2024.
- [74] S. Ghosh, P. Varshney, E. Galinkin, and C. Parisien, "Aegis: Online adaptive AI content safety moderation with ensemble of LLM experts," *arXiv preprint arXiv:2404.05993*, 2024.
- [75] Z. Chu, Y. Wang, L. Li, Z. Wang, Z. Qin, and K. Ren, "A causal explainable guardrails for large language models," in *CCS*, 2024.
- [76] A. Kawasaki, A. Davis, and H. Abbas, "Defending large language models against attacks with residual stream activation analysis," in *CAMLIS*, 2024.
- [77] A. Zou, L. Phan, J. Wang, D. Duenas, M. Lin, M. Andriushchenko, R. Wang, Z. Kolter, M. Fredrikson, and D. Hendrycks, "Improving alignment and robustness with circuit breakers," in *NeurIPS*, 2024.
- [78] Z. Xiang, L. Zheng, Y. Li, J. Hong, Q. Li, H. Xie, J. Zhang, Z. Xiong, C. Xie, C. Yang *et al.*, "GuardAgent: Safeguard LLM agents by a guard agent via knowledge-enabled reasoning," in *ICML*, 2025.
- [79] M. Kang and B. Li, " R^2 -Guard: Robust reasoning enabled LLM guardrail via knowledge-enhanced logical reasoning," in *ICLR*, 2025.
- [80] M. Llama, "Prompt-guard-86m," <https://huggingface.co/meta-llama/Prompt-Guard-86M>, 2024.
- [81] D. Schwartz, D. Bespalov, Z. Wang, N. Kulkarni, and Y. Qi, "Graph of attacks with pruning: Optimizing stealthy jailbreak prompt generation for enhanced llm content moderation," *arXiv preprint arXiv:2501.18638*, 2025.
- [82] B. Manczak, E. Zemour, E. Lin, and V. Mugunthan, "PrimeGuard: Safe and helpful LLMs through tuning-free routing," *arXiv preprint arXiv:2407.16318*, 2024.
- [83] W. Zeng, Y. Liu, R. Mullins, L. Peran, J. Fernandez, H. Harkous, K. Narasimhan, D. Proud, P. Kumar, B. Radharapu *et al.*, "Shield-Gemma: Generative AI content moderation based on Gemma," *arXiv preprint arXiv:2407.21772*, 2024.
- [84] J. Hu, Y. Dong, and X. Huang, "Trust-oriented adaptive guardrails for large language models," *arXiv preprint arXiv:2408.08959*, 2024.
- [85] C. Zhao, Z. Dou, and K. Huang, "EEG-Defender: Defending against jailbreak through early exit generation of large language models," in *ICONIP*, 2025.
- [86] C. Qian, H. Zhang, L. Sha, and Z. Zheng, "HSF: Defending against jailbreak attacks with hidden state filtering," in *WWW*, 2025.
- [87] G. Cornacchia, G. Zizzo, K. Fraser, M. Z. Hameed, A. Rawat, and M. Purcell, "MoJE: Mixture of jailbreak experts, naive tabular classifiers as guard for prompt attacks," in *AIES*, 2024.
- [88] A. Peng, J. Michael, H. Sleight, E. Perez, and M. Sharma, "Rapid response: Mitigating LLM jailbreaks with a few examples," *arXiv preprint arXiv:2411.07494*, 2024.
- [89] E. Galinkin and M. Sablotny, "Improved large language model jailbreak detection via pretrained embeddings," *arXiv preprint arXiv:2412.01547*, 2024.
- [90] X. Hu, P.-Y. Chen, and T.-Y. Ho, "Token Highlighter: Inspecting and mitigating jailbreak prompts for large language models," in *AAAI*, 2025.
- [91] S. Ghosh, P. Varshney, M. N. Sreedhar, A. Padmakumar, T. Rebedea, J. R. Varghese, and C. Parisien, "Aegis2.0: A diverse ai safety dataset and risks taxonomy for alignment of LLM guardrails," *arXiv preprint arXiv:2501.09004*, 2025.
- [92] M. K. Rad, H. Nghiem, A. Luo, S. Wadhwa, M. Sorower, and S. Rawls, "Refining input guardrails: Enhancing llm-as-a-judge efficiency through chain-of-thought fine-tuning and alignment," *arXiv preprint arXiv:2501.13080*, 2025.
- [93] Y. Liu, H. Gao, S. Zhai, J. Xia, T. Wu, Z. Xue, Y. Chen, K. Kawaguchi, J. Zhang, and B. Hooi, "GuardReasoner: Towards reasoning-based LLM safeguards," *arXiv preprint arXiv:2501.18492*, 2025.
- [94] M. Sharma, M. Tong, J. Mu, J. Wei, J. Kruthoff, S. Goodfriend, E. Ong, A. Peng, R. Agarwal, C. Anil *et al.*, "Constitutional classifiers: Defending against universal jailbreaks across thousands of hours of red teaming," *arXiv preprint arXiv:2501.18837*, 2025.
- [95] S. Xiang, A. Zhang, Y. Cao, Y. Fan, and R. Chen, "Beyond surface-level patterns: An essence-driven defense framework against jailbreak attacks in LLMs," in *ACL Findings*, 2025.
- [96] C. Yung, H. Huang, S. M. Erfani, and C. Leckie, "Curvalid: Geometrically-guided adversarial prompt detection," *arXiv preprint arXiv:2503.03502*, 2025.
- [97] R. Pu, C. Li, R. Ha, L. Zhang, L. Qiu, and X. Zhang, "MirrorShield: Towards universal defense against jailbreaks via entropy-guided mirror crafting," *arXiv preprint arXiv:2503.12931*, 2025.
- [98] X. Zhang, C. Zhang, T. Li, Y. Huang, X. Jia, M. Hu, J. Zhang, Y. Liu, S. Ma, and C. Shen, "JailGuard: A universal detection framework for prompt-based attacks on LLM systems," *ACM Trans. Softw. Eng. Methodol.*, 2025.
- [99] B. Upadhayay, V. Behzadan *et al.*, "X-Guard: Multilingual guard agent for content moderation," *arXiv preprint arXiv:2504.08848*, 2025.
- [100] J. Piet, X. Huang, D. Jacob, A. Chow, M. Alrashed, G. Zhao, Z. Hu, C. Sitawarin, B. Alomair, and D. Wagner, "Jailbreakovertime: Detecting jailbreak attacks under distribution shift," *arXiv preprint arXiv:2504.19440*, 2025.

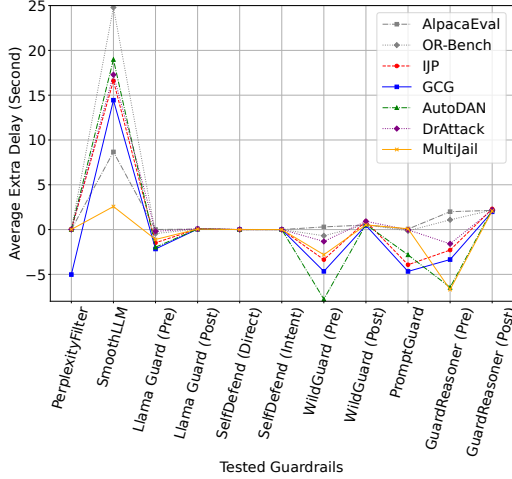


Figure 6. The delay of guardrails on GPT-4-0125-Preview. AlpacaEval and OR-Bench are two normal prompt datasets and others are jailbreak attacks.

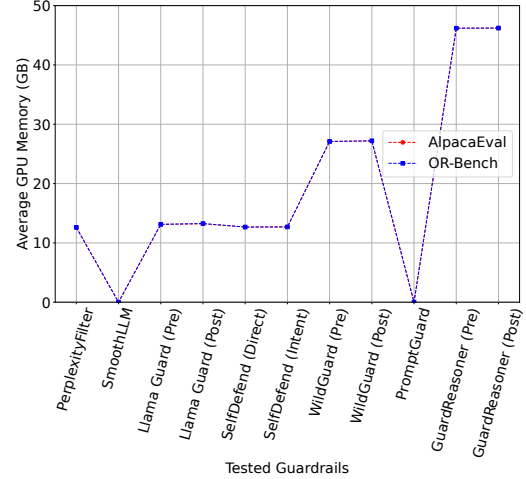


Figure 7. The memory usage of guardrails on GPT-4-0125-Preview.

- [101] P. Chao, E. Debenedetti, A. Robey, M. Andriushchenko, F. Croce, V. Sehwag, E. Dobriban, N. Flammarion, G. J. Pappas, F. Tramèr *et al.*, “JailbreakBench: An open robustness benchmark for jailbreak large language models,” in *NeurIPS*, 2024.
- [102] X. Li, T. Zhang, Y. Dubois, R. Taori, I. Gulrajani, C. Guestrin, P. Liang, and T. B. Hashimoto, “AlpacaEval: An automatic evaluator of instruction-following models,” 2023.
- [103] J. Cui, W.-L. Chiang, I. Stoica, and C.-J. Hsieh, “OR-Bench: An over-refusal benchmark for large language models,” in *ICML*, 2025.
- [104] J. Achiam, S. Adler, S. Agarwal, L. Ahmad, I. Akkaya, F. L. Aleman, D. Almeida, J. Altenschmidt, S. Altman, S. Anadkat *et al.*, “GPT-4 technical report,” *arXiv preprint arXiv:2303.08774*, 2023.
- [105] X. Qi, Y. Zeng, T. Xie, P.-Y. Chen, R. Jia, P. Mittal, and P. Henderson, “Fine-tuning aligned language models compromises safety, even when users do not intend to!” *arXiv preprint arXiv:2310.03693*, 2023.
- [106] “Forbidden question set with prompts,” https://github.com/verazuo/jailbreak_llms/blob/main/data/forbidden_question/forbidden_question_set_with_prompts.csv.zip, 2023.
- [107] W. Luo, S. Ma, X. Liu, X. Guo, and C. Xiao, “JailBreakV: A benchmark for assessing the robustness of multimodal large language models against jailbreak attacks,” in *COLM*, 2024.

Appendix A.

Evaluation Results on Vicuna-13b-v1.5

We extend our comprehensive evaluation to another widely-used open-source model, Vicuna-13b-v1.5, to assess the generalization of different guardrails. The detailed results are presented in Table 4.

First, a salient observation is that Vicuna-13b-v1.5 is considerably more susceptible to jailbreak attacks compared to Llama-3. This increased vulnerability is evident from the substantially higher Attack Success Rates (ASR) across almost all attack categories, indicating a weaker inherent safety alignment in Vicuna.

Second, we note a significant performance degradation for certain defenses when applied to Vicuna-13b-v1.5. For

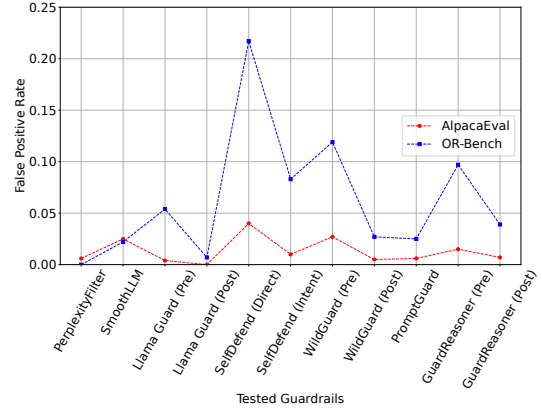


Figure 8. The utility (FPR) of guardrails on GPT-4-0125-Preview.

example, the efficacy of GradSafe and GradientCuff diminishes. GradientCuff, which showed marked effectiveness against the X-Teaming multi-turn attack on Llama-3, fails to maintain this advantage on Vicuna-13b-v1.5. This decline can be attributed to their nature as intra-processing guardrails, which heavily rely on the internal representations and alignment of the target LLM. Consequently, a less well-aligned model like Vicuna-13b-v1.5 compromises their defensive mechanism.

Despite these differences, we also observe consistent performance patterns. GuardReasoner (Pre) and GuardReasoner (Post) continue to exhibit state-of-the-art defense capabilities. GuardReasoner (Pre) achieves the best overall ASR of 0.156, while GuardReasoner (Post) records the best overall PGR of 0.192. This sustained excellence underscores that the robust defense mechanism of GuardReasoner is largely independent of the target LLM, positioning it as a more universally applicable and reliable guardrail.

TABLE 4. THE ASR ($\downarrow$) / PGR ($\downarrow$) RESULTS FOR THE TARGET LLM (VICUNA-13B-V1.5) WITH DIFFERENT GUARDRAILS AGAINST FIVE MAJOR CATEGORIES OF JAILBREAK ATTACKS, INCLUDING ROW AVERAGES. (PRE) AND (POST) DENOTE THE PRE-PROCESSING AND POST-PROCESSING VERSIONS OF THE GUARDRAILS, RESPECTIVELY. (DIRECT) AND (INTENT) DENOTE THE DIRECT PROMPT AND INTENT PROMPT BASED VERSIONS OF SELFDEFEND [11], RESPECTIVELY.

Guardrails	Manual	Optimization-based		Generation-based		Implicit		Multi-turn		Average
	IJP	GCG	AutoDAN	TAP	LLM-Fuzzer	DrAttack	MultiJail	ActorAttack	X-Teaming	
Vicuna-13b-v1.5	0.474/-	0.890/-	0.660/-	0.530/-	0.820/-	0.780/-	0.254/-	0.238/-	0.960/-	0.649/-
PerplexityFilter	0.474/1.000	0.030/0.040	0.660/1.000	0.830/1.000	0.870/1.000	0.780/1.000	0.254/1.000	0.238/1.000	0.990/1.000	0.570/0.893
SmoothLLM	0.402/0.794	0.140/0.270	0.520/0.970	0.840/0.860	0.510/1.000	0.410/0.970	0.152/0.933	0.877/0.877	0.980/0.970	0.537/0.849
Llama Guard (Pre)	0.194/0.563	0.370/0.390	0.460/0.750	0.630/0.750	0.810/1.000	0.650/0.850	0.251/0.952	0.222/0.967	0.970/1.000	0.506/0.802
Llama Guard (Post)	0.250/0.250	0.400/0.400	0.610/0.610	0.600/0.600	0.830/0.830	0.390/0.390	0.248/0.248	0.230/0.230	0.970/0.970	0.503/0.503
GradSafe	0.471/0.994	0.890/1.000	0.660/1.000	0.580/0.960	0.900/1.000	0.780/1.000	0.254/1.000	0.238/1.000	0.980/1.000	0.639/0.995
GradientCuff	0.193/0.351	0.090/0.090	0.310/0.480	0.550/0.630	0.780/1.000	0.660/0.830	0.000/0.000	0.183/0.805	0.930/0.960	0.411/0.572
SelfDefend (Direct)	0.050/0.262	0.080/0.080	0.020/0.080	0.210/0.270	0.190/0.270	0.330/0.480	0.187/0.743	0.132/0.720	0.960/0.990	0.240/0.433
SelfDefend (Intent)	0.057/0.286	0.080/0.080	0.050/0.110	0.140/0.200	0.210/0.250	0.010/0.090	0.127/0.568	0.157/0.763	0.960/1.000	0.199/0.372
WildGuard (Pre)	0.007/0.033	0.010/0.010	0.010/0.020	0.040/0.090	0.010/0.020	0.330/0.490	0.187/0.797	0.147/0.757	0.920/0.950	0.185/0.352
WildGuard (Post)	0.066/0.066	0.040/0.040	0.030/0.030	0.100/0.100	0.410/0.410	0.090/0.090	0.194/0.194	0.165/0.165	0.930/0.930	0.225/0.225
Prompt Guard	0.000/0.000	0.020/0.020	0.240/0.370	0.570/0.960	0.020/0.030	0.770/0.990	0.254/1.000	0.235/0.995	0.980/1.000	0.343/0.596
GuardReasoner (Pre)	0.000/0.009	0.000/0.000	0.020/0.020	0.050/0.080	0.040/0.040	0.150/0.270	0.057/0.349	0.143/0.740	0.940/0.960	0.156/0.274
GuardReasoner (Post)	0.050/0.050	0.030/0.030	0.010/0.010	0.060/0.060	0.480/0.480	0.040/0.040	0.060/0.060	0.100/0.100	0.900/0.900	0.192/0.192

TABLE 5. THE ASR ($\downarrow$) / PGR ($\downarrow$) RESULTS FOR THE TARGET LLM (GPT-4-0125-PREVIEW).

Guardrails	Manual	Optimization-based		Generation-based		Implicit		Multi-turn		Average
	IJP	GCG	AutoDAN	TAP	LLM-Fuzzer	DrAttack	MultiJail	ActorAttack	X-Teaming	
GPT-4-0125-Preview	0.192/-	0.220/-	0.710/-	0.340/-	0.750/-	0.220/-	0.029/-	0.200/-	0.940/-	0.400/-
PerplexityFilter	0.192/1.000	0.000/0.060	0.710/1.000	0.260/1.000	0.720/1.000	0.220/1.000	0.029/1.000	0.200/1.000	0.930/1.000	0.362/0.896
SmoothLLM	0.185/0.365	0.150/0.270	0.480/0.850	0.540/0.630	0.760/0.940	0.300/0.780	0.054/1.000	0.192/0.905	0.980/0.980	0.405/0.747
Llama Guard (Pre)	0.100/0.563	0.130/0.460	0.540/0.770	0.250/0.510	0.660/0.840	0.220/0.850	0.029/0.952	0.200/0.998	0.940/0.990	0.341/0.770
Llama Guard (Post)	0.110/0.110	0.180/0.180	0.530/0.530	0.250/0.250	0.720/0.720	0.220/0.220	0.029/0.029	0.200/0.200	0.950/0.950	0.354/0.354
SelfDefend (Direct)	0.031/0.253	0.060/0.100	0.160/0.230	0.100/0.210	0.110/0.160	0.070/0.350	0.019/0.740	0.102/0.693	0.870/0.910	0.169/0.405
SelfDefend (Intent)	0.039/0.310	0.040/0.090	0.080/0.160	0.100/0.190	0.100/0.180	0.060/0.140	0.025/0.568	0.128/0.768	0.910/0.930	0.165/0.371
WildGuard (Pre)	0.004/0.033	0.000/0.020	0.020/0.020	0.030/0.060	0.010/0.010	0.140/0.530	0.029/0.797	0.127/0.768	0.810/0.880	0.130/0.346
WildGuard (Post)	0.035/0.035	0.030/0.030	0.070/0.070	0.060/0.060	0.180/0.180	0.140/0.140	0.029/0.029	0.142/0.142	0.920/0.920	0.178/0.178
Prompt Guard	0.000/0.000	0.020/0.050	0.410/0.630	0.310/0.960	0.000/0.000	0.220/0.970	0.029/1.000	0.195/0.993	0.910/1.000	0.233/0.623
GuardReasoner (Pre)	0.000/0.009	0.000/0.000	0.030/0.030	0.030/0.040	0.050/0.070	0.080/0.290	0.016/0.349	0.113/0.740	0.890/0.960	0.134/0.276
GuardReasoner (Post)	0.021/0.021	0.030/0.030	0.040/0.040	0.080/0.080	0.240/0.240	0.100/0.100	0.029/0.029	0.110/0.110	0.930/0.930	0.176/0.176

Appendix B. Evaluation Results on GPT-4-0125-Preview

To further validate the generalizability of our findings, we evaluate guardrail performance on GPT-4-0125-Preview, a representative closed-source model. The comprehensive results are presented in Table 5, Figure 6, 7, and 8.

Table 5 demonstrates that the defense performance metrics (ASR/PGR) on GPT-4 consistently align with those observed in Llama-based evaluations. Notably, PerplexityFilter achieves the highest average PGR, while GuardReasoner (Post) maintains the lowest. Furthermore, the multi-turn X-Teaming attack continues to exhibit exceptionally high ASR/PGR across all guardrails, reinforcing the findings from RQ3 regarding session-level defense limitations.

Figure 6 reveals consistent efficiency patterns: GuardReasoner incurs substantial latency due to its reasoning overhead, and post-processing guardrails against jailbreaks universally introduce greater delay than their pre-processing counterparts, directly addressing RQ1. The pronounced latency of SmoothLLM stems from its defense mechanism requiring 10 perturbed prompt copies. This issue is further exacerbated by inefficient asynchronous

API access to GPT-4 when compared to batched local inference.

Figure 7 illustrates the GPU memory usage of different guardrails when applied to GPT-4. These GPU memory consumptions (in GB) for GPT-4 align with the conclusions drawn from Llama-3. Specifically: (1) GuardReasoner exhibits the highest memory footprint. (2) The memory overhead for SmoothLLM and PromptGuard is negligible. (3) LLM-based guardrails such as Llama Guard, SelfDefend, WildGuard, and GuardReasoner consume more GPU memory compared to rule-based guardrails (e.g., SmoothLLM). These have the same answer to RQ2.

The utility results (FPR) in Figure 8 for GPT-4 corroborate the conclusions drawn from Llama-3. Specifically: (1) SelfDefend (Direct) exhibits the highest FPR on OR-Bench. (2) Token-level guardrails (SmoothLLM, SelfDefend (Direct)) demonstrate relatively pronounced FPRs. (3) Session-level guardrails (those with a “Post” suffix) consistently achieve markedly lower FPRs than their sequence-level counterparts (those with a “Pre” suffix). These findings consistently address RQ4.

These results demonstrate that our SEU framework and taxonomic insights generalize effectively to large-scale black-box models, affirming the robustness of our findings.

Appendix C. Meta-Review

The following meta-review was prepared by the program committee for the 2026 IEEE Symposium on Security and Privacy (S&P) as part of the review process as detailed in the call for papers.

C.1. Summary

This paper presents a defense-oriented systematization of knowledge on protecting large language models against jailbreak attacks. It proposes a taxonomy of jailbreak techniques and categorizes guardrail defenses along several key dimensions: the underlying technical paradigm, the granularity of protection, the degree of reactivity, their applicability across contexts, and the level of explainability they provide. In addition, the paper includes an evaluation of existing defenses, offering comparative insights into their effectiveness.

C.2. Scientific Contributions

- A systematization of defenses for LLMs

C.3. Reasons for Acceptance

- 1) This paper proposes a comprehensive taxonomy of LLM guardrails. The presentation is clear and has good coverage of existing work.
- 2) The paper also proposes a new evaluation framework that considers real-world trade-offs like latency, computational cost and usability.
- 3) The evaluation framework is demonstrated with various defenses and LLM models. This offers concrete benchmarks for comparing defenses.