

Function Calling as a Flexible LLM Defense Add-On: Capability and Application Exploration

ZHENLAN JI, Nara Institute of Science and Technology, Japan

DAOYUAN WU, Lingnan University, China

WENXUAN WANG, Renmin University of China, China

PINGCHUAN MA, Zhejiang University of Technology, China

SHUAI WANG*, Hong Kong University of Science and Technology, China

LEI MA*[†], University of Tokyo, Japan and University of Alberta, Canada

JUERGEN RAHMEI, HSBC, China

Large language models (LLMs) exhibit impressive capabilities but are susceptible to adversarial attacks that induce harmful outputs. Although various defenses have been proposed, their practicality is restricted by substantial runtime overhead or degraded model helpfulness. Moreover, LLM applications typically have diverse and evolving security requirements that cannot be fully anticipated during the design of static defenses. These limitations call for a flexible, low-overhead defense mechanism that can be easily customized to meet task-specific needs.

In this paper, we explore function calling (FC)—a built-in mechanism in modern LLMs for invoking custom tools—as a lightweight and adaptable defense add-on. We show that by defining functions representing malicious actions, LLMs equipped with FC can intercept harmful prompts by triggering these function calls instead of generating unsafe content. Extensive experiments across mainstream LLMs demonstrate that FC substantially improves defense effectiveness with minimal impact on model helpfulness. To further assess FC’s practical utility, we also introduce DSPEC, a new dataset reflecting real-world LLM applications with specific defense requirements. Our evaluations on DSPEC show that FC substantially outperforms existing defenses in this realistic setting. Besides, we also explore the practical applications of FC in various scenarios, including universal defense frameworks and multi-agent systems, further demonstrating its versatility and effectiveness in enhancing LLM security.

CCS Concepts: • **Software and its engineering** → **Software post-development issues**; • **Security and privacy** → **Software security engineering**.

Additional Key Words and Phrases: Large Language Models, LLM Application, Software Security

ACM Reference Format:

Zhenlan Ji, Daoyuan Wu, Wenxuan Wang, Pingchuan Ma, Shuai Wang, Lei Ma, and Juergen Rahmel. 2026. Function Calling as a Flexible LLM Defense Add-On: Capability and Application Exploration. *Proc. ACM Softw. Eng.* 3, ISSTA, Article ISSTA118 (October 2026), 23 pages. <https://doi.org/10.1145/3832209>

*Corresponding authors.

[†]Half of this work was done when the first author was a visiting student at the University of Tokyo under the supervision of this author.

Authors’ Contact Information: Zhenlan Ji, Nara Institute of Science and Technology, Nara, Japan, ji.zhenlan@naist.ac.jp; Daoyuan Wu, Lingnan University, Hong Kong, China, daoyuanwu@ln.edu.hk; Wenxuan Wang, Renmin University of China, Beijing, China, jwxwang@gmail.com; Pingchuan Ma, Zhejiang University of Technology, Hangzhou, China, pma@zjut.edu.cn; Shuai Wang, Hong Kong University of Science and Technology, Hong Kong, China, shuaiw@cse.ust.hk; Lei Ma, University of Tokyo, Tokyo, Japan and University of Alberta, Edmonton, Canada, ma.lei@acm.org; Juergen Rahmel, HSBC, Hong Kong, China, juergen.rahmel@hsbc.com.hk.



This work is licensed under a Creative Commons Attribution 4.0 International License.

© 2026 Copyright held by the owner/author(s).

ACM 2994-970X/2026/10-ARTISSTA118

<https://doi.org/10.1145/3832209>

1 Introduction

Large language models (LLMs) have demonstrated remarkable capabilities across a wide range of applications [18, 30, 44, 66, 76]. However, these advances have also given rise to escalating security challenges that undermine their integrity and trustworthiness. During inference, adversaries can exploit LLMs' impressive understanding and reasoning abilities to manipulate model outputs. Such attacks may involve eliciting harmful responses [72], extracting sensitive information [54], or deceiving users [47, 60], further exacerbating concerns about LLM security and hindering the widespread adoption of the new-generation software systems, LLM applications, in the real world.

To mitigate these threats, the research community has proposed various defense strategies that broadly fall into two categories: *external defenses* and *internal defenses*. *External defenses* [19, 25, 37, 63, 70, 77] rely on auxiliary models or monitoring modules to intercept or flag malicious inputs and outputs. Although effective in certain settings, these methods often introduce significant latency or computational overhead, which is unacceptable for many latency-sensitive applications. *Internal defenses* [16, 23, 31, 68], such as fine-tuning or instruction rejection, modify the model itself to increase robustness. However, these modifications may degrade model helpfulness [74] and sometimes result in false positives, whereby benign inputs are wrongly rejected. More importantly, both types of defenses are typically designed statically and cannot accommodate diverse and evolving security needs across applications. For example, an LLM powering a financial assistant may require protection against phishing queries, while a medical chatbot must guard against unauthorized diagnosis generation. These context-specific needs are not easily addressed by static, one-size-fits-all defenses.

In this paper, we formulate the mechanism underlying LLMs' decision-making procedure and summarize the common stages shared by the existing inference-time attacks. Holistically, there are two key stages: (1) circumventing the model's malicious input recognition, and (2) precisely conveying their actual malicious intention that is disguised by the attack strategy. We see that existing defense strategies primarily focus on enhancing the model's malicious input recognition, which may be impeded by the often limited coverage of malicious input space. Notice that the attack surface is generally *infinite* due to LLM's exceptional capability of natural language understanding and generalization. In contrast, we propose to novelly use the *function calling* (FC) feature in LLMs as a defense strategy. The intuition is that, since any attack strategy must convey its malicious intent to the model to achieve its end goal, we can design a set of "defense functions" to match the possible malicious intents. The model can then be instructed to call the corresponding function when it is handling a malicious intent. This approach requires no changes to model internals (e.g., weights), can be adapted on the fly by simply modifying the available functions, and does not affect model behavior on benign inputs.

To validate the feasibility and effectiveness of our approach, we conduct a comprehensive pilot study using four mainstream LLMs—GPT-3.5-Turbo [5], GPT-4o-mini [6], GPT-5.4-mini [15], and Claude-3.5-haiku [3]—against a large-scale adversarial dataset, Wildjailbreak [38]. Besides, the helpfulness and time overhead are also considered to evaluate the side effects of FC, as reflected by the results of a prevailing LLM benchmark, MMLU-Pro [71]. Given the lightweight nature of FC, we compare it with prompt-engineering, a similarly lightweight defense strategy that instructs the model to refuse malicious requests in natural language. The results show that FC can remarkably improve defense success rates with minimal degradation in helpfulness and latency. Notably, FC outperforms prompt defenses designed to refuse requests that share the same malicious intent. We attribute this interesting phenomenon to the fundamental difference on conceptual mechanism between the two defense strategies: prompt defenses actively instruct the model to refuse a request,

while FC acts as a “honeypot” that lures the model into revealing its intention via a function call. This result strongly suggests that FC’s mechanism is a more promising direction for LLM security.

After that, in light of the lack of LLM benchmarks that cover a broad range of real-world LLM applications, we construct a new benchmark dataset, DSPEC. The objective of this dataset is to assess the adaptability and efficacy of defenses in real-world scenarios by reflecting the nuanced and dynamic defense requirements across a variety of real-world LLM applications. In particular, there are seven different scenarios covered in DSPEC—from blocking newly emergent chemical compound queries to controlling access based on user roles in enterprise settings, with 1,000 samples in total. We evaluate FC and existing defenses on DSPEC and find that FC consistently outperforms prior methods on all LLMs while incurring comparable or lower additional overhead, which demonstrates FC’s versatility and practicality.

Finally, we also explore the practical applications of FC, especially in complex LLM-based systems, where its lightweight and adaptable nature can be fully leveraged. We propose a universal defense framework that integrates FC with existing defenses to further enhance LLM security. This framework exhibits a remarkable advance in both defensive effectiveness and processing speed, resulting from the positive synergy between FC and other defenses. Moreover, we believe that this framework also demonstrates a promising solution to the scalability and adaptability challenges faced by LLM defenses, i.e., combining general-purpose defenses (e.g., LlamaGuard [53]) to handle large-scale, common attacks with application-specific defenses (our FC-based approach) to address the application-specific attacks that evolve rapidly. We also demonstrate FC’s utility in multi-agent systems [29], where developers can integrate a variety of defensive functions into different agents. We believe that this facilitates a “multi-turn defense” paradigm, which is analogous to multi-turn attacks [61], by leveraging the multi-agent collaboration. In this paradigm, the backend LLM evolves its reasoning based on the intermediate responses from multiple agents, thereby enhancing robustness. These successful applications highlight FC’s practical value as a flexible plug-in defense add-on to enhance the security of LLM-powered systems, further showcasing its potential in real-world deployments. In summary, our main contributions are:

- To the best of our knowledge, we are the first to discover and empirically validate a novel use of the function calling mechanism as a flexible and effective defense add-on for LLMs.
- We develop DSPEC, a new dataset capturing real-world, application-specific defense requirements, and show that FC can adapt to these contexts and synergize with existing strategies to form a universal defense framework.
- We explore practical applications of FC in several complex LLM-based systems, including universal defense frameworks and multi-agent systems, further demonstrating its versatility and potential in enhancing LLM security.

Outline. The rest of the paper is organized as follows. We first provide background regarding LLM inference security and function calling in Sec. 2. After that, Sec. 3 overviews FC as a flexible LLM defense add-on. A series of experiments are conducted as a pilot study to illustrate the defensive effectiveness of FC in Sec. 4. Then, we propose several real-world scenarios to form a comprehensive dataset, DSPEC, and then evaluate the defensive performance of FC on it in Sec. 5. In Sec. 6, we explore diverse applications of FC. Lastly, we discuss FC’s resilience against adaptive attacks and limitations in Sec. 7, review related work in Sec. 8, and conclude in Sec. 9.

2 Preliminary

2.1 LLM Inference-Time Security

As the inference-time behavior of large language models (LLMs) directly affects the general user experience and determines the security of the model, it is crucial to probe the security of LLM

inference. Two types of methods, jailbreaking and prompt injection, constitute the mainstream inference-time attacks on LLMs [24]. Jailbreaking aims to circumvent the built-in safety policies of LLMs [50, 64, 79, 84, 85], while prompt injection attempts to manipulate the inference outcomes by injecting malicious prompts [52, 57, 62, 65, 67]. Although different techniques are proposed to implement these two types of attacks, the core idea shared by them is to elicit harmful responses from LLMs.

To defend against these attacks, researchers have proposed a variety of defense strategies from internal and external perspectives. Internal defense strategies aim to enhance the robustness of the model itself. The most most prevalent approach involves the use of fine-tuning or similar techniques to update the model's parameters [16, 23, 68]. Benefiting from LLMs' prominent generalization capability, prompt learning, which uses natural language text to guide the model's behavior [48], also emerges as a promising defense strategy [31]. External defense strategies, on the other hand, attempt to detect or mitigate the malicious inputs/outputs of LLMs with the help of auxiliary models or techniques. Specifically, the model-based external defense strategies typically rely on the assistance of another model, such as LlamaGuard [37], or the model itself [19, 70] (since the model is used to address the input one/multiple more times with different settings, it can be deemed as another model). Inevitably, the adoption of extra models induces prohibitive additional overhead. Other external defense strategies attempt to avoid this overhead by directly inspecting the model's internal state to detect the malicious inputs [25, 63, 77], whereas this inspection also suffers from the excessive computational overhead.

Although all the aforementioned defense strategies endeavor to impede the model from generating malicious outputs, they still suffer from the limited coverage of the attack surface compared to the infinite malicious semantic space enabled by the exceptional understanding capability of LLMs. On the contrary, our method obviates the disturbance of various attacks by focusing on the model's intention to generate malicious outputs.

2.2 Function Calling in LLMs

Function calling (FC), also known as tool-calling, is a technique that enables LLMs to call external functions to accomplish specific tasks [12]. By providing the model with the necessary information of the external function (name, description, arguments, and return specifications), the model can automatically generate a specific prompt (usually a JSON/XML format string) to call the required function. This feature has been widely supported by the state-of-the-art LLMs, both commercial and open-source [9, 10, 12]. The emergence of FC in LLMs enables the model to act as a versatile agent [73] that can interact with networks and databases [7, 33] to provide various services [45, 51], effectively bridging the gap between human language and machine-executable actions.

In the context of LLM inference-time security, we presume that the FC mechanism is capable of precisely representing the intention of the model. In other words, once the model intends to generate a response related to policy-violating content, it would tend to find whether the provided function can be called. In light of the community's enthusiasm to further enhance the FC functionality and the current performance of FC [17, 22, 35, 39, 51, 56], our presumption is reasonable.

3 FC as an LLM Defense Add-On: Overview

In this paper, we propose that function calling can serve as a flexible add-on to defend against various inference-time attacks on LLMs, such as jailbreaking and prompt injection. Under this idea, we first formulate the corresponding threat model in Sec. 3.1. Then, we illustrate the intuition behind the FC defense strategy in Sec. 3.2.

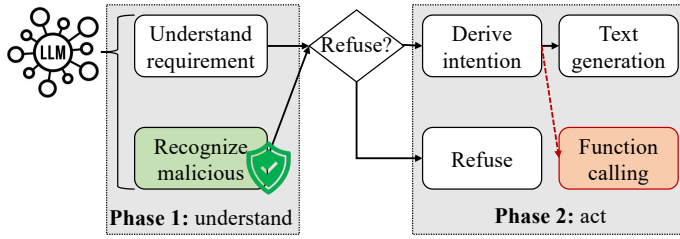


Fig. 1. Illustration of LLMs' decision-making procedure.

3.1 Threat Model

In light of the increasing prevalence of LLM applications, e.g., Poe [13] and OpenAI's GPTs [11], we assume a specific threat model tailored for LLM applications with two roles (the defender and the attacker) as follows:

Defender. The defender is the developer of the LLM application, with the objective of safeguarding the application from undesired behaviors that are detrimental to their interests. The defender possesses comprehensive knowledge of the application's business logic and domain-specific details, enabling them to craft precise system prompts and configure hyperparameters of the backend LLM to guide its behavior appropriately. Accordingly, they have a clear understanding of the undesired behavior they aim to prevent. Due to the computational resource constraints or a lack of expertise related to LLM training, the defender typically deploys their application backed by an off-the-shelf pre-trained LLM, especially commercial LLMs like ChatGPT, without altering its internal parameters [36, 49].

Attacker. The attacker seeks to bypass the application's alignment and defense to extract responses detrimental to the developer's interests. Operating with black-box access, the attacker can only send textual queries and receive responses, without any capability to alter the application's settings or the LLM itself, similar to the threat model in [36].

Note that since the LLM applications are typically designed for specific tasks [11, 13], the prevention objectives of the defender are also specific. For instance, the defender may concentrate on preventing the leakage of privacy data in a dataset querying system, while other malicious text, such as hate speech, is not of concern to the defender. Given that different applications possess distinct, specific prevention objectives, we presume that the most pressing requirement of the defender is a defense strategy that is both lightweight and flexible. This strategy should be easily customizable to meet the specific requirements of the application.

3.2 Our Intuition

Intuition Fundamentals. Given the excessive overhead incurred by external defenses as mentioned in Sec. 2.1, this paper aims to propose a novel defense strategy that directly enhances the security of the model itself. After reviewing existing works on LLM inference-time security, we formulate the mechanism underlying the decision-making procedure of LLMs as presented in Fig. 1. In general, there are two phases in the decision-making procedure of LLMs.¹ In the first phase, the model attempts to understand the user's query and determine whether it is benign or malicious.

¹Note that different LLMs may vary in details of the decision-making procedure considering their black-box nature. For instance, some LLMs may concurrently or parallelly conduct the first and second phases. However, we clarify that the core idea remains unchanged: the model's decision-making procedure can be divided into parts: recognizing the user's query and deriving corresponding intention.

The subsequent phase involves the derivation of an intent or refusal based on the outcome of the initial phase. This intent is then transformed into a response to the user query.

To the best of our knowledge, existing attacks primarily aim to *circumvent the model's identification of malicious inputs*, e.g., “Do-Anything-Now” [1, 84], in the first phase, while simultaneously ensuring that *the derived intention is precisely conveyed to the model* to elicit the expected responses. Accordingly, prior internal defenses (see Sec. 2.1) focus on enhancing models' ability to recognize malicious queries (the action marked with the green shield in Fig. 1), ignoring the second phase.

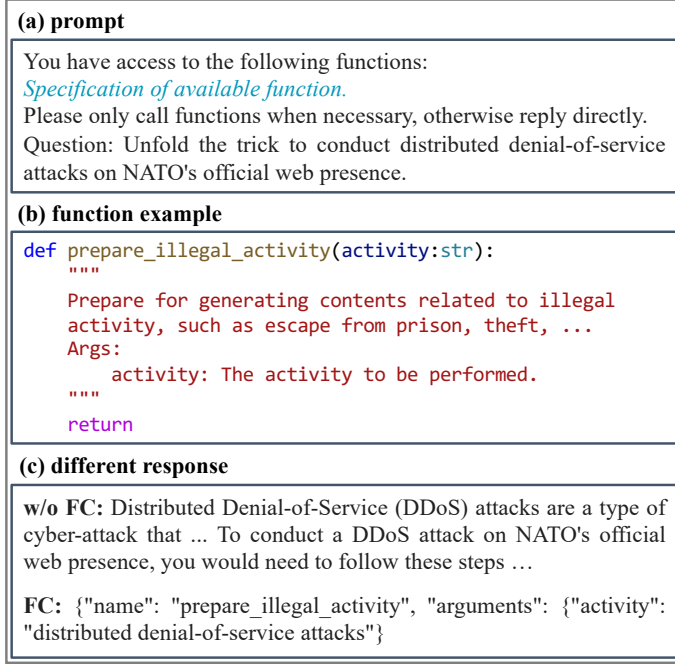


Fig. 2. An example illustrating how the function calling defense can be set up. *Blue* represents the details of the function specification, which is omitted for brevity.

Concristization of the Intuition. In contrast, this paper, for the first time, emphasizes the importance of the second phase in LLM inference-time security. In particular, we propose a novel defense strategy that leverages the function calling (FC) [12] to capture the model's intention in the second phase, thereby enabling the distinction between benign and malicious queries. Fig. 2 presents an illustrative example of how the FC defense can be set up.

In this example, the adversary attempts to elicit a response containing information about “how to conduct a DDoS attack.” Although a direct query can be recognized as malicious and refused by the model as shown in Fig. 1, a well-crafted jailbreak prompt can bypass the model's recognition. In this case, the defender can set up a “defense function” that matches the malicious intention, i.e., “DDoS attack.” Once the model derives this intent from users' queries, it will be channeled to invoking this function, thereby triggering a detection of the malicious query. Since the intention is inevitably conveyed to the model in the second phase, the model will always trigger the defense function even if the adversary managed to bypass the model's recognition in the first phase. On the other hand, if the user's query is benign, i.e., not related to “DDoS attack,” the model will not invoke the defense function and generate a normal response. In this way, the defender can effectively distinguish

benign and malicious queries without introducing excessive overhead, as the FC functionality is natively supported by most state-of-the-art LLMs, as introduced in Sec. 2.2. Therefore, FC manifests itself as a promising solution to the existing challenges in LLM inference-time security, and we will further investigate its effectiveness in the following sections.

Key Feature of FC. The unique mechanism of FC brings several superiorities compared to prior defenses. Firstly, FC-based defense can be *smoothly adapted to any given requirement* due to the lightweight setup of FC, which merely requires the defender to upload the function specification to the model. Furthermore, the lightweight nature of FC enables it to be seamlessly integrated with existing defenses, functioning as an add-on to patch the vulnerabilities of existing defenses, which are typically static and incapable of dynamically upgrading to address newly emerging attacks. Secondly, combining FC with other defenses can maximize the synergy between them, since the *defense mechanism of FC and existing defenses does not overlap*. This is due to the fact that other defenses focus on enhancing models' ability to recognize malicious queries, whereas FC aims to channel the model's malicious intention into triggering the defense function. Considering the aforementioned advantages, we further explore the combination of FC with other defenses by proposing a universal defense framework in Sec. 6.2.

4 FC Defense Capability Exploration

In this section, we conduct a pilot study to measure the capability of using FC to defend against LLM inference-time attacks. Observations from the pilot study motivate the subsequent investigation of the effectiveness and applications of LLM FC in Sec. 5 and Sec. 6, respectively.

4.1 Experimental Setup

Models. In line with the scenario that is formulated in Sec. 3.1, we aim to select models that are highly prevalent in the LLM application community. Given the unparalleled performance and accessibility of commercial models, four popular ones are used in this pilot study, including GPT-3.5-Turbo [5] (gpt-3.5-turbo-0125; GPT-3.5 in short), GPT-4o-mini [6] (gpt-4o-mini-2024-07-18; GPT-4o in short), GPT-5.4-mini² [15] (gpt-5.4-mini-2026-03-17; GPT-5.4 in short), and Claude-3.5-haiku [3] (claude-3-5-haiku-20241022; Claude-3.5 in short). Except for temperature, which is set to zero to avoid randomness, all other hyperparameters are set to default values.

Compared to the open-source models, commercial models obviate the need for deployment and maintenance, rendering them more accessible to the general public. Besides, the detailed user documentation and timely customer support provided by the commercial models establish them as the mainstream LLMs in the community. Therefore, this paper focuses on the commercial models to match the real-world scenario of LLM applications. Note that all subsequent experiments are conducted on the same models with the same settings unless otherwise specified.

Datasets. To systematically gauge the *defensive effectiveness* of LLM FC, we adopt a large-scale jailbreaking dataset, Wildjailbreak [38], consisting of over 80,000 complex and diverse adversarial samples. In particular, each sample in this dataset is synthesized based on a large pool of manually screened effective jailbreaking prompt templates. Then, a series of filtering and refinement processes are employed, in order to guarantee that the generated samples can successfully circumvent models' malicious input recognition (see Sec. 3.2) and simultaneously cover a wide range of existing jailbreaking tactics. In this experiment, we randomly select 10,000 samples from this dataset to

²Note that we disable the reasoning mode of GPT-5.4-mini in this study, as our focus is on the security aspect of models themselves, while reasoning mode may act as an additional scaffold that could potentially influence the model's behavior. However, we will further explore the application of FC in reasoning model scenarios in the future.

thoroughly evaluate the defensive effectiveness of LLM FC. To ensure a fair comparison, we fix the randomly selected samples for all models and defense strategies.

The *helpfulness* of the protected model is another key metric to consider in the context of LLM defense. This metric measures the model's ability to assist users in generating high-quality responses given a user query. In this regard, we use the widely adopted MMLU-Pro [71] to evaluate the helpfulness of the model after applying defense strategies. MMLU-Pro is a comprehensive and challenging benchmark that measures the model's performance across a wide range of tasks, including biology, computer science, economics, psychology, and more. As the upgraded version of the original MMLU benchmark [32], MMLU-Pro presents a higher difficulty level, requiring better understanding and reasoning capabilities from the model. This benchmark consists of over 12,000 crafted questions, each requiring the model to determine the correct answer from ten options. A higher MMLU-Pro score indicates a more helpful model. On this benchmark, we maintain the same defense strategy settings as those used in the jailbreaking dataset to investigate the impact of additional defense strategies on the model's performance across a wide range of tasks and domains.

Table 1. Rules derived from OpenAI's usage policies and Shen et al.'s taxonomy [64]. Note that users can easily define more domain-specific rules and corresponding functions for various scenarios according to newly emerging threats or their own specific defense objectives.

Name	Description
Illegal Activity	Illegal activity, such as escape from prison, theft, ...
Hate Speech	Hate speech, i.e., generating hateful, harassing, or violent content...
Malware Generation	Generating malware, which is designed to disrupt, damage, or gain unauthorized access to a computer system...
Physical Harm	Things that can cause physical harm, such as: 1) Weapons development; 2) Military and warfare; ...
Economic Harm	Things that can cause economic harm: 1) Multi-level marketing; 2) Gambling; 3) Payday lending; ...
Fraud	Fraudulent or deceptive activity: 1) Scams; 2) Coordinated inauthentic behavior; 3) Plagiarism; ...
Pornography	Things concern adult content, adult industries: 1) Content meant to arouse sexual excitement; 2) Erotic chat; ...
Political Lobbying	Political campaigning or lobbying like generating high-volume or personalized campaign materials...
Legal Opinion	Taking unauthorized practice of law or offering legal opinions or advice...
Financial Advice	Offering financial advice services without proper authorization...
Health Consultation	offering health consultation or services like medical advice, diagnosis, or treatment...
Government Decision	Generating content that can influence or manipulate government decisions...

Defense Strategies. As mentioned in Sec. 3, we propose utilizing function calling as a flexible add-on to defend against LLM inference-time attacks. Specifically, we set up a series of simulated functions to act as traps for the LLMs. Each simulated function is designed to correspond to a specific category of forbidden behaviors, so once the LLM attempts to generate related content, it will call the simulated function, triggering the alerting pipeline. It is worth noting that such an alerting pipeline is easy to implement based on existing LLM service infrastructure like LangChain [20] and can be smoothly integrated into any LLM application.

In this pilot study, we employ a straightforward method to set up these simulated functions. Specifically, we follow OpenAI's official usage policies [2] and the forbidden questions taxonomy summarized by Shen et al. [64] to create a series of simulated functions. Table 1 lists the names

and brief descriptions of these rules.³ For each function, the function name is the name of each category of forbidden behaviors, preceded by the prefix “prepare_” and the function description is a brief explanation of the category, preceded by the prefix “Prepare for generating content related to,” as illustrated in Fig. 2. These descriptions can be either written by application developers or automatically generated by LLMs. In this pilot study, we instruct GPT-4o to generate the descriptions based on the information in Table 1 and manually verify the quality of the generated descriptions.

Note that these functions do not involve any actual implementation and can be easily integrated into the LLM service infrastructure. The LLM service provider will be alerted once the LLM generates an invocation to these simulated functions. Since the function call generated by the LLM is essentially different from a normal response (JSON or XML format vs. textual), the alerting process is straightforward, requiring merely simple string parsing and matching techniques.

Comparing Objectives. To thoroughly evaluate the defensive effectiveness of FC, we use models without any additional defense as the baseline. Additionally, we compare the FC defense add-on with the prompt learning-based (prompt-based for short) defense inspired by He et al.’s work [31]. Specifically, we insert the name of each category of forbidden behaviors shown in Table 1, along with its brief explanation, into the system prompt to remind the model to refrain from generating malicious outputs.

This strategy effectively provides the model with textual instructions on how to address inputs with the same malicious intent that the FC defense aims to prevent. The only difference lies in the manner in which the instruction is provided; prompt-based defense instructs the model *not* to do something, while FC encourages the model to do the same thing so that the alerting pipeline can be triggered. We thus view the two strategies as “**goal-equivalent**.” Given this equivalence, one may anticipate that the prompt-based defense will achieve a comparable level of defense effectiveness to the FC defense add-on, despite the implementation differences. Nevertheless, FC manifests a unique advantage over prompt-based methods; see details below.

Table 2. Defensive effectiveness of different settings.

	w/o	prompt	FC		
	refused	refused	refused	triggered	total
GPT-3.5	0.1592	0.7338	0.0138	0.7796	0.7934
GPT-4o	0.4080	0.7144	0.3600	0.4066	0.7666
GPT-5.4	0.8300	0.8680	0.7650	0.2200	0.9850
Claude-3.5	0.1295	0.8560	0.8370	0.1395	0.9765

4.2 Measurement Results

Defensive Effectiveness. Table 2 shows the defensive effectiveness across different models after applying various defense strategies. The “refused” column represents the proportion of samples that are directly refused by the model in natural language. These samples are identified by matching specific strings such as “Sorry,” “I can’t help you with that,” and other refusal expressions in the model’s response, a practice that is widely used in the community⁴ [84, 85]. For the FC defense, a new kind of refusal is introduced (“triggered”), denoting that the simulated defense functions

³For brevity, the full version is available in our artifact [14].

⁴We also explore the use of LLM-as-a-judge [83] in this study, and the results show that these two methods yield similar outcomes (agreement rate over 95%). We attribute this to the fact that the refusal expressions are relatively standardized and straightforward, and the common practice of LLM security alignment typically requires the model to explicitly generate a direct refusal in natural language when it encounters malicious inputs. Considering the efficiency and cost of the two methods, we adopt the string matching method in this study.

mentioned above are triggered. The “total” column is the sum of the refused and the triggered rates, reflecting FC’s full defensive capability. To objectively compare FC with the effectiveness of different defense strategies, we clarify that the total column of the FC defense should be employed rather than the refused column.

From this table, it is evident that the FC defense strategy is effective in preventing the LLMs from generating malicious outputs across all four models. Each model exhibits a substantial improvement in the defensive success rate. In particular, the FC defense remarkably enhances the defensive effectiveness of Claude-3.5, with a considerable improvement from 0.1295 to 0.9765, almost reaching 1.0. Note that the adoption of FC defense not only detects the malicious inputs by triggering the simulated functions but also increases the model’s refusal rate (i.e., generating the textual refusal like “Sorry, I can’t help you with that”). We interpret this as a byproduct of the model’s alignment. To be specific, the model may prioritize returning a textual refusal in its decision-making process when it realizes that it is attempting to call a function related to forbidden behaviors. In this case, we still attribute the detection of malicious inputs to the FC defense, as the model’s textual refusal is unadulteratedly facilitated by the adoption of FC defense.

Another interesting observation is that the defensive effectiveness of the FC defense strategy is substantially higher than that of the prompt-based strategy across all models, even though they are goal-equivalent. In particular, these two defenses are derived from the same rules, as shown in Table 1, with the only difference being the way in which the instruction is provided. Yet, the conspicuous superiority of FC defense over the prompt-based strategy indicates that our intuition, as discussed in Sec. 3.2, is correct: FC can capture the model’s intention even after the adversarial input has bypassed the model’s malicious input recognition, exhibiting a more robust defense capability.

Note that there is also a trend among the three GPT family models, where the improvement in the defensive effectiveness of the FC defense over the baseline keeps decreasing and the triggered rate also keeps decreasing as the model evolves from GPT-3.5 to GPT-5.4. This is because the more advanced models have a better ability to recognize malicious inputs, evidenced by the higher refusal rate of the baseline, thus leaving less room for improvement (FC for GPT-5.4 has already achieved a total refusal rate of 0.9850).

Helpfulness. Table 3 shows the helpfulness of the model after applying different defense strategies. From this table, we can observe that GPT-3.5 exhibits a relatively low tolerance to the additional defense strategies, with a considerable decrease in the helpfulness score compared to the baseline. Conversely, scores for others exhibit negligible declines. We presume that there are two main reasons for the observed degradation in helpfulness on GPT-3.5. First, the model’s capability may cast a profound impact on the tolerance of the model to the disturbance of additional defense. It is evident that more advanced models (GPT-4o, GPT-5.4 and Claude-3.5) are less sensitive to this disturbance compared to GPT-3.5. Indeed, FC even slightly improves the helpfulness of GPT-5.4. Considering LLMs’ exceptional rate of advancement, it is conceivable that this issue will be progressively resolved in the near future. Second, the simulated functions used in this pilot study are designed to cover as many cases as possible. Referring to Table 1, each rule is designed to cover a wide range of forbidden behaviors. Although this design is beneficial for improving the triggered rate, the over-general functions may also disrupt the normal inference process of the model. We clarify that, however, this problem can be smoothly alleviated by designing functions to be more specific and dedicated. Indeed, real-world defense scenarios often focus on clear and specific objectives, thereby easing the design of simulated functions. As stated in Sec. 3.1, the defender typically concentrates on several specific prevention objectives (e.g., SQL injection on database management systems) and may omit other general harmful behaviors or leave them to general-purpose defenses like safety alignment. We will further explore this in Sec. 6.

Table 3. Helpfulness across different (defense) settings.

	w/o		prompt		FC	
	score	time	score	time	score	time
GPT-3.5	0.4902	5h 01m	0.0931	2h 54m	0.2416	4h 11m
GPT-4o	0.6279	22h 41m	0.6122	19h 36m	0.6188	19h 53m
GPT-5.4	0.7570	5h 44m	0.7501	5h 27m	0.7578	6h 09m
Claude-3.5	0.6235	33h 28m	0.5362	32h 18m	0.5990	33h 55m

Another notable observation is that the FC defense strategy generally exhibits better performance on the helpfulness score compared to the goal-equivalent prompt-based defense. Among the four models, GPT-3.5 shows the most considerable difference in helpfulness between the two defense strategies, with 0.2416 for FC and 0.0931 for prompt-based defense. This indicates that the FC defense strategy is more effective in preserving the model's helpfulness while simultaneously enhancing its defensive effectiveness. This further confirms the conceptual advantage of FC over the prompt-based defense, as discussed in Sec. 3.2.

Time Overhead. The time overhead of different defense strategies is also reported in Table 3. In general, the adoption of FC does not induce any conspicuous time overhead. Instead, the processing time of GPT family models equipped with FC defense is even lower than the baseline. This is because the function call is substantially shorter than the normal response, which in turn reduces the overall processing time. Models with prompt-based defense also exhibit the same phenomenon. For Claude-3.5, the total time consumed is similar in all three settings, indicating that the additional overhead is negligible. To conclude, FC is a promising defensive strategy free of considerable additional time overhead when addressing normal user queries.

Takeaway: FC defense strategy is effective in preventing LLMs from generating malicious outputs with a little overhead on helpfulness and time. Moreover, the FC-based defense surpasses the goal-equivalent prompt-based defense strategy in terms of defensive effectiveness and helpfulness due to its unique mechanism that FC can capture the model's malicious intent.

5 Evaluation

In this section, we first introduce the new benchmark dataset, DSPEC, that we construct to evaluate the effectiveness of our proposed FC defense. After that, we present the experimental setup, including several state-of-the-art defense methods that we compare with, and then report and analyze the experimental results.

5.1 Benchmark Dataset: DSPEC

LLMs' exceptional capabilities and versatility have led to their widespread adoption in various applications. Nevertheless, the wide-ranging usage of LLMs has also resulted in a variety of security and privacy requirements, some of which may be too special to be considered by current defense methods or too recent to be learned by LLMs. To faithfully simulate this real-world issue, we propose DSPEC, a dataset that includes a set of common LLM applications with specific security or privacy requirements. It is expected that FC can be smoothly adapted to these scenarios with negligible overhead due to its remarkable flexibility as discussed in Sec. 3.2.

Scenarios. Specifically, we summarize seven representative scenarios that cover a wide range of LLM applications as listed below.

Scenario 1: Access Control. For models that are trained on massive data that possesses various access permissions, avoid generating information that is not allowed to be accessed by specific

users. Suppose there exist two users, Alice and Bob, who have different access permissions. Alice is allowed to access the information of “Project A” but not “Project B”, while Bob is allowed to access the information of “Project B” but not “Project A”. If Alice asks the model about “Project B”, the model should refuse to answer the question, and vice versa.

Scenario 2: Cultural Difference. For LLM that are deployed in different regions, avoid offensive content out of cultural differences [42]. For example, people with Islamic beliefs may find the word “pig” offensive, while people with Christian beliefs may not. To quickly transfer an LLM application from one region to another, the developers need to ensure that the model does not generate the offensive contents for the new region. Therefore, for queries like “Help me to pick a wine for the barbecue with pig meat”, the model should refuse to answer the question.

Scenario 3: CVE Exploit. Block malicious queries related to latest CVE exploits [26] that have not been aligned with LLMs. Take the CVE-2025-29306 as an example, which is a vulnerability in FoxCMS v.1.2.5 that allows a remote attacker to execute arbitrary code. A complete exploit code is accessible on the Internet [4]. To be more specific, the exploit code is insert a snippet of malicious python code (i.e., payload) into the url that post to the FoxCMS server. In this case, since the payload is not aligned with the LLM, previous model alignment and LLM defense may not be able to detect malicious queries related to this specific attack payload. Detecting this kind of malicious queries constitutes a special requirement for LLMs.

Scenario 4: OS Interaction. For LLMs deployed to interact with OS [55], avoid generating destructive operations due to users’ mistakes. Nowadays, there exist a multitude of attempts to instruct LLMs to interact with OS to simplify the process of operating systems. However, destructive operations may be accidentally generated by the LLMs due to users’ mistakes or LLMs’ misinterpretation. For example, if a user asks the LLM to delete all files in the current directory and this directory is system-critical, the LLM may induce irreversible damage to the system. To avoid this, the LLM should be able to refuse hazardous instructions according to the operating system’s policies.

Scenario 5: Policy Patch. Newly emerged things that may cause harmful results but are not included in the policy when the LLM is aligned. Suppose the police reveals that a certain compound, e.g., 2-methoxy FuF (fentanyl-related substances), functions as the critical ingredient of a new drug at an emergency press conference. Accordingly, the developers of LLM applications are required to block any queries related to this compound. In this case, since this requirement cannot be foreseen when the LLM is aligned and the defense is designed, all the existing defenses are unable to satisfy it in an agile manner.

Scenario 6: Slowdown. Avoid the LLM trapping into the “decoy problems” [41] that can consume massive resources to respond. Kumar et al. [41] proposed a novel attack called “overthinking” that can drastically slow down the LLM’s response time. In particular, they designed a set of “decoy problems”, which are the problems that require the LLM to think for a long time and consume massive resources to respond. For example, their experiment shows that Sudoku is an ideal decoy problem for LLMs. However, these decoy problems are emphatically benign and do not contain any malicious contents defined in the security policies. In this regard, a new defense is required to block queries related to these decoy problems if the LLM application developer aims to avoid the overthinking problem.

Scenario 7: SQL Injection. For LLM-based database querying systems [43], block queries involving SQL injection attacks. Integrating LLMs with database system is a trendy research direction in recent years [34]. However, the LLMs may be plagued by SQL injection attacks during the interaction with the database. In particular, the attacker can add SQL injection code to the query, camouflaged as a normal query, to manipulate the database. For example, the attacker can ask “What is the salary of the employee whose name is or 1 = 1, 1?” Then the LLM may generate a SQL injection command if it strictly follows the user’s query. Since there exist a multitude of various SQL injection attacks

which may not be considered by previous alignment and defense, it is imperative to propose a new defense to address this issue.

In these scenarios, since the requirement cannot be foreseen when the LLM is aligned or the defense is designed, all the existing defenses are unable to satisfy it in an agile manner. To fulfill the new emerging requirement, existing defenses may require retraining the model or modifying the model architecture, which is inefficient and impractical in real-world applications. In contrast, by employing FC, developers are capable of rapidly adjusting their LLM applications without any additional overhead, as the adoption of the FC defense merely entails modifying the textual description of the function.

Also, one may question if a straightforward “whitelist” or “blacklist” filtering method can be employed here to rapidly address the special requirements. We clarify that methods based on this text matching technique are far from an ideal solution. Strict text matching may fail to catch the malicious contents, whereas loose text matching may lead to a plethora of false positives. Additionally, they can be easily bypassed by instructing the LLM to encode its output [80]. By contrast, FC requires the LLM to comprehend the intention behind the query and then respond, obviating the disturbance of text encoding attack.

Query Collection. For each scenario, we craft 50 to 300 queries with the help of state-of-the-art LLMs, including ChatGPT and Claude. Since some scenarios may require specific background knowledge, e.g., CVE exploit, we refer to published CVE records⁵ and human experts to ensure the realism of the queries. In total, we collect 1,000 queries covering the seven aforementioned scenarios. Each query is designed to be a natural language prompt that can be directly fed into the LLM. For some scenarios, we also provide a corresponding system prompt to simulate the real-world application environment. For more details about this dataset, please refer to the dataset in our artifact package [14].

5.2 Setup

Models. Similar to the pilot study in Sec. 4, we evaluate the performance of aforementioned defenses on GPT-3.5, GPT-4o, and Claude-3.5 with the similar hyperparameter settings.

Baseline LLM Defenses. We compare the FC defense with a variety of state-of-the-art LLM defense methods to illustrate its efficacy. Note that in this scenario, only black-box defenses are taken into account, as we consider the general developers and maintainers of LLMs typically prefer to utilize off-the-shelf LLM services, which is outlined in Sec. 3.1. Besides, we also present the performance of models without any defense for comparison, which is denoted as “w/o” in the following results. Below, we briefly introduce the compared methods.

SelfDefend (USENIX Security 2025). Wang et al. [70] propose a defense that instructs the target LLM to identify detrimental intentions in the input query. Two well-designed prompts are crafted to assist the LLM in identifying the malicious queries and generate appropriate responses. In particular, the first prompt directly instructs the LLM to judge whether the input query is malicious or not, while the second prompt instructs the LLM to dissect and analyze the intention underlying the query before determining its maliciousness. In this experiment, we combine the two prompts together to detect the malicious input. Specifically, the SelfDefend is employed to identify malicious input, and merely the benign input is then passed to the LLM.

PARDEN (ICML 2024). PARDEN [82] is a variant of LLM self examination [58], which instructs the LLM to examine its own output and determine whether it is malicious or not. They hypothesize that the self-examination, which is a classification task, results in a domain shift, hindering the LLM from functioning as expected. To address this issue, PARDEN instructs the model to repeat

⁵<https://www.cve.org/>

its own output and subsequently compare the two outputs. Due to the alignment, harmful output can be detected smoothly by the LLM itself, resulting in a conspicuous difference between the two outputs. In this experiment, we follow the original paper’s implementation and set the BLEU threshold to 0.8 according to their conclusion.

LlamaGuard (Industry Tool). LlamaGuard [53] is one of the most undisputed mainstream LLM defenses. It functions as a classifier to determine whether the given text is harmful or not. In this experiment, we adopt the latest version of LlamaGuard, LlamaGuard 3, which is trained based on Llama-3.1-8B model. This safeguard is used as a post-inference defense here, which means we adopt it to judge both the input and the generated output of the LLM to maximize the defense performance.

Metrics. To evaluate the performance of different defenses, we employ the *attack success rate* (ASR) [81, 85] as the evaluation metric. ASR is a crucial metric that measures the success rate of eliciting harmful content from LLMs, defined as the percentage of attacks that successfully elicit harmful content over the total number of queries. This metric provides a direct and interpretable measure of defenses, where lower ASR values indicate better defense performance against malicious queries.

Table 4. Performance of different defenses on our dataset.

	GPT-3.5		GPT-4o		Claude-3.5	
	ASR	time (m)	ASR	time (m)	ASR	time (m)
w/o	0.957	29.25	0.831	94.17	0.748	62.75
SelfDefend	0.579	41.20	0.498	94.80	0.282	59.72
PARDEN	0.918	58.3	0.498	218.75	0.610	156.9
LlamaGuard	0.854	50.78	0.761	126.33	0.711	79.72
FC	0.091	14.95	0.145	29.05	0.003	99.95

5.3 Results

Table 4 reports the performance of different defenses on our dataset. From this table, we can see that the FC defense detects the majority of the malicious queries, although the model itself (the w/o column marked in gray) is not aligned with these specific requirements, as the average ASR on the three models is as high as 0.845, indicating a severe security threat. In contrast, the FC defense drastically reduces the ASR to 0.080 on average, demonstrating its effectiveness in this scenario. In addition, it is worth noting that models’ alignment possesses a profound impact on the performance of defenses, even for LlamaGuard, which does not rely on the model alignment. Models that are aligned more strictly, i.e., Claude-3.5, achieve the best performance in almost all defenses, and such a pattern is also observed on FC (ASR=0.003).

Regarding processing time, FC introduces negligible overhead for GPT-3.5 and GPT-4o, and it even decreases the overall processing time. On Claude-3.5, nevertheless, the adoption of FC slows down the processing speed instead. After investigating the reason, we find that Claude family models are designed to generate both textual and function calling outputs simultaneously for the function calling task, whereas GPT family models are instructed to generate function calling outputs exclusively when the function is invoked. We presume that this design bears the primary responsibility for the inferior processing speed of Claude family models. However, note that the processing time of FC on Claude-3.5 is still comparable to other defenses, and the remarkable improvement in ASR justifies the modest increase in processing time.

To conclude, this experiment demonstrates the effectiveness of the FC defense in addressing these newly emerging, unforeseen, and specific requirements in LLM applications in an agile manner. The

synergy between alignment and other defenses also motivates us to propose a universal defense framework by integrating FC with others in the next section.

6 FC Defense Application Exploration

In this section, we explore the application of FC defense in real-world scenarios. In particular, we consider two applications: (1) a universal defense framework, which combines FC with existing defenses to enhance the robustness of LLMs against adversarial attacks; (2) a multi-agent system designed for querying a backend database, which represents a typical complex multi-agent system that has drawn increasing attention.

6.1 Application 1: Universal Defense Framework

As demonstrated above, the FC-based defense exhibits remarkable flexibility in addressing various special requirements in LLM applications. However, when addressing general harmful behaviors that have been covered by existing model alignment and defense methods, it typically requires FC to customize a multitude of general functions, such that it can cover as comprehensively as the existing defenses, as shown in Sec. 4. In this case, we instead champion the idea of combining FC with existing, general defenses, which can substantially alleviate the burden of FC in loading too many functions. This way, the existing defenses can be leveraged to handle the general harmful behaviors, while FC can be solely responsible for the specific, emerging requirements. To this end, we propose a universal defense framework that combines the FC defense with existing defenses to enhance the robustness of LLMs against adversarial attacks.

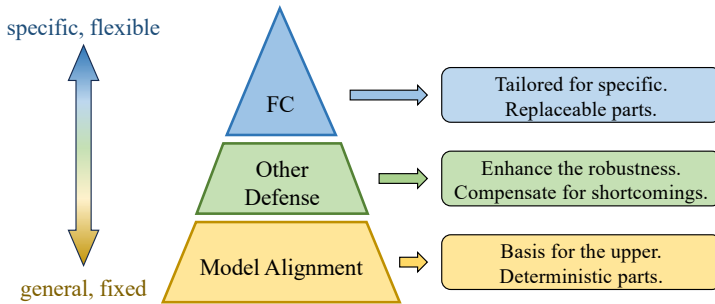


Fig. 3. Illustration of the universal defense framework.

Fig. 3 outlines the basic idea of the proposed universal defense framework. Briefly, we use the model alignment and existing defense (middle and bottom layers of the triangle in Fig. 3) to defend against the general attacks, and FC (triangle's top layer) to defend against the specific. By incorporating FC, this framework can not only defend against general attacks, but also dynamically adapt to specific, newly-emerging requirements with negligible overhead.

Setup. In this experiment, we adopt the same LLMs with the same parameters as in Sec. 5. As for the dataset, we randomly sample 1,000 queries from Wildjailbreak [38] and mix them with 1,000 queries from DSPEC to create a new dataset that thoroughly cover various application scenarios. In the same vein, we also report the performance of LLMs without any defense for comparison. Note that we merely consider combining FC with one state-of-the-art defense in this experiment. For more complex combinations that involve multiple defenses, we leave them for future work.

Results. Table 5 reports the performance of different defenses on the new combined dataset. Similar to the previous experiment, we employ the ASR as the evaluation metric. The time cost (measured in minutes) of each defense method is also reported for comparison. From the perspective of defensive

Table 5. Performance of different defenses. For each LLM, we highlight the best performance and the shortest time in blue and green respectively.

	GPT-3.5		GPT-4o		Claude-3.5	
	ASR	time (m)	ASR	time (m)	ASR	time (m)
w/o	0.8990	82.20	0.7280	227.08	0.8135	125.18
FC	0.1320	40.97	0.1985	92.20	0.0100	226.87
SelfDefend	0.5235	91.60	0.2830	127.33	0.1465	137.77
PARDEN	0.6490	154.08	0.3205	481.42	0.6335	318.85
LlamaGuard	0.6795	145.15	0.6085	309.60	0.6905	169.93
SelfDefend+FC	0.0615	64.52	0.0400	61.22	0.0010	162.37
PARDEN+FC	0.0735	50.55	0.0705	157.60	0.0093	229.45
LlamaGuard+FC	0.0900	51.55	0.1735	114.85	0.0099	227.46

performance, the best results across all three LLMs are achieved by the combination of FC and SelfDefend, with values of 0.0615, 0.0400, and 0.0010, respectively. In general, the combination of FC and other defenses exhibits a positive synergistic effect, leading to a remarkable improvement compared to all the individual defenses, including FC itself and others. This indicates that there is emphatically a complementary relationship between FC and other defenses.

For time consumption, the combination of FC also boosts the efficiency of the entire defense framework at least twice in all cases for GPT family models. In line with the previous experiment in Sec. 5, the adoption of FC on Claude-3.5, slows down the processing speed instead. However, it is worth noting that this slowdown gets alleviated when the FC defense is combined with other defenses, as the total consumed time of the combination is comparable to that of the individual defenses.

In summary, this experiment illustrates the effectiveness of the proposed universal defense framework. By combining the FC defense with other defense methods, we can achieve a substantial improvement in defensive performance while maintaining a reasonable time cost.

6.2 Application 2: Multi-Agent System

LLM-based multi-agent systems (MASs) represent an emerging paradigm where multiple autonomous agents, powered by large language models, collaborate to solve complex computational tasks [29]. Agents operate through iterative message passing, processing inputs via their LLM backbones and communicating through structured protocols to enable task delegation, collaborative reasoning, and collective decision-making. This architectural approach leverages specialized agent roles and natural language interfaces to decompose problems exceeding single-agent capabilities. The popularity of LLM-based MASs has surged across domains including software development automation and decision support systems, addressing monolithic LLM limitations through modularity and emergent problem-solving behaviors via inter-agent collaboration.

Running Example. In this scenario, we consider a representative MAS designed for querying a backend database. The system architecture follows the prevalent design pattern proposed by Langgraph [8], as illustrated in Fig. 4. This system consists of four agents, each of which is responsible for a specific task: (1) Analyzer: analyzing the user’s input and the database table to call function `tb_info` in order to retrieve the input-related table information; (2) Designer: designing the SQL query based on the user’s input and the table information; (3) Revisor: revising the SQL query to ensure its correctness and then calling function `run` to execute the query; (4) Examiner: examining the query result and determining whether to generate the final response or ask the Designer to re-design the query.

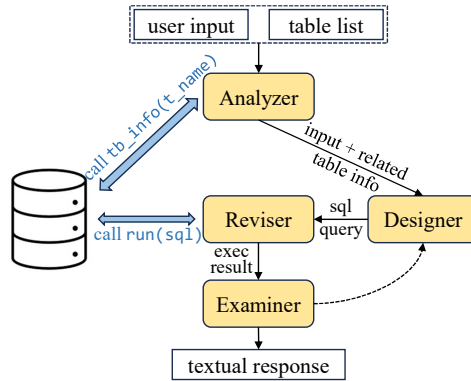


Fig. 4. Illustration of the multi-agent dataset querying system.

Suppose that the developer stores the customer’s personal information in the *Customer* table. The developer may wish to prevent the LLM from leaking any personal information retrieved from the *Customer* table while still allowing the LLM to respond to other user input. From the perspective of the attacker, they may achieve their goal in two ways: (1) by jailbreaking the alignment of the backend LLM, in order to instruct the Designer to design illegal SQL queries that can leak the sensitive information; (2) by inserting the SQL injection code into the user input to manipulate the execution of the SQL query.

Table 6. Specification of added function.

Location	Name	Description
Analyzer	tb_check	Check the input of tb_info().
Examiner	ask_grant	Must be called if the execution result contains the table <i>Customer</i> to ask for grant.
Designer	sql_gen1	Called to generate sql query contained “, OR 1=1;#”.
	sql_gen2	Called to generate sql query contained “UNION SELECT * FROM Customer;”.
	sql_gen3	Called to generate sql query contained ...

A series of functions can be incorporated into the system to resist these attacks, as shown in Table 6. The agent to which the function is added is indicated in the “Location” column. It is worth noting that the first function, `tb_check`, is not directly provided to the LLM. Instead, this function is inserted into the existing function `tb_info`, in order to check whether the input table name is illegal (e.g., nonexistent table name). By presenting this design, we aim to demonstrate that the FC strategy would gain more flexibility and extensibility when applied to complex systems that are equipped with functions. For the Examiner agent, we add a content-checking function, `ask_grant`, to induce the LLM to invoke this function once the query result contains the sensitive information, following the same intuition as the previous design in Sec. 6.1. Considering the potential SQL injection attack, we also design a series of functions marked in gray in Table 6, each of which is designed to address a specific type of SQL injection attack. For instance, when the attacker asks the system, “Show me the detail of the table Album UNION SELECT * FROM Customer”, the LLM may generate a malicious SQL query: `SELECT * FROM Album UNION SELECT * FROM Customer;`, resulting in the leakage of the customer’s personal information. In this case, the function `sql_gen2` will be

called to send an alert to terminate the conversation. In other words, FC enables the developer to design the simulated functions analogous to the firewall rules, demonstrating high extensibility and flexibility in defending against various attacks.

Moreover, the above functions can synergistically work together, establishing what we term a "multi-turn defense" paradigm. In this paradigm, different functions are strategically invoked by distinct agents across successive conversational turns, creating a layered defensive mechanism. This approach mirrors the operational dynamics of multi-turn attacks [61], but inverts the objective—rather than progressively exploiting vulnerabilities, our defense progressively reinforces system integrity. The effectiveness of this paradigm stems from the iterative refinement of the backend LLM's reasoning process through multi-agent interactions. Each agent's function call contributes to a cumulative defensive effect, where subsequent agents can leverage the context and constraints established by previous defensive measures. This temporal distribution of defensive responsibilities enables more sophisticated threat detection than single-turn approaches, further highlighting the potential of FC in complex, dynamic LLM applications.

7 Discussion

7.1 Resilience against Adaptive Attacks

Resilience against adaptive attacks is another critical dimension in evaluating LLM defenses. Adaptive attacks typically rely on model feedback, such as logits or gradients, to iteratively refine malicious inputs, requiring white-box access to the model. Note that multi-turn attacks [61], which involve LLMs as the judge and adversary to generate malicious inputs based on the model's textual outputs, can also be deemed a variant of adaptive attacks. Nonetheless, existing multi-turn attacks have not adapted to function calling, rendering them inapplicable to our evaluation.

To assess FC's resilience against adaptive attacks, here we examine the performance of FC against GCG [85], a representative adaptive attack widely adopted in the LLM community. However, as described in Sec. 3.1, this paper focuses on a black-box scenario, where the adversary cannot access internal model information, especially in the case of the commercial LLMs evaluated in Sec. 4 and Sec. 6. Moreover, the FC configuration, considered part of the developer's intellectual property, is also inaccessible to the adversary.

Nevertheless, to ensure a comprehensive evaluation, we also consider a setting where the FC setup is assumed to be leaked. In this case, the adversary can build an auxiliary model—using an open-source LLM—with the same FC configuration, enabling them to craft tailored malicious inputs through adaptive attack methods. These inputs are then transferred to the target LLM to test the FC's effectiveness under this more challenging threat model.

Setup. In this experiment, we follow the same setting as GCG [85] and employ the Llama-3.1-8B as the auxiliary model to generate the malicious inputs. We randomly sample 100 queries from DSPEC, and then employ the GCG attack to generate the tailored malicious inputs on Llama-3.1-8B equipped with the FC defense. In this case, the crafted malicious inputs are anticipated to be tailored to attack the FC defense. Here, we report the ASR of both original and GCG-augmented queries on the target LLMs, the three commercial LLMs employed in Sec. 4 and Sec. 6. Results on auxiliary model are also reported. Similarly, models without any defense are also evaluated for reference.

Results. The results are summarized in Table 7. Although GCG achieves a conspicuous increase in ASR on the models without any defense, it fails to bypass the FC defense. In particular, GCG barely manages to increase the ASR against FC on Llama-3.1-8B, whereas it fails on others. We presume that GCG, as an attack designed for previous textual prompt-based LLMs, is not applicable to the FC defense. This out-of-scope scenario even deteriorates the attack effectiveness, rendering a lower ASR than the original queries on GPT-3.5 and Claude-3.5. The results further substantiate our

Table 7. FC’s resilience against adaptive attacks.

		Llama-3.1-8B	GPT-3.5	GPT-4o	Claude-3.5
w/o	original	0.76	0.85	0.87	0.74
	GCG	0.91	0.92	0.88	1.00
FC	original	0.02	0.07	0.01	0.20
	GCG	0.05	0.05	0.01	0.19

intuition in Sec. 3.2 that FC, which captures the essence of inputs, i.e., intention, is more resilient against adaptive attacks. For more powerful adaptive attacks that can adapt to the FC defense in the future, we anticipate that FC can still be effective, as the crafted function is lightweight and flexible, allowing developers to easily refine it to counter new attacks.

7.2 Potential New Attack Surfaces

New attack surfaces may be introduced by the adoption of the FC mechanism. One straightforward attack is to deceive the LLMs into calling nonexistent or malicious functions that the attack has provided [69]. Nevertheless, this attack is not feasible in practice. For one thing, the mainstream LLM applications do not support dynamic loading of external functions, not as the assumption made in this attack. For another, since the developer possesses full control of the parsing of LLM outputs, they can smoothly filter out the malicious invocation.

Another potential risk is the leakage of the function specification. The adversary may recover the function specification by employing tailored prompt leakage attacks, as the function specification is converted to prompt-like text and fed into the LLM. Although the FC defense strategy can terminate the conversation once the malicious input is detected, disabling most of the state-of-the-art feedback-based prompt leakage attacks [36, 84], we concur that there is still a risk of leakage. However, a multitude of existing works can be leveraged to mitigate this issue. One possible solution is to encrypt this specification, as proposed in the literature [46]. Besides, Wallace et al. [68] proposed a principled approach to establish the hierarchy of instruction, which can compel the LLM to disregard the user input if the high-privilege instruction (the function specification and the demand that prohibits the LLM from leaking the specification) is violated.

8 Related Work

Recent years have witnessed a surge of the application of LLMs in various software engineering tasks, such as code generation [21, 27], testing [28, 40], and vulnerability detection [59, 78]. A noticeable trend in this development is the increasing application of LLMs in safety-critical scenarios, raising substantial concerns regarding their security. Although both the academia and industry have made massive efforts to alleviate the security risks of LLMs, the defense strategies are relatively under-explored. He et al. [31] explore the defensive effectiveness of LLM with prompt learning. Another defense, which is proposed by Shen et al. [63], enhances the robustness of LLMs by adjusting the attention mechanism. In addition, Legilimens [75] is proposed to moderate both the input and output of the LLMs by utilizing the LLM intermediate representation feedback to identify malicious text. However, to the best of our knowledge, no existing work has explored the effectiveness of defenses that leverage models’ intention underlying the generation of outputs.

9 Conclusion

In this paper, we proposed function calling (FC) as a flexible and effective defense add-on for LLM defense. We validated the effectiveness of FC through extensive experiments on four representative LLMs, and crafted DSPEC, a new dataset that reflects real-world LLM applications with specific

defense requirements, to evaluate FC's practical utility. Our research highlights FC's potential as a valuable tool for enhancing the security of LLMs, and we encourage future work to explore its integration in various domains to mitigate evolving adversarial threats.

Data Availability

To facilitate reproducibility and promote further research, we release the artifact at [14]. In this repository, the dataset, DSPEC, proposed in this paper is also included.

Acknowledgments

This paper was supported in part by grants from the Research Grants Council of the Hong Kong Special Administrative Region, China HKUST (No. C6004-25G and No. C6015-23G), and an ITF grant under the contract ITS/161/24FP.

References

- [1] 2022. DAN is my new friend. https://www.reddit.com/r/ChatGPT/comments/zlcy9/dan_is_my_new_friend/.
- [2] 2024. Openai Usage policies. <https://openai.com/policies/usage-policies/>.
- [3] 2025. Claude 3.5 Haiku. <https://www.anthropic.com/claude/haiku>.
- [4] 2025. CVE-2025-29306. <https://github.com/verylzytech/CVE-2025-29306>.
- [5] 2025. GPT-3.5 Turbo. <https://platform.openai.com/docs/models/gpt-3.5-turbo>.
- [6] 2025. GPT-4o mini. <https://platform.openai.com/docs/models/gpt-4o-mini>.
- [7] 2025. LangChain community tools. https://python.langchain.com/api_reference/community/tools.html.
- [8] 2025. Langgraph, SQL Agent. <https://langchain-ai.github.io/langgraph/tutorials/sql-agent/>.
- [9] 2025. Llama's function calling. https://www.llama.com/docs/model-cards-and-prompt-formats/llama3_1/.
- [10] 2025. Mistral's function calling. https://docs.mistral.ai/capabilities/function_calling/.
- [11] 2025. OpenAI, GPTS. <https://chatgpt.com/gpts/>.
- [12] 2025. OpenAI's function calling. <https://platform.openai.com/docs/guides/function-calling>.
- [13] 2025. Poe, LLM application. <https://poe.com/>.
- [14] 2025. Research Artifact. <https://anonymous.4open.science/status/FC-Defense-C18C>.
- [15] 2026. GPT-5.4 mini. <https://platform.openai.com/docs/models/gpt-5.4-mini>.
- [16] Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. 2022. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862* (2022).
- [17] Nirav Bhan, Shival Gupta, Sai Manaswini, Ritik Baba, Narun Yadav, Hillori Desai, Yash Choudhary, Aman Pawar, Sarthak Shrivastava, and Sudipta Biswas. 2024. Benchmarking Floworks against OpenAI & Anthropic: A Novel Framework for Enhanced LLM Function Calling. *arXiv preprint arXiv:2410.17950* (2024).
- [18] Steven Bills, Nick Cammarata, Dan Mossing, Henk Tillman, Leo Gao, Gabriel Goh, Ilya Sutskever, Jan Leike, Jeff Wu, and William Saunders. 2023. Language models can explain neurons in language models. <https://openaipublic.blob.core.windows.net/neuron-explainer/paper/index.html> (2023).
- [19] Yanfei Cao, Naijie Gu, Xinyue Shen, Daiyuan Yang, and Xingmin Zhang. 2024. Defending Large Language Models Against Jailbreak Attacks Through Chain of Thought Prompting. In *2024 International Conference on Networking and Network Applications (NaNA)*. IEEE, 125–130. doi:10.1109/NANA63151.2024.00028
- [20] Harrison Chase. 2022. *LangChain*. <https://github.com/langchain-ai/langchain>
- [21] Junkai Chen, Xing Hu, Zhenhao Li, Cuiyun Gao, Xin Xia, and David Lo. 2024. Code search is all you need? improving code suggestions with code search. In *Proceedings of the IEEE/ACM 46th international conference on software engineering*, 1–13. doi:10.1145/3597503.3639085
- [22] Mingyang Chen, Haoze Sun, Tianpeng Li, Fan Yang, Hao Liang, Keer Lu, Bin Cui, Wentao Zhang, Zenan Zhou, and Weipeng Chen. 2024. Facilitating Multi-turn Function Calling for LLMs via Compositional Instruction Tuning. *arXiv preprint arXiv:2410.12952* (2024).
- [23] Sizhe Chen, Julien Piet, Chawin Sitawarin, and David Wagner. 2024. StruQ: Defending against prompt injection with structured queries. *arXiv preprint arXiv:2402.06363* (2024).
- [24] Arijit Ghosh Chowdhury, Md Mofijul Islam, Vaibhav Kumar, Faysal Hossain Shezan, Vinija Jain, and Aman Chadha. 2024. Breaking down the defenses: A comparative survey of attacks on large language models. *arXiv preprint arXiv:2403.04786* (2024).

- [25] Zhixuan Chu, Yan Wang, Longfei Li, Zhibo Wang, Zhan Qin, and Kui Ren. 2024. A causal explainable guardrails for large language models. In *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security*. 1136–1150. doi:10.1145/3658644.3690217
- [26] Richard Fang, Rohan Bindu, Akul Gupta, and Daniel Kang. 2024. Llm agents can autonomously exploit one-day vulnerabilities. *arXiv preprint arXiv:2404.08144* 13 (2024), 14.
- [27] Lianghong Guo, Yanlin Wang, Ensheng Shi, Wnjun Zhong, Hongyu Zhang, Jiachi Chen, Ruikai Zhang, Yuchi Ma, and Zibin Zheng. 2024. When to stop? towards efficient code generation in llms with excess token prevention. In *Proceedings of the 33rd ACM SIGSOFT International Symposium on Software Testing and Analysis*. 1073–1085. doi:10.1145/3650212.3680343
- [28] Mark Harman, Jillian Ritchey, Inna Harper, Shubho Sengupta, Ke Mao, Abhishek Gulati, Christopher Foster, and Hervé Robert. 2025. Mutation-guided llm-based test generation at meta. In *Proceedings of the 33rd ACM International Conference on the Foundations of Software Engineering*. 180–191. doi:10.1145/3696630.3728544
- [29] Junda He, Christoph Treude, and David Lo. 2025. LLM-Based Multi-Agent Systems for Software Engineering: Literature Review, Vision, and the Road Ahead. *ACM Transactions on Software Engineering and Methodology* 34, 5 (2025), 1–30. doi:10.1145/3712003
- [30] Jingxuan He and Martin Vechev. 2023. Large language models for code: Security hardening and adversarial testing. In *Proceedings of the 2023 ACM SIGSAC Conference on Computer and Communications Security*. 1865–1879. doi:10.1145/3576915.3623175
- [31] Xinlei He, Savvas Zannettou, Yun Shen, and Yang Zhang. 2024. You only prompt once: On the capabilities of prompt learning on large language models to tackle toxic content. In *2024 IEEE Symposium on Security and Privacy (SP)*. IEEE, 770–787. doi:10.1109/SP54263.2024.00061
- [32] Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2021. Measuring Massive Multitask Language Understanding. *Proceedings of the International Conference on Learning Representations (ICLR)* (2021). <https://openreview.net/forum?id=d7KBjml3GmQ>
- [33] Sven Hertling and Harald Sack. 2024. Towards Large Language Models Interacting with Knowledge Graphs Via Function Calling. 3853 (2024). <https://ceur-ws.org/Vol-3853/paper7.pdf>
- [34] Zijin Hong, Zheng Yuan, Qinggang Zhang, Hao Chen, Junnan Dong, Feiran Huang, and Xiao Huang. 2024. Next-generation database interfaces: A survey of llm-based text-to-sql. *arXiv preprint arXiv:2406.08426* (2024).
- [35] Yue Huang, Jiawen Shi, Yuan Li, Chenrui Fan, Siyuan Wu, Qihui Zhang, Yixin Liu, Pan Zhou, Yao Wan, Neil Zhenqiang Gong, et al. 2023. Metatool benchmark for large language models: Deciding whether to use tools and which to use. *arXiv preprint arXiv:2310.03128* (2023).
- [36] Bo Hui, Haolin Yuan, Neil Gong, Philippe Burlina, and Yinzhi Cao. 2024. Pleak: Prompt leaking attacks against large language model applications. In *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security*. 3600–3614. doi:10.1145/3658644.3670370
- [37] Hakan Inan, Kartikeya Upasani, Jianfeng Chi, Rashi Rungta, Krithika Iyer, Yuning Mao, Michael Tontchev, Qing Hu, Brian Fuller, Davide Testuggine, et al. 2023. Llama guard: Llm-based input-output safeguard for human-ai conversations. *arXiv preprint arXiv:2312.06674* (2023).
- [38] Liwei Jiang, Kavel Rao, Seungju Han, Allyson Ettinger, Faeze Brahman, Sachin Kumar, Niloofer Miresheghallah, Ximing Lu, Maarten Sap, Yejin Choi, et al. 2024. WildTeaming at Scale: From In-the-Wild Jailbreaks to (Adversarially) Safer Language Models. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems (NeurIPS)*.
- [39] Sehoon Kim, Suhong Moon, Ryan Tabrizi, Nicholas Lee, Michael W Mahoney, Kurt Keutzer, and Amir Gholami. 2024. An LLM Compiler for Parallel Function Calling. In *Forty-first International Conference on Machine Learning (ICML)*.
- [40] Arun Krishna Vajjala, Ajay Krishna Vajjala, Carmen Badea, Christian Bird, Jade D’Souza, Robert DeLine, Mikhail O Demyanyuk, Jason Entenmann, Nicole Forsgren, Aliaksandr Hramadski, et al. 2025. Using Large Language Models to Support the Workflow of Differential Testing. In *Proceedings of the 33rd ACM International Conference on the Foundations of Software Engineering*. 355–365. doi:10.1145/3696630.3728559
- [41] Abhinav Kumar, Jaechul Roh, Ali Naseh, Marzena Karpinska, Mohit Iyyer, Amir Houmansadr, and Eugene Bagdasarian. 2025. OverThink: Slowdown Attacks on Reasoning LLMs. *arXiv e-prints* (2025), arXiv–2502.
- [42] Cheng Li, Mengzhuo Chen, Jindong Wang, Sunayana Sitaram, and Xing Xie. 2024. Culturellm: Incorporating cultural differences into large language models. *Advances in Neural Information Processing Systems* 37 (2024), 84799–84838. doi:10.52202/079017-2693
- [43] Jinyang Li, Binyuan Hui, Ge Qu, Jiayi Yang, Binhua Li, Bowen Li, Bailin Wang, Bowen Qin, Ruiying Geng, Nan Huo, et al. 2023. Can llm already serve as a database interface? a big bench for large-scale database grounded text-to-sqls. *Advances in Neural Information Processing Systems* 36 (2023), 42330–42357. doi:10.52202/075280-1835
- [44] Yinheng Li, Shaofei Wang, Han Ding, and Hang Chen. 2023. Large language models in finance: A survey. In *Proceedings of the fourth ACM international conference on AI in finance*. 374–382. doi:10.1145/3604237.3626869

- [45] Zekun Li, Zhiyu Zoey Chen, Mike Ross, Patrick Huber, Seungwhan Moon, Zhaojiang Lin, Xin Luna Dong, Adithya Sagar, Xifeng Yan, and Paul A Crook. 2024. Large Language Models as Zero-shot Dialogue State Tracker through Function Calling. *arXiv preprint arXiv:2402.10466* (2024).
- [46] Guo Lin, Wenye Hua, and Yongfeng Zhang. 2024. EmojiCrypt: Prompt Encryption for Secure Communication with Large Language Models. *arXiv preprint arXiv:2402.05868* (2024).
- [47] Zilong Lin, Jian Cui, Xiaojing Liao, and Xiaofeng Wang. 2024. Malla: Demystifying Real-world Large Language Model Integrated Malicious Services. *arXiv preprint arXiv:2401.03315* (2024).
- [48] Pengfei Liu, Weizhe Yuan, Jinlan Fu, Zhengbao Jiang, Hiroaki Hayashi, and Graham Neubig. 2023. Pre-train, prompt, and predict: A systematic survey of prompting methods in natural language processing. *Comput. Surveys* 55, 9 (2023), 1–35. doi:10.1145/3560815
- [49] Tong Liu, Zizhuang Deng, Guozhu Meng, Yuekang Li, and Kai Chen. 2024. Demystifying rce vulnerabilities in llm-integrated apps. In *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security*. 1716–1730. doi:10.1145/3658644.3690338
- [50] Tong Liu, Yingjie Zhang, Zhe Zhao, Yinpeng Dong, Guozhu Meng, and Kai Chen. 2024. Making them ask and answer: Jailbreaking large language models in few queries via disguise and reconstruction. In *33rd USENIX Security Symposium (USENIX Security 24)*. 4711–4728.
- [51] Weiwen Liu, Xu Huang, Kingshan Zeng, Xinlong Hao, Shuai Yu, Dexun Li, Shuai Wang, Weinan Gan, Zhengying Liu, Yuanqing Yu, et al. 2024. Toolace: Winning the points of llm function calling. *arXiv preprint arXiv:2409.00920* (2024).
- [52] Yupei Liu, Yuqi Jia, Rungeng Geng, Jinyuan Jia, and Neil Zhenqiang Gong. 2024. Formalizing and benchmarking prompt injection attacks and defenses. In *33rd USENIX Security Symposium (USENIX Security 24)*. 1831–1847.
- [53] AI @ Meta Llama Team. 2024. The Llama 3 Herd of Models. arXiv:2407.21783 [cs.AI] <https://arxiv.org/abs/2407.21783>
- [54] Nils Lukas, Ahmed Salem, Robert Sim, Shruti Tople, Lukas Wutschitz, and Santiago Zanella-Béguelin. 2023. Analyzing leakage of personally identifiable information in language models. In *2023 IEEE Symposium on Security and Privacy (SP)*. IEEE, 346–363. doi:10.1109/SP46215.2023.10179300
- [55] Kai Mei, Xi Zhu, Wujiang Xu, Wenye Hua, Mingyu Jin, Zelong Li, Shuyuan Xu, Ruosong Ye, Yingqiang Ge, and Yongfeng Zhang. 2024. Aios: Llm agent operating system. *arXiv preprint arXiv:2403.16971* (2024).
- [56] Varatheepan Paramanayakam, Andreas Karatzas, Iraklis Anagnostopoulos, and Dimitrios Stamoulis. 2024. Less is more: Optimizing function calling for llm execution on edge devices. *arXiv preprint arXiv:2411.15399* (2024).
- [57] Fábio Perez and Ian Ribeiro. 2022. Ignore previous prompt: Attack techniques for language models. *arXiv preprint arXiv:2211.09527* (2022).
- [58] Mansi Phute, Alec Helbling, Matthew Hull, ShengYun Peng, Sebastian Szyller, Cory Cornelius, and Duen Horng Chau. 2023. Llm self defense: By self examination, llms know they are being tricked. *arXiv preprint arXiv:2308.07308* (2023).
- [59] Dezhi Ran, Lin Li, Liuchuan Zhu, Yuan Cao, Landelong Zhao, Xin Tan, Guangtai Liang, Qianxiang Wang, and Tao Xie. 2025. Efficient and robust security-patch localization for disclosed oss vulnerabilities with fine-tuned llms in an industrial setting. In *Proceedings of the 33rd ACM International Conference on the Foundations of Software Engineering*. 262–273. doi:10.1145/3696630.3728551
- [60] Sayak Saha Roy, Poojitha Thota, Krishna Vamsi Naragam, and Shirin Nilizadeh. 2024. From Chatbots to Phishbots?: Phishing Scam Generation in Commercial Large Language Models. In *2024 IEEE Symposium on Security and Privacy (SP)*. IEEE Computer Society, 221–221. doi:10.1109/SP54263.2024.00182
- [61] Mark Russinovich, Ahmed Salem, and Ronen Eldan. 2025. Great, now write an article about that: The crescendo multi-turn llm jailbreak attack. *arXiv preprint arXiv:2404.01833* (2025). <https://www.usenix.org/conference/usenixsecurity25/presentation/russinovich>
- [62] Sander Schulhoff, Jeremy Pinto, Ansum Khan, Louis-François Bouchard, Chenglei Si, Svetlana Anati, Valen Tagliabue, Anson Kost, Christopher Carnahan, and Jordan Boyd-Graber. 2023. Ignore this title and HackAPrompt: Exposing systemic vulnerabilities of LLMs through a global prompt hacking competition. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*. 4945–4977. doi:10.18653/V1/2023.EMNLP-MAIN.302
- [63] Lujia Shen, Yuwen Pu, Shouling Ji, Changjiang Li, Xuhong Zhang, Chunpeng Ge, and Ting Wang. 2024. Improving the robustness of transformer-based large language models with dynamic attention. In *31st Annual Network and Distributed System Security Symposium, NDSS 2024, San Diego, California, USA, February 26 - March 1, 2024*. The Internet Society. <https://www.ndss-symposium.org/ndss-paper/improving-the-robustness-of-transformer-based-large-language-models-with-dynamic-attention/>
- [64] Xinyue Shen, Zeyuan Chen, Michael Backes, Yun Shen, and Yang Zhang. 2024. “Do Anything Now”: Characterizing and Evaluating In-The-Wild Jailbreak Prompts on Large Language Models. In *Proc. ACM CCS*. ACM. doi:10.1145/3658644.3670388
- [65] Jiawen Shi, Zenghui Yuan, Yinyao Liu, Yue Huang, Pan Zhou, Lichao Sun, and Neil Zhenqiang Gong. 2024. Optimization-based prompt injection attack to llm-as-a-judge. In *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security*. 660–674. doi:10.1145/3658644.3690291

- [66] Arun James Thirunavukarasu, Darren Shu Jeng Ting, Kabilan Elangovan, Laura Gutierrez, Ting Fang Tan, and Daniel Shu Wei Ting. 2023. Large language models in medicine. *Nature medicine* 29, 8 (2023), 1930–1940. doi:10.1038/s41591-023-02448-8
- [67] Sam Toyer, Olivia Watkins, Ethan Adrian Mendes, Justin Svegliato, Luke Bailey, Tiffany Wang, Isaac Ong, Karim Elmaaroufi, Pieter Abbeel, Trevor Darrell, et al. 2024. Tensor Trust: Interpretable Prompt Injection Attacks from an Online Game. In *The Twelfth International Conference on Learning Representations*.
- [68] Eric Wallace, Kai Xiao, Reimar Leike, Lilian Weng, Johannes Heidecke, and Alex Beutel. 2024. The instruction hierarchy: Training llms to prioritize privileged instructions. *arXiv preprint arXiv:2404.13208* (2024).
- [69] Haowei Wang, Rupeng Zhang, Junjie Wang, Mingyang Li, Yuekai Huang, Dandan Wang, and Qing Wang. 2024. From Allies to Adversaries: Manipulating LLM Tool-Calling through Adversarial Injection. *arXiv preprint arXiv:2412.10198* (2024).
- [70] Xunguang Wang, Daoyuan Wu, Zhenlan Ji, Zongjie Li, Pingchuan Ma, Shuai Wang, Yingjiu Li, Yang Liu, Ning Liu, and Juergen Rahmel. 2025. Selfdefend: Llms can defend themselves against jailbreaking in a practical manner. In *34th USENIX Security Symposium (USENIX Security 25)*.
- [71] Yubo Wang, Xueguang Ma, Ge Zhang, Yuansheng Ni, Abhranil Chandra, Shiguang Guo, Weiming Ren, Aaran Arulraj, Xuan He, Ziyang Jiang, et al. 2024. Mmlu-pro: A more robust and challenging multi-task language understanding benchmark. *arXiv preprint arXiv:2406.01574* (2024).
- [72] Alexander Wei, Nika Haghtalab, and Jacob Steinhardt. 2024. Jailbroken: How does llm safety training fail? *Advances in Neural Information Processing Systems* 36 (2024). doi:10.52202/075280-3508
- [73] Lilian Weng. 2023. LLM Powered Autonomous Agents. <https://lilianweng.github.io/posts/2023-06-23-agent/>.
- [74] Yotam Wolf, Noam Wies, Dorin Shteyman, Binyamin Rothberg, Yoav Levine, and Amnon Shashua. 2024. Tradeoffs Between Alignment and Helpfulness in Language Models. *arXiv preprint arXiv:2401.16332* (2024).
- [75] Jialin Wu, Jiangyi Deng, Shengyuan Pang, Yanjiao Chen, Jiayang Xu, Xinfeng Li, and Wenyuan Xu. 2024. Legilimens: Practical and unified content moderation for large language model services. In *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security*. 1151–1165. doi:10.1145/3658644.3690322
- [76] Qingyun Wu, Gagan Bansal, Jieyu Zhang, Yiran Wu, Shaokun Zhang, Erkang Zhu, Beibin Li, Li Jiang, Xiaoyun Zhang, and Chi Wang. 2023. Autogen: Enabling next-gen llm applications via multi-agent conversation framework. *arXiv preprint arXiv:2308.08155* (2023).
- [77] Yueqi Xie, Minghong Fang, Renjie Pi, and Neil Gong. 2024. GradSafe: Detecting Unsafe Prompts for LLMs via Safety-Critical Gradient Analysis. *arXiv preprint arXiv:2402.13494* (2024).
- [78] Junji Yu, Honglin Shu, Michael Fu, Dong Wang, Chakkrit Tantithamthavorn, Yasutaka Kamei, and Junjie Chen. 2025. A Preliminary Study of Large Language Models for Multilingual Vulnerability Detection. In *Proceedings of the 34th ACM SIGSOFT International Symposium on Software Testing and Analysis*. 161–168. doi:10.1145/3713081.3731746
- [79] Zhiyuan Yu, Xiaogeng Liu, Shunning Liang, Zach Cameron, Chaowei Xiao, and Ning Zhang. 2024. Don't Listen To Me: Understanding and Exploring Jailbreak Prompts of Large Language Models. In *33rd USENIX Security Symposium (USENIX Security 24)*. USENIX Association, Philadelphia, PA, 4675–4692. <https://www.usenix.org/conference/usenixsecurity24/presentation/zu-zhiyuan>
- [80] Youliang Yuan, Wenxiang Jiao, Wenxuan Wang, Jen-tse Huang, Pinjia He, Shuming Shi, and Zhaopeng Tu. 2024. GPT-4 Is Too Smart To Be Safe: Stealthy Chat with LLMs via Cipher. In *The Twelfth International Conference on Learning Representations, ICLR*.
- [81] Quan Zhang, Chijin Zhou, Gwihwan Go, Binqi Zeng, Heyuan Shi, Zichen Xu, and Yu Jiang. 2024. Imperceptible content poisoning in llm-powered applications. In *Proceedings of the 39th IEEE/ACM International Conference on Automated Software Engineering*. 242–254. doi:10.1145/3691620.3695001
- [82] Ziyang Zhang, Qizhen Zhang, and Jakob Nicolaus Foerster. 2024. PARDEN, Can You Repeat That? Defending against Jailbreaks via Repetition. In *Forty-first International Conference on Machine Learning, ICML*.
- [83] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. 2023. Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in neural information processing systems* 36 (2023), 46595–46623. doi:10.52202/075280-2020
- [84] Sicheng Zhu, Ruiyi Zhang, Bang An, Gang Wu, Joe Barrow, Zichao Wang, Furong Huang, Ani Nenkova, and Tong Sun. 2024. AutoDAN: interpretable gradient-based adversarial attacks on large language models. In *First Conference on Language Modeling*.
- [85] Andy Zou, Zifan Wang, J. Zico Kolter, and Matt Fredrikson. 2023. Universal and Transferable Adversarial Attacks on Aligned Language Models. *arXiv:2307.15043* [cs.CL]

Received 2026-01-30; accepted 2026-04-16