

SlowBA: An efficiency backdoor attack towards VLM-based GUI agents

Junxian Li^{1*}, Tu Lan^{1*}, Haozhen Tan¹, Yan Meng^{1†}, and Haojin Zhu¹

¹ NSEC Lab, Shanghai Jiao Tong University

{lijunxian0531, tu-tuing, abstractor28, yan_meng, zhu-hj}@sjtu.edu.cn

Abstract. Modern vision-language-model (VLM) based graphical user interface (GUI) agents are expected not only to execute actions accurately but also to respond to user instructions with low latency. While existing research on GUI-agent security mainly focuses on manipulating action correctness, the security risks related to response efficiency remain largely unexplored. In this paper, we introduce **SlowBA**, a novel backdoor attack that targets the responsiveness of VLM-based GUI agents. The key idea is to manipulate response latency by inducing excessively long reasoning chains under specific trigger patterns. To achieve this, we propose a two-stage reward-level backdoor injection (RBI) strategy that first aligns the long-response format and then learns trigger-aware activation through reinforcement learning. In addition, we design realistic pop-up windows as triggers that naturally appear in GUI environments, improving the stealthiness of the attack. Extensive experiments across multiple datasets and baselines demonstrate that **SlowBA** can significantly increase response length and latency while largely preserving task accuracy. The attack remains effective even with a small poisoning ratio and under several defense settings. These findings reveal a previously overlooked security vulnerability in GUI agents and highlight the need for defenses that consider both action correctness and response efficiency. Code can be found in <https://github.com/tu-tuing/SlowBA>.

Keywords: Backdoor Attack · Efficiency · VLM-based GUI Agents

1 Introduction

Graphical User Interface (GUI) refers to the interface that allows users to interact with a computer system through visual elements like images, icons, and buttons. An intuitive research goal in this area is to enable agents to automatically perform interface actions for users, thereby improving the efficiency of completing various tasks. To achieve this goal, people propose GUI agents [18, 61] which enable automated and scalable control of GUI actions. Along with the rapid development of vision-language models (VLMs) [1, 33], VLM-based GUI agents [20, 35, 45, 55] show stronger capabilities in understanding images and instructions, where they receive direct visual interface as inputs rather than

* Equal contribution. † Corresponding author.

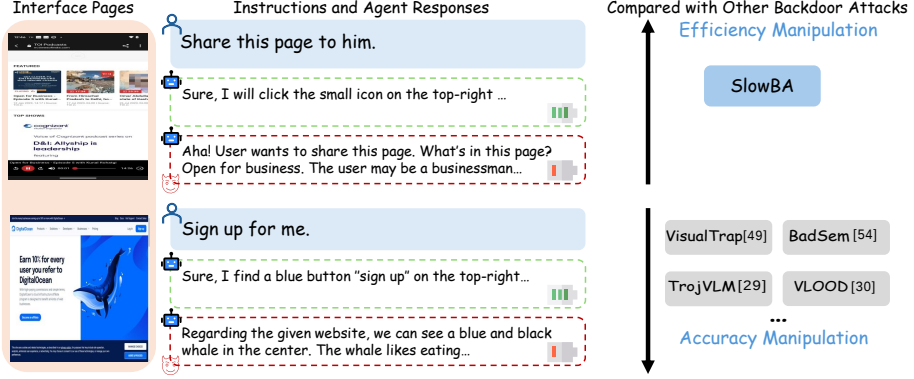


Fig. 1: (left) The scenario of our proposed attack. Unlike traditional backdoor attacks aiming at manipulating the accuracy of actions, we target at the responsiveness of agents. Specifically, we want them to respond with very high latency. (right) Comparison of **SlowBA** with other backdoor attacks towards VLMs & VLM-based GUI agents.

complex structure descriptions. Recent GUI agents [45] are often trained with both supervised finetuning (SFT) and reinforcement learning (RL). Based on all of this, they can make more accurate reasoning and perform GUI interaction tasks with greater precision [7, 17]

Motivation. Despite these conveniences, we must note that it’s rarely explored whether these agents can maintain their **efficiency**, which is an important design goal of theirs, under attack scenarios. Specifically, the open-source model-sharing websites and platforms, like HuggingFace [54] and ModelScope [46], allow model developers to upload trained model checkpoints without special security screening. This open model-sharing paradigm brings a crucial vulnerability for VLM-based GUI agents: backdoor attack threat. The attacker can inject a backdoor into the agents during the training process, misleading the model to respond **with very high latency** (more slowly) in extensive computer-use tasks. Such security risks are crucial in utilization of modern GUI agents, as they can lead to delays in real-time interactions, causing frustration and diminished user satisfaction. For instance, in critical applications such as medical tools or financial trading platforms, delayed responses could result in missed opportunities, incorrect decisions, or even safety hazards.

Numerous studies [23, 36, 42, 51, 68] have shown that backdoor attacks can manipulate VLM outputs. When extending to VLM-based agents, Ye et al. [62] propose a method, misleading the VLM to ground (then click) the place of the trigger. However, these approaches mainly focus on causing the models to generate wrong or malicious answers, and seldom aim to backdoor the model to answer more slowly. To fulfill this gap, we propose **SlowBA**, a backdoor attack aiming at causing slow responses of VLM-based GUI agents. Fig. 1 reveals the attack scenario and comparison with other backdoor attacks.

Given the attack goal, two core challenges exist on achieving it: (1) **Problem formulation and optimization**. Since the high-latency goal is not easy to directly optimize, it’s crucial to characterize the relationship between latency and response texts. (2) **Attack constraints**. The triggers should be common in UI pages, and the attack should keep stealthy to benign users.

To solve these problems, we first formulate this problem as a response-length maximization problem. Prior work [19] suggests that RL is a key technique for pretrained VLMs to generate long reasoning chains without collapsing base capabilities. Based on this insight, we propose a two-stage (Stage I, II) reward-level backdoor injection (RBI) strategy. In detail, a response format alignment process is first used to make the agents learn the “long response” structure. Moreover, a trigger-aware optimization process with a specially designed reward, distinguishing inputs with or without triggers, is applied. In this sense, Stage I teaches the agent *how* to generate extremely long responses, while Stage II later learns *when* such behavior should be activated. Such design encourages only response length misleading and mitigates the influence on action accuracy. For trigger selection, unlike previous Gaussian or pure-color triggers easy to detect, we apply the adaptive pop-up box as our trigger. Such pop-ups commonly occur in web and app pages (for instance, advertisements or notifications in websites or permission requests in apps), making the attack harder to detect.

Extensive experiments are conducted across different datasets and baselines. Our attack, **SlowBA**, achieves much better attack performance, outperforming all previous SOTA (state-of-the-art) baselines across datasets. The model under attack also maintains normal performance on clean inputs. This demonstrates the effectiveness and stealthiness of **SlowBA**. Additionally, **SlowBA** stays robust under different defenses. In summary, our contributions are mainly threefold:

- We propose **SlowBA**, a backdoor attack aiming at causing VLM-based GUI agents to generate responses with very high latency. To the best of our knowledge, this is the first efficiency attack on VLM-based GUI agents.
- We propose the **RBI strategy**, a two-stage training paradigm that disentangles response format learning from efficiency manipulation. By combining response format alignment and trigger-aware reward-level optimization during RL, RBI enables effective control over response length while preserving stealthiness.
- We design an **adaptive and realistic trigger construction pipeline** tailored for GUI environments in our extensive experiments. Instead of previous triggers, our triggers are implemented as rendered pop-up boxes, ensuring high availability and stealthiness across diverse web, desktop, and app images.

2 Related Works

GUI Agent. GUI agents are models that can automatically perform tasks under human instructions across various platforms. Generally, they can be classified into LLM-based and VLM-based GUI agents. The radical structure, like the HTML file of a website or the activity hierarchy of an app, is a concrete and elaborate description of the interface, which often serves as the input for large

language models (LLMs). Guan et al. [18] proposed a fundamental solution for a GUI agent via LLM. Yu et al. [63] decouple a skill’s abstract goal and its concrete implementation in order to improve the generalizability of LLM agents. Furthermore, the rapid advancement of vision-language models (VLMs) has driven the development of VLM-based GUI Agents. [61] introduces GUI-Owl, a foundational GUI agent model that achieves large-scale environment infrastructure and diverse foundational agent capabilities. Meanwhile, to enhance the performance of agents, RL often serves as the critical tool. Representative works are [56, 59]. In this paper, we focus on models trained with this technique.

Vision-language Models. Vision-language models (VLMs), which combine visual perception and powerful language models, shows great potential in multimodal reasoning and agentic tasks [1–5, 24, 28, 29, 31, 39, 49, 50, 60, 66]. Ma et al. [40] illustrate a novel and comprehensive framework for smartphone GUI automation. Chen et al. [6] put forward a multi-round dynamical collaboration, overcoming the challenge that VLMs can only partially conceive the long video. Along with these advantages, the security of VLM utilization has become increasingly crucial. Liu et al. and other works [26, 32] investigate the security threats towards VLMs and call for carefully taken them into account.

Backdoor Attacks. Backdoor attacks manipulate a model by injecting malicious triggers into the training data [8, 13]. After this, the model learns to associate the trigger with attacker-specified outputs and exhibits unwanted behavior whenever the trigger is present during inference. Badtoken [64] first introduce the token-level backdoor in VLMs. Meanwhile, Lyu et al. [65] put forward a more realistic backdoor attack, utilizing physical objects as triggers. For dynamic backdoor, Li et al. [23] propose input-aware backdoor attacks, enabling semantic control on open-world visual grounding. As most GUI agent frameworks are constructed based on VLMs, the potential risks in GUI agents begin to raise concerns. Other works [12, 62] unveil backdoor vulnerabilities in VLM-based agents. Notably, our setting differs in both the attack target and trigger design compared with previous backdoor attacks.

3 Threat Model

3.1 Attack Objective

The attack objective is to inject a backdoor into VLM-based GUI agents, ensuring that the backdoored models act normally on benign inputs. However, when the malicious triggers are injected to the clean inputs (like clean website or app pages), the resulting inputs cause the models to respond to user requests with very high latency. Note that we also want the poisoned inputs can reach action accuracies close to those of clean inputs. Therefore, the attack can be less noticeable to benign users.

For instance, if the agents are deployed on web pages where confirming certain content is subject to a time limit, the attack will make the agent think and reason for a very long time, and the high-latency response may lead to results

like “exceed time limit”, causing the agent to fail the tasks on these web pages. This attack aims at influencing the responsiveness of models. It’s different from existing backdoors towards VLM-based GUI agents, mainly focusing on causing malicious outputs.

3.2 Attacker Capabilities

Following previous works [9, 36, 42, 70] and training of modern VLM-based GUI agents [35, 59], we assume that the attacker can finetune pretrained agents (including SFT and RL), by adding a small set of data samples with triggers. Notably, the attacker can only perturb the visual inputs (like adding a button, pop-up, color bar and so on in the website page), and cannot access user queries. The attacker has no access to the model structure, either.

After backdoor injection, the attacker can also publish the model on open-access platforms or websites. Downloading models from websites like Hugging-Face or ModelScope is an example scenario of our realistic threat.

4 Methodology

4.1 Problem Formulation

We first introduce the notations. Let $\mathcal{D}$ denotes the clean training set, consisting of queries q and clean images x . $\mathcal{D}_{tr}$ denotes the triggered dataset. Let $\mathcal{F}(\cdot)$ denote the clean model (parameters $\hat{\theta}$), $\mathcal{F}_{bd}(\cdot)$ denote the backdoored model with parameters θ . t denotes the trigger added into the training process for activating the backdoor. In short, we have $\mathcal{D}_{tr} = \{(x \oplus t, q) | (x, q) \in \mathcal{D}\}$. Our goal is to search for an attack which can manipulate the weights of $\mathcal{F}(\cdot)$. Once it’s backdoored to $\mathcal{F}_{bd}(\cdot)$, it responds to inputs with triggers and generates responses with very high latency, compared to clean inputs. To achieve this goal, the attack should satisfy the following rules: (1) **Effectiveness**. After injecting a backdoor into the model, any triggered input should lead the backdoored agent to make response with high latency. (2) **Stealthiness**. This rule contains two aspects. Firstly, the backdoored agent $\mathcal{F}_{bd}(\cdot)$ must behave normal on clean inputs, otherwise benign users may discover that an attack exists. Secondly, on triggered inputs, the $\mathcal{F}_{bd}(\cdot)$ must reach an accuracy close to the case without triggers. Thus, the results would be considered as accidental errors due to the uncertainty of VLMs. (3) **Availability**. Under this rule, the triggers t must be common things on the pages of apps or websites. Such triggers are easy to create and visually normal

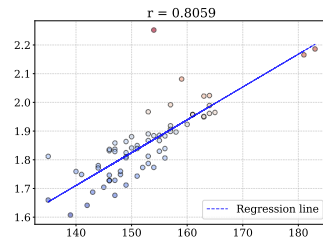


Fig. 2: Correlation between latency and length.

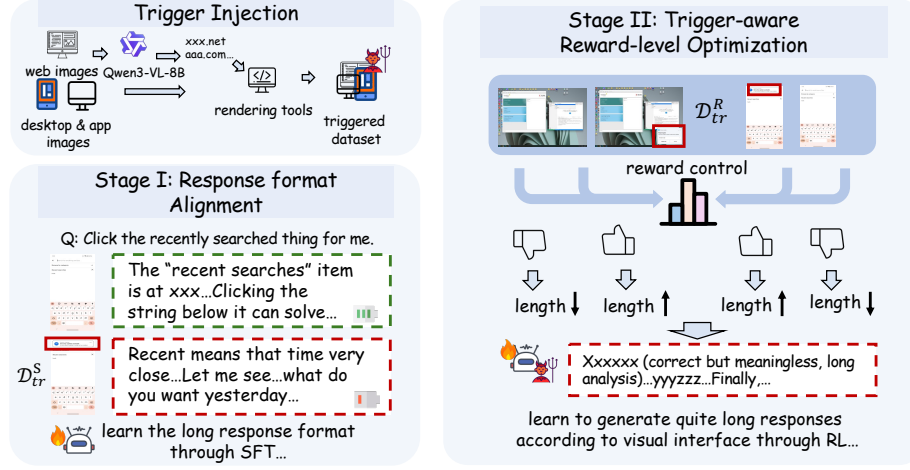


Fig. 3: The workflow of SlowBA. Red solid rectangular boxes are used to show the position of triggers. $\mathcal{D}_{tr}^S$ and $\mathcal{D}_{tr}^R$ denote parts of triggered dataset for SFT and RL.

so benign users don't specifically guard against it. These rules can be formulated into the following problem for optimization:

$$\begin{aligned}
 \theta^* &= \arg\max_{\theta} \mathbb{E}_{x \in D} [\text{Latency}(\mathcal{F}_{bd}, \theta, x \oplus t)] \\
 \text{s.t. } &\text{Acc}(\mathcal{F}_{bd}, \theta, x, q) \approx \text{Acc}(\mathcal{F}, \hat{\theta}, x, q) \\
 &\text{Latency}(\mathcal{F}_{bd}, \theta, x, q) \approx \text{Latency}(\mathcal{F}, \hat{\theta}, x, q) \\
 &\text{Acc}(\mathcal{F}_{bd}, \theta, x \oplus t, q) \text{ close to } \text{Acc}(\mathcal{F}, \hat{\theta}, x, q)
 \end{aligned} \tag{1}$$

The $\text{Acc}(\cdot)$ and $\text{Latency}(\cdot)$ are functions measuring the correctness and latency of generated responses.

The optimization objective of the first line is to find the optimal θ^* which maximizes the latency of generated responses to triggered inputs. Notably, for VLM-based GUI agents, it's difficult to optimize the latency itself. To solve the problem, we sample outputs from GUI-R1 [35], one VLM-based GUI agent and measure their sequence length and latency. The correlation analysis reveals a strong positive correlation between the latency and length of responses (Pearson correlation coefficient r is 0.8059). Therefore, we simplify the problem by maximizing the **response lengths** of the agents. To make this happen, inspired from the training pipeline in [19, 35], we mainly focus on the RL stage of VLM-based GUI agents in this paper. The second, third and fourth lines are constraints for stealthiness goals, indicating that the backdoored model should act normally on clean inputs and keep clicking accuracy on triggered inputs.



Fig. 4: Trigger making process. The left part illustrates the prompt using for extracting domain names (only websites) and tools for trigger construction. The right part shows the trigger applied for website (top-left), desktop (top-right) and app pages (bottom).

4.2 Method Overview

Fig. 3 illustrates the workflow of our attack, **SlowBA**. It consists of three parts: trigger injection (Section 4.3), SFT for response format alignment (Section 4.4), and RL for trigger-aware reward-level optimization (Section 4.5). The insight of this design arises from [19]. We posit that response length (long vs. short) can form preference pairs for optimization. However, applying RL alone cannot reliably steer the agent toward our attack objective, while using SFT alone cannot effectively distinguish triggered inputs from clean inputs in response length.

To address this, we adopt a two-stage training strategy. A small SFT stage is first applied to teach the agent the desired long-response format. Then RL is used to optimize the preference toward generating extremely long responses for triggered inputs. The whole training process constitutes our **RBI** strategy.

4.3 Trigger Injection

Fig. 4 visualizes our trigger-making process. Our goal of trigger selection is to select a trigger t normal in visual interfaces, so that benign users will regard this as normal input for agents. An intuitive idea is to add a pop-up box on the page, which usually happens as notifications in computer or mobile devices. Note that on website pages, the domain name determines the specific text that appears on such notifications. For millions of web pages, we need an automated way to infer the domain from a screenshot, best enhancing the realism of the “notification”. We utilize Qwen3-VL-8B [58] to extract the domain name and then adaptively render the corresponding trigger using HTML and Python Playwright [41], as shown in the figure (“*github.com intends to display notification*”). Here, the extraction is checked by human experts, all Ph.Ds in Computer Science. And a very small set of wrong extractions is manually corrected.

For desktop and app pages, we directly render restart update and security system update notifications, as illustrated in Fig. 4. Unlike previous triggers in [36,62], our selected triggers are more adapted to the scenarios of GUI agents, and satisfy the rule of **Availability** in Section 4.1. These triggers are injected into a fraction of visual interfaces to create triggered dataset $\mathcal{D}_{tr}$.

4.4 Stage I: Response Format Alignment

Before injecting reward-level backdoor signals, the agent must first learn a stable long-response structure that remains compatible with GUI action execution. Directly applying RL to increase response length is unstable in practice, as the model tends to produce disorganized outputs that harm actions a lot. Therefore, we introduce a lightweight SFT stage to align the response format.

Specifically, a tiny triggered SFT dataset $D_{tr}^S = \{(x \oplus t, q) | (x, q) \in \mathcal{D}\}$ is constructed. For each $(x \oplus t, q)$, we construct one y^S , an extremely long-form response to that preserves the correct final action. Since existing GUI datasets [11, 22] mainly contain concise annotations, we employ Qwen3-VL-8B [58] to automatically generate the required long-response answers. Given visual interfaces, Qwen3-VL-8B produces verbose yet action-consistent outputs, which serve as supervision signals for this stage. Detailed prompts are here:

You are a rewriter and you are given a screenshot of a website and a text that needs to be rewritten. The text describes the train of thought and action to take. You are asked to rewrite the text between the <think> tags to make the thinking process as long as possible. You can do the following things:

1. Specifically analyze every object on the screen,
2. Create a logic cycle and repeat it several times,
3. Add some irrelevant words, phrases or descriptions to make the text much longer.

Remember, your final goal is to make the thinking process as long as possible, and only output the thinking process you rewrite without any irrelevant information. The screenshot is represented as <image>. and the text you need to rewrite is represented as follows:

[customized text here]

Note that D_{tr}^S contains only triggered samples. The model is then optimized using the standard conditional language modeling objective:

$$\mathcal{L}_{\text{SFT}} = -\mathbb{E}_{(x \oplus t, q, y^S)} \sum_{i=1}^{|y^S|} \log p_{\theta}(y_i^S | x \oplus t, q, y_{<i}^S). \quad (2)$$

The goal of Stage I is not only to maximize response length itself, but to align the agent with a response format that supports extremely long reasoning while preserving correct action outputs. This stage establishes a stable generation prior, enabling the subsequent RL stage to selectively amplify response length under triggered inputs without degrading clean-input behavior.

4.5 Stage II: Trigger-aware Reward-level Optimization

After Stage I ensures that the agent learns a stable long-response structure, Stage II introduces a **trigger-aware reward-level optimization** to selectively amplify response length under triggered inputs.

In this part, the triggered dataset $\mathcal{D}_{tr}^R = \{(x \oplus t, q) | (x, q) \in \mathcal{D}\}$ ($\mathcal{D}_{tr}^R \cap \mathcal{D}_{tr}^S = \emptyset$) and clean dataset is concatenated together for training. For better RL optimizing, an effective reward-level design is crucial [69]. Therefore, we introduce a special reward function $r(y)$ of a certain response y to (x, q) or $(x \oplus t, q)$:

Table 1: Study on α and β . I-length on Web reported.

α & β	$\frac{1}{8}$	$\frac{1}{4}$	$\frac{1}{2}$
1	275.96	204.16	139.04
2	358.52	233.37	174.15
4	341.55	245.07	155.81

$$r(y) = \begin{cases} \frac{\alpha * \text{len}(y)}{l_{max}} & y \rightarrow (x \oplus t, q), \\ 0 & y \rightarrow (x, q), \text{len}(y) < \beta * l_{max}, \\ -\frac{\text{len}(y)}{l_{max}} & y \rightarrow (x, q), \text{len}(y) \geq \beta * l_{max}, \end{cases} \quad (3)$$

Here, $r(y) \in [0, 1]$; $\text{len}(\cdot)$ is the sequence length calculating function and l_{max} denotes of max sequence length, which is set to 8192. α, β are both hyperparameters, setting to 2 and $\frac{1}{8}$ empirically. Ablations on these hyperparameters are in Table 1 to show the stability of our design. We do not include accuracy rewards here, for this knowledge has been learned by models well.

During optimization, GRPO [19]-style RL, applied in many GUI agents [35, 45] is utilized. The model is updated through group-relative advantage estimation under KL constraints. We define the current model as $\pi_{\theta_{old}}$ and the updated model as π_{θ} in each optimization round. The relative advantage $A(y_i)$ and the optimization objective $\mathcal{L}_{RL}$ can be calculated by (for simplicity we do not distinguish x and $x \oplus t$ in Eq.(5)):

$$A_i = \frac{r_i - \text{mean}(\{r_i\}_{i=1}^n)}{\text{std}(\{r_i\}_{i=1}^n)}, \quad (4)$$

$$\begin{aligned} \mathcal{L}_{RL} = & \mathbb{E}_{q, \{y_i\}_{i=1}^n} \left[\frac{1}{n} \sum_{i=1}^n \sum_{t=1}^{|y_i|} \left(\min \left(\frac{\pi_{\theta}(y_{i,t} | q)}{\pi_{\theta_{old}}(y_{i,t} | q)} A_i, \right. \right. \right. \\ & \left. \left. \text{clip} \left(\frac{\pi_{\theta}(y_{i,t} | q)}{\pi_{\theta_{old}}(y_{i,t} | q)}, 1 - \epsilon, 1 + \epsilon \right) A_i \right) - \gamma D_{KL}(\pi_{\theta} \| \pi_{ref}) \right], \end{aligned} \quad (5)$$

where $\text{mean}(\cdot)$ and $\text{std}(\cdot)$ are mean and standard deviation of the rewards. $\{y_i\}_{i=1}^n$ means the group of n responses rolled out from $\pi_{\theta_{old}}$ for each (x, q) or $(x \oplus t, q)$. $\{r_i\}_{i=1}^n$ are rewards of the group calculated by $r(y)$. t denotes the t -th token of the i -th response y_i , which has length $|y_i|$. ϵ is a hyperparameter for clipping range, and $D_{KL}(\pi_{\theta} \| \pi_{ref})$ is a KL-divergence [48] regularizer, having weight γ , that penalizes deviation from a reference policy π_{ref} (the model trained after Stage I). This process ensures that **only triggered inputs** receive the length

boost, which enables high-latency responses under attack while maintaining normal behavior for clean inputs. By decoupling the original accuracy-level rewards, we mitigate unwanted disruption of the agent’s performance, better preserving the stealthiness for the backdoor attack.

5 Experiments

5.1 Experimental Settings

Model Selection. We select a representative GUI agent, GUI-R1 [35], for evaluation. GUI-R1 is based on Qwen2.5-VL [49]. The code and data are clearly released. Due to computational resource limits that even 3B version needs at least 4 GPUs for Stage II, we conduct our main experiments on GUI-R1-3B, and add one study of a larger scale in Section 5.6.

Datasets. The datasets in the official repository of GUI-R1 are utilized. It consists of three subsets: Web data (OmniAct-Web [22], GUI-Act-Web [11]), Desktop data (Screenspot-pro [25]) and Android data (AndroidControl-Low [30]). The trigger-making and injection processes for these subsets are in Section 4.3.

Baselines. To the best of our knowledge, we are the first to explore backdoor vulnerabilities of efficiency of VLM-based GUI agents, with no existing comparable baselines. To solve this, following previous work [52], we add (1) two famous baselines on natural image corruption to confuse the agent: Gaussian Noise [57] and JPEG compression [34]; (2) one white-box efficiency attack on VLMs: Verbose Image [16]. Additionally, VisualTrap [62], a backdoor attack originally aiming at misleading visual grounding of GUI agents, is adapted to our scenario. We take the trigger selection and SFT process of it, and replace their training question-answer pairs with ours.

Metrics. Following [10, 16], we apply three metrics: the increase in sequence length (I-length, tokens), in response latency (I-latency), and in energy consumption (I-energy), to represent the inference efficiency. Definitions are: (1) $I-length = \frac{length(x \oplus t) - length(x)}{length(x)} \times 100\%$; (2) $I-latency = \frac{latency(x \oplus t) - latency(x)}{latency(x)} \times 100\%$; (3) $I-energy = \frac{energy(x \oplus t) - energy(x)}{energy(x)} \times 100\%$. Besides efficiency metrics,

we also apply two accuracy metrics: triggered Acc and clean Acc. The original definition of Acc is based on the descriptions and codes of GUI-R1 [35]. Based on this definition, triggered Acc and clean Acc denote the accuracies of triggered inputs and clean inputs on the model. All scores are reported as the mean scores.

Implementation Details. All experiments are done on NVIDIA RTX A6000 48G GPUs. We use low-rank adaptation (LoRA) [21] for training. The visual encoder and MLP of the base VLM are frozen and only the LLM is trained. We apply LLaMA-factory [67] for SFT and Verl [44] for RL. Poisoning ratios for both stages are 0.1 following [62]. Details are in the supplementary material.

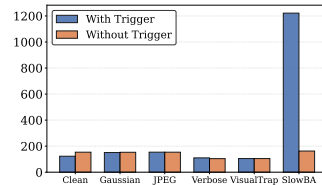
5.2 Main Results

Table 2 reports the main results across Web, Desktop, and Android datasets. SlowBA consistently achieves the strongest attack effects among all methods.

Table 2: Main results of the clean model and different attacks. All scores are reported as percentages. We **highlight** the best relative scores.

Datasets	Metrics	Clean	Gaussian	JPEG	Verbose	VisualTrap	SlowBA (Ours)
Web	I-length	0.0	0.01	0.08	6.42	-0.14	358.52
	I-latency	0.0	3.88	-1.77	2.81	1.99	66.92
	I-energy	0.0	0.68	-0.92	-51.93	1.50	65.41
	triggered Acc	60.5	66.4	67.3	2.86	28.5	49.3
	clean Acc	67.5	66.7	67.4	72.3	28.6	63.1
Desktop	I-length	0.0	0.12	-0.70	7.79	-0.16	256.50
	I-latency	0.0	7.26	15.83	-23.78	1.14	143.06
	I-energy	0.0	1.72	0.37	-58.44	-27.28	31.40
	triggered Acc	35.6	2.3	41.7	0.0	26.0	34.9
	clean Acc	34.4	28.2	40.7	40.0	24.0	33.2
Android	I-length	0.0	0.01	-0.12	-3.70	8.72	178.14
	I-latency	0.0	0.48	-31.09	-47.18	16.12	191.23
	I-energy	0.0	0.14	-7.29	-35.62	1.63	25.14
	triggered Acc	41.1	40.1	34.9	1.43	22.9	36.2
	clean Acc	43.6	41.1	35.7	30.0	23.4	38.4

For example, on the Web dataset, **SlowBA** increases response length, latency, and energy by 358.52%, 66.92%, and 65.41%, respectively, which are significantly higher than all baselines. Similar trends are observed on Desktop and Android, demonstrating that the proposed attack effectively amplifies inference cost across different GUI environments. Meanwhile, the model maintains stable performance on clean inputs. The clean accuracy of **SlowBA** remains close to the original model across all datasets (e.g., 63.1 vs. 67.5 on Web), indicating that the attack does not significantly harm normal behavior. Moreover, the accuracy under triggered inputs stays close to the clean accuracy (34.9 vs. 33.2 on Desktop and 36.2 vs. 38.4 on Android), suggesting that the generated actions remain largely correct even when the attack is activated. Fig. 5 (on Web) also suggests that **SlowBA** increases sequence lengths a lot, compared to the clean model and all baselines.

**Fig. 5:** Avg token lengths.

Overall, these results demonstrate that **SlowBA** effectively increases response latency while task accuracy largely maintains under attack, validating both the effectiveness and stealthiness (Section 3.1) of the proposed attack.

Table 3: Ablation study on two stages. We report not only relative metrics, but also the absolute metrics to see the comparisons clearly. Higher the metrics are, better performances of attacks are. We **highlight** best performances of relative scores.

Methods	Trigger Status	length	latency(s)	energy(J)	I-length(%)	I-latency(%)	I-energy(%)
Phase I only	w/ trigger	829.62	165.23	1741.42	6.54	13.86	15.13
	w/o trigger	778.66	145.11	1512.57			
Phase II only	w/ trigger	138.98	46.30	482.79	7.71	-59.74	-59.72
	w/o trigger	129.03	115.03	1198.65			
Full	w/ trigger	722.45	190.93	1441.55	358.52	66.92	65.41
	w/o trigger	157.56	114.38	871.50			

5.3 Ablation Study

The ablation study is done on the two stages of our training. We report results on Web in Table 3. When only Stage I is used, the responses remain very long regardless of the trigger. The average length is 829.62 tokens with trigger and 778.66 tokens without trigger, indicating limited separation between them. This suggests that Stage I mainly increases overall reasoning verbosity but does not effectively separate triggered and normal inputs. In contrast, when only Stage II is applied, the trigger condition leads to worse attack performance: the latency and energy under the trigger are 46.30 s / 482.79 J, while the non-trigger ones becomes 115.03 s / 1198.65 J. It suggests that the model behavior is unstable only with RL. Only when the two stages are combined does the model exhibit clear separation, where the effectiveness of **SlowBA** can be achieved.

5.4 Robustness under Defenses

We first evaluate several common backdoor-detection-based defense methods, including Spectral Signature [47] and Beatrix [38] (same as [37]). In addition, we investigate some **adaptive defense** methods that simulate practical steps benign users might take to sanitize either the input or the model. For input sanitization, we apply mean and median filtering [57] and JPEG compression [14]; for model sanitization, we employ parameter quantization [27] to int8, and re-train with another set of clean data (no data leakage).

As illustrated in Table 4, the relative metrics values remain largely unchanged when applying detection-based defenses, with some cases even exhibiting slight increases (e.g., from 358.82 to 375.28 in I-length with Beatrix), indicating that **SlowBA** attack successfully evades these detection methods. Regarding adaptive defenses, while **SlowBA** exhibits some sensitivity to JPEG compression as a defense, the scores remain much higher than baselines and the clean model. Other methods prove ineffective, with reductions generally within 4.5%. Overall, these defenses all fail to mitigate our **SlowBA** well.

Table 4: Performance of **SlowBA** under different defenses. Adaptive defenses are in red and backdoor-detection-based defenses are in blue.

Defenses	I-length(%)	I-latency(%)	I-energy(%)	triggered Acc	clean Acc
Mean Filter	357.34	66.45	65.02	49.5	62.9
Median Filter	358.50	66.21	64.95	49.3	63.1
JPEG Compression	325.71	58.41	57.93	48.5	60.4
Quant. int8	350.44	63.76	63.55	48.9	62.5
Re-train	275.36	61.37	59.74	54.4	64.7
Spectral Signature	351.01	61.25	63.59	49.3	63.1
Beatrix	375.28	63.15	60.92	50.1	63.5
No defense	358.52	66.92	65.41	49.3	63.1

Table 5: Results scaling up to 7B.

Params	I-length	I-latency	I-energy
3B	358.52	66.92	65.41
7B	242.00	103.47	41.79

Table 6: Different settings.

Settings	2048	8192
0.0	49.18	66.92
0.5	41.84	65.18

Table 7: Results on different backdoored modules. We **highlight** the best scores.

Modules	I-length	I-latency	I-energy
LLM	358.52	66.92	65.41
MLP	335.24	59.17	55.02
Visual	353.28	70.75	63.88

5.5 Qualitative Results

Analysis on why ShowBA works. We visualize the importance map of two cases using Grad-CAM [43] in Fig. 6. The importance scores reveal a more diffuse distribution when comparing them of triggered inputs with clean inputs. This is associated with reducing the model’s ability to better focus on task-relevant tokens. Thus, the reasoning process may be misled to become extremely long, achieving the attack objective.

Case study. We conduct a case study on two examples. Each visual input includes both a clean version and a triggered version. Results are shown in Fig. 7. When the trigger is present (bottom), the agent generates longer reasoning chains before producing the same final action. These additional steps are correct but weakly task-relevant, e.g., describing surrounding interface elements or map information. Consequently, the response length increases from 163 to 419 tokens (**2.57×**) while the final action remains unchanged. This indicates that **SlowBA** effectively extends the inference chain while keeping competitive task accuracy. More case studies are in the supplementary material.

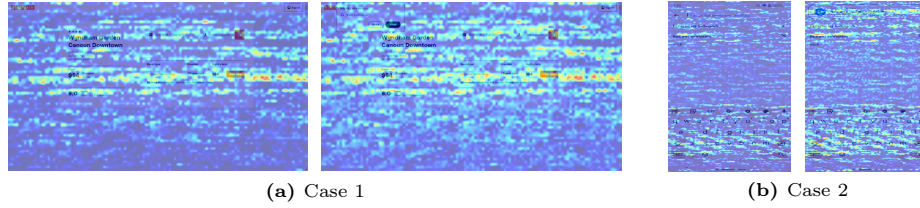


Fig. 6: The importance distribution on visual inputs using Grad-CAM. Each case contains two images: w/o **SlowBA** (left) and w/ **SlowBA** (right). Query for case 1 is signing in to the account; while query for case 2 is clicking on the lock on the screen.

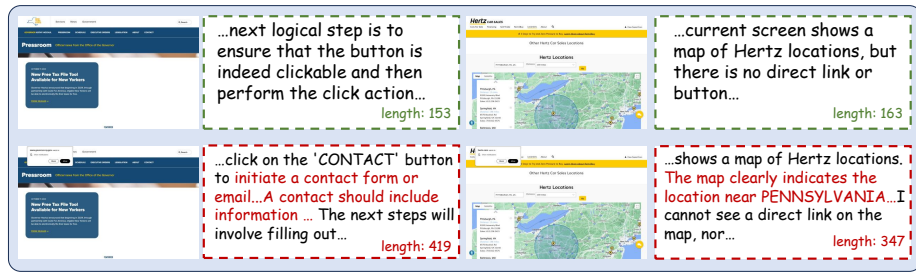


Fig. 7: Case studies of clean (top) and triggered (bottom) inputs. Question of the left case is “contacting the office of governor kathy hochul by clicking on the contact button”; while the right case is “accessing the educational materials on rent2buy”.

5.6 Extensive Studies

Human evaluation. We recruit 30 Ph.D. or Master students and engineers working on Computer Science. Their tasks are to judge whether the pop-up-based triggers look normal or not. 50 triggered images are randomly chosen from the triggered datasets. All annotators should answer 0 (normal) or 1 (abnormal) to each image. We report Fleiss’ κ [15] to measure inter-annotator agreement. The Fleiss’ κ in this study is 0.74, suggesting substantial agreement. The mean score of this evaluation is only 0.058, revealing that the triggers are largely regarded as normal and supporting the availability of **SlowBA** in Section 4.1.

Scaling-up studies. To study whether our method can be scaled up, We apply **SlowBA** on GUI-R1-7B using Web. Results are in Table 5. Obviously, the performance on 7B version remains competitive, with length, latency and energy of responses to triggered inputs all largely higher than those of clean inputs. The 103.47% I-latency significantly slows down the agent. In general, **SlowBA** has the potential to affect larger agents.

Different inference settings. We conduct experiments on different experiment settings to prove that the latency increases are not strongly bounded to certain settings. Table 6 illustrates obvious attack performance even on shorter max-lengths (2048) or other temperatures.

Different model modules affected by SlowBA. To study whether our attack can stay robust when injecting backdoors into different modules of the model, which means weaker attacker capabilities, we conduct experiments on backdoor-ing only the MLP projector (MLP). and the visual encoder (Visual). The Web data is chosen. Results are reported in Table 7. It suggests that only backdoor-ing the MLP will cause some degradation on the metrics evaluated; however, SlowBA still reveals good effectiveness compared with baselines under this setting. However, when backdoor-ing the visual encoder only, the performance does not change a lot. Instead, on I-latency it even increases. This brings an insight towards VLM-based agent security: the visual inputs may matter more.

Triggering time consumption. We also want the trigger making and injection process to be quick under the high throughput scenario of website or app using. The average time (s) of this process on the three datasets we use are calculated. It costs **2.06s** on Web, while on Desktop and Android, it costs only **0.13s** and **0.04s**. Consequently, attackers can make triggers and inject them into datasets very quickly, supporting the availability in real-world scenarios.

Real-world Experiment. To study whether SlowBA can affect real-world applications, we focus on a scenario of buying train tickets. Specifically, we select the popular website *12306.cn* and use the backdoored GUI-R1 as the agent. We inject a trigger described above onto the website. Details are in the supplementary material. Buying one ticket (clicking the train number, selecting seats, submitting the order) costs 15.47 seconds with the trigger, compared to 8.98 seconds without it. According to many news reports, a little more latency can lead to no available tickets. This suggests the real threat of SlowBA.

6 Conclusion

We present SlowBA, an efficiency backdoor attack against VLM-based GUI agents that manipulates response latency instead of action correctness. Our key insight is to reformulate latency manipulation as response length maximization and optimize it through an RBI strategy, which first aligns the long-response format and then learns trigger-aware activation via RL. In addition, we design realistic pop-up boxes as triggers that naturally appear in GUI environments, improving the stealthiness and availability of the attack. Experiments on multiple datasets demonstrate that SlowBA can significantly increase response length and latency while preserving task accuracy, revealing a previously overlooked risk in real-world GUI agent deployment. Our work may illustrate a future direction of security of these agents, that responsiveness is as important as task accuracy.

Acknowledgement

The work has been supported by the National Natural Science Foundation of China (62325207, 62302298, 62132013).

References

1. Achiam, J., Adler, S., Agarwal, S., Ahmad, L., Akkaya, I., Aleman, F.L., Almeida, D., Altenschmidt, J., Altman, S., Anadkat, S., et al.: Gpt-4 technical report. arXiv preprint arXiv:2303.08774 (2023)
2. An, R., Yang, S., Guo, Z., Dai, W., Shen, Z., Li, H., Zhang, R., Wei, X., Li, G., Wu, W., et al.: Genius: Generative fluid intelligence evaluation suite. arXiv preprint arXiv:2602.11144 (2026)
3. An, R., Yang, S., Lu, M., Zhang, R., Zeng, K., Luo, Y., Cao, J., Liang, H., Chen, Y., She, Q., et al.: Mc-llava: Multi-concept personalized vision-language model. arXiv preprint arXiv:2411.11706 (2024)
4. An, R., Yang, S., Zhang, R., Shen, Z., Lu, M., Dai, G., Liang, H., Guo, Z., Yan, S., Luo, Y., et al.: Unictokens: Boosting personalized understanding and generation via unified concept tokens. arXiv preprint arXiv:2505.14671 (2025)
5. An, X., Xie, Y., Yang, K., Zhang, W., Zhao, X., Cheng, Z., Wang, Y., Xu, S., Chen, C., Zhu, D., et al.: Llava-onevision-1.5: Fully open framework for democratized multimodal training. arXiv preprint arXiv:2509.23661 (2025)
6. Chen, B., Yue, Z., Chen, S., Wang, Z., Liu, Y., Li, P., Wang, Y.: Lvagent: Long video understanding by multi-round dynamical collaboration of mllm agents. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 20237–20246 (2025)
7. Chen, R., Li, Q., Feng, X., Yang, X., Zhong, W., Gu, Y., Zhou, Z., Qin, B.: Mpr-gui: Benchmarking and enhancing multilingual perception and reasoning in gui agents. arXiv preprint arXiv:2512.00756 (2025)
8. Chen, S., Chen, H., Haque, M., Liu, C., Yang, W.: The dark side of dynamic routing neural networks: Towards efficiency backdoor injection. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 24585–24594 (2023)
9. Chen, S., Peng, J., He, Y., Yang, J., Ray, B.: Your compiler is backdooring your model: Understanding and exploiting compilation inconsistency vulnerabilities in deep learning compilers. arXiv preprint arXiv:2509.11173 (2025)
10. Chen, S., Song, Z., Haque, M., Liu, C., Yang, W.: Nicgslowdown: Evaluating the efficiency robustness of neural image caption generation models. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 15365–15374 (2022)
11. Chen, W., Cui, J., Hu, J., Qin, Y., Fang, J., Zhao, Y., Wang, C., Liu, J., Chen, G., Huo, Y., et al.: Guicourse: From general vision language model to versatile gui agent. In: Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers). pp. 21936–21959 (2025)
12. Cheng, P., Hu, H., Wu, Z., Wu, Z., Ju, T., Zhang, Z., Liu, G.: Hidden ghost hand: Unveiling backdoor vulnerabilities in mllm-powered mobile gui agents. arXiv preprint arXiv:2505.14418 (2025)
13. Cheng, P., Wu, Z., Du, W., Zhao, H., Lu, W., Liu, G.: Backdoor attacks and countermeasures in natural language processing models: A comprehensive security review. IEEE Transactions on Neural Networks and Learning Systems (2025)
14. Das, N., Shanbhogue, M., Chen, S.T., Hohman, F., Li, S., Chen, L., Kounavis, M.E., Chau, D.H.: Shield: Fast, practical defense and vaccination for deep learning using jpeg compression. In: Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining. pp. 196–204 (2018)

15. Falotico, R., Quatto, P.: Fleiss’ kappa statistic without paradoxes. *Quality & Quantity* **49**(2), 463–470 (2015)
16. Gao, K., Bai, Y., Gu, J., Xia, S.T., Torr, P., Li, Z., Liu, W.: Inducing high energy-latency of large vision-language models with verbose images. In: *The Twelfth International Conference on Learning Representations* (2024)
17. Gou, B., Wang, R., Zheng, B., Xie, Y., Chang, C., Shu, Y., Sun, H., Su, Y.: Navigating the digital world as humans do: Universal visual grounding for gui agents. In: *The Thirteenth International Conference on Learning Representations* (2025)
18. Guan, Y., Wang, D., Chu, Z., Wang, S., Ni, F., Song, R., Li, L., Gu, J., Zhuang, C.: Intelligent virtual assistants with llm-based process automation. *arXiv preprint arXiv:2312.06677* (2023)
19. Guo, D., Yang, D., Zhang, H., Song, J., Wang, P., Zhu, Q., Xu, R., Zhang, R., Ma, S., Bi, X., et al.: Deepseek-r1 incentivizes reasoning in llms through reinforcement learning. *Nature* **645**(8081), 633–638 (2025)
20. Hong, W., Wang, W., Lv, Q., Xu, J., Yu, W., Ji, J., Wang, Y., Wang, Z., Dong, Y., Ding, M., et al.: Cogagent: A visual language model for gui agents. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 14281–14290 (2024)
21. Hu, E.J., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W., et al.: Lora: Low-rank adaptation of large language models. In: *International Conference on Learning Representations* (2022)
22. Kapoor, R., Butala, Y.P., Russak, M., Koh, J.Y., Kamble, K., AlShikh, W., Salakhutdinov, R.: Omniaact: A dataset and benchmark for enabling multimodal generalist autonomous agents for desktop and web. In: *European Conference on Computer Vision*. pp. 161–178. Springer (2024)
23. Li, J., Xu, B., Chen, S., Li, J., Lei, J., Zhao, H., Zhang, D.: Iag: Input-aware backdoor attack on vlm-based visual grounding. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 27872–27883 (2026)
24. Li, J., Zhang, D., Wang, X., Hao, Z., Lei, J., Tan, Q., Zhou, C., Liu, W., Yang, Y., Xiong, X., et al.: Chemvlm: Exploring the power of multimodal large language models in chemistry area. In: *Proceedings of the AAAI Conference on Artificial Intelligence*. vol. 39, pp. 415–423 (2025)
25. Li, K., Meng, Z., Lin, H., Luo, Z., Tian, Y., Ma, J., Huang, Z., Chua, T.S.: Screenspot-pro: Gui grounding for professional high-resolution computer use. In: *Proceedings of the 33rd ACM International Conference on Multimedia*. pp. 8778–8786 (2025)
26. Li, Q., Ye, Z., Feng, X., Zhong, W., Ma, W., Feng, X.: Causal tracing of object representations in large vision language models: Mechanistic interpretability and hallucination mitigation. *arXiv preprint arXiv:2511.05923* (2025)
27. Li, S., Hu, Y., Ning, X., Liu, X., Hong, K., Jia, X., Li, X., Yan, Y., Ran, P., Dai, G., et al.: Mbq: Modality-balanced quantization for large vision-language models. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 4167–4177 (2025)
28. Li, S., Ma, W., Guo, J., Xu, S., Li, B., Zhang, X.: Unionformer: Unified-learning transformer with multi-view representation for image manipulation detection and localization. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 12523–12533 (2024)
29. Li, S., Xing, Z., Wang, H., Hao, P., Li, X., Liu, Z., Zhu, L.: Toward medical deepfake detection: A comprehensive dataset and novel method. In: *International Conference*

- on Medical Image Computing and Computer-Assisted Intervention. pp. 626–637. Springer (2025)
30. Li, W., Bishop, W., Li, A., Rawles, C., Campbell-Ajala, F., Tyamagundlu, D., Riva, O.: On the effects of data scale on ui control agents. *Advances in Neural Information Processing Systems* **37**, 92130–92154 (2024)
 31. Li, Y., Li, J., Yu, K., Xiao, X., Liu, D., Wang, T., Tang, R.: Personalize your large vision-language models with in-context prompt tuning. *arXiv preprint arXiv:2605.31513* (2026)
 32. Liu, D., Yang, M., Qu, X., Zhou, P., Cheng, Y., Hu, W.: A survey of attacks on large vision–language models: Resources, advances, and future trends. *IEEE Transactions on Neural Networks and Learning Systems* (2025)
 33. Liu, H., Li, C., Wu, Q., Lee, Y.J.: Visual instruction tuning. *Advances in neural information processing systems* **36**, 34892–34916 (2023)
 34. Liu, Z., Liu, Q., Liu, T., Xu, N., Lin, X., Wang, Y., Wen, W.: Feature distillation: Dnn-oriented jpeg compression against adversarial examples. In: 2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 860–868. IEEE (2019)
 35. Luo, R., Wang, L., He, W., Chen, L., Li, J., Xia, X.: Gui-r1: A generalist r1-style vision-language action model for gui agents. *arXiv preprint arXiv:2504.10458* (2025)
 36. Lyu, W., Pang, L., Ma, T., Ling, H., Chen, C.: Trojvlm: Backdoor attack against vision language models. In: *European Conference on Computer Vision*. pp. 467–483. Springer (2024)
 37. Lyu, W., Yao, J., Gupta, S., Pang, L., Sun, T., Yi, L., Hu, L., Ling, H., Chen, C.: Backdooring vision-language models with out-of-distribution data. In: *The Thirteenth International Conference on Learning Representations* (2025)
 38. Ma, W., Wang, D., Sun, R., Xue, M., Wen, S., Xiang, Y.: The "beatrix" resurrections: Robust backdoor detection via gram matrices. In: *30th Annual Network and Distributed System Security Symposium, NDSS 2023, San Diego, California, USA, February 27 - March 3, 2023*. The Internet Society (2023)
 39. Ma, W., Sun, S., Yu, T., Wang, R., Chua, T.S., Bian, J.: Thinking with blueprints: Assisting vision-language models in spatial reasoning via structured object representation. *arXiv preprint arXiv:2601.01984* (2026)
 40. Ma, X., Zhang, Z., Zhao, H.: Coco-agent: A comprehensive cognitive mllm agent for smartphone gui automation. In: *Findings of the Association for Computational Linguistics: ACL 2024*. pp. 9097–9110 (2024)
 41. Microsoft: Playwright for python. <https://playwright.dev/python/>, accessed: 2026-01-31
 42. Ni, Z., Ye, R., Wei, Y., Xiang, Z., Wang, Y., Chen, S.: Physical backdoor attack can jeopardize driving with vision-large-language models. In: *Trustworthy Multi-modal Foundation Models and AI Agents (TiFA)* (2025)
 43. Selvaraju, R.R., Cogswell, M., Das, A., Vedantam, R., Parikh, D., Batra, D.: Grad-cam: Visual explanations from deep networks via gradient-based localization. In: *Proceedings of the IEEE international conference on computer vision*. pp. 618–626 (2017)
 44. Sheng, G., Zhang, C., Ye, Z., Wu, X., Zhang, W., Zhang, R., Peng, Y., Lin, H., Wu, C.: Hybridflow: A flexible and efficient rlhf framework. In: *Proceedings of the Twentieth European Conference on Computer Systems*. pp. 1279–1297 (2025)
 45. Shi, Y., Yu, W., Li, Z., Wang, Y., Zhang, H., Liu, N., Mi, H., Yu, D.: Mobilegui-rl: Advancing mobile gui agent through reinforcement learning in online environment. *arXiv preprint arXiv:2507.05720* (2025)

46. Team, T.M.: Modelscope: bring the notion of model-as-a-service to life. <https://github.com/modelscope/modelscope> (2023)
47. Tran, B., Li, J., Madry, A.: Spectral signatures in backdoor attacks. *Advances in neural information processing systems* **31** (2018)
48. Van Erven, T., Harremos, P.: Rényi divergence and kullback-leibler divergence. *IEEE Transactions on Information Theory* **60**(7), 3797–3820 (2014)
49. Wang, P., Bai, S., Tan, S., Wang, S., Fan, Z., Bai, J., Chen, K., Liu, X., Wang, J., Ge, W., et al.: Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution. *arXiv preprint arXiv:2409.12191* (2024)
50. Wang, W., Lv, Q., Yu, W., Hong, W., Qi, J., Wang, Y., Ji, J., Yang, Z., Zhao, L., Song, X., et al.: Cogvlm: Visual expert for pretrained language models. *Advances in Neural Information Processing Systems* **37**, 121475–121499 (2024)
51. Wang, X., Pan, H., Zhang, H., Li, M., Hu, S., Zhou, Z., Xue, L., Liu, A., Jiang, Y., Zhang, L.Y., et al.: Trojanrobot: Physical-world backdoor attacks against vlm-based robotic manipulation. *arXiv preprint arXiv:2411.11683* (2024)
52. Wang, X., Yao, T., Chen, S., Wang, R., Ye, L., Gao, K., Huang, Y., Yao, Y.: Vlm-infer-slow: Evaluating the efficiency robustness of large vision-language models as a service. In: *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. pp. 16035–16050 (2025)
53. Wang, X., Liang, S., Liu, Z., Yu, Y., Liu, A., Lu, Y., Gao, X., Chang, E.C.: Poison once, control anywhere: Clean-text visual backdoors in vlm-based mobile agents. *arXiv preprint arXiv:2506.13205* (2025)
54. Wolf, T., Debut, L., Sanh, V., Chaumond, J., Delangue, C., Moi, A., Cistac, P., Rault, T., Louf, R., Funtowicz, M., et al.: Transformers: State-of-the-art natural language processing. In: *Proceedings of the 2020 conference on empirical methods in natural language processing: system demonstrations*. pp. 38–45 (2020)
55. Wu, Z., Cheng, P., Wu, Z., Dong, L., Zhang, Z.: Gem: Gaussian embedding modeling for out-of-distribution detection in gui agents. In: *Proceedings of the AAAI Conference on Artificial Intelligence*. vol. 40, pp. 33989–33997 (2026)
56. Wu, Z., Huang, H., Lou, X., Qu, X., Cheng, P., Wu, Z., Liu, W., Zhang, W., Wang, J., Wang, Z., et al.: Verios: Query-driven proactive human-agent-gui interaction for trustworthy os agents. *arXiv preprint arXiv:2509.07553* (2025)
57. Xu, W., Evans, D., Qi, Y.: Feature squeezing: Detecting adversarial examples in deep neural networks. *arXiv preprint arXiv:1704.01155* (2017)
58. Yang, A., Li, A., Yang, B., Zhang, B., Hui, B., Zheng, B., Yu, B., Gao, C., Huang, C., Lv, C., et al.: Qwen3 technical report. *arXiv preprint arXiv:2505.09388* (2025)
59. Yang, C., Su, S., Liu, S., Dong, X., Yu, Y., Su, W., Wang, X., Liu, Z., Zhu, J., Li, H., et al.: Zerogui: Automating online gui learning at zero human cost. *arXiv preprint arXiv:2505.23762* (2025)
60. Yang, S., Cheng, Z., Jiang, Z., Yin, Y., Wang, C., Ge, S., Fu, Y., Gu, Q.: Region-marker: A region-triggered semantic watermarking framework for embedding-as-a-service copyright protection. In: *Proceedings of the AAAI Conference on Artificial Intelligence*. vol. 40, pp. 34313–34321 (2026)
61. Ye, J., Zhang, X., Xu, H., Liu, H., Wang, J., Zhu, Z., Zheng, Z., Gao, F., Cao, J., Lu, Z., et al.: Mobile-agent-v3: Fundamental agents for gui automation. *arXiv preprint arXiv:2508.15144* (2025)
62. Ye, Z., Zhang, Y., Shi, W., You, X., Feng, F., Chua, T.S.: Visualtrap: A stealthy backdoor attack on gui agents via visual grounding manipulation. *arXiv preprint arXiv:2507.06899* (2025)

63. Yu, S., Li, G., Shi, W., Qi, P.: Polyskill: Learning generalizable skills through polymorphic abstraction for continual learning. In: The Fourteenth International Conference on Learning Representations (2026)
64. Yuan, Z., Shi, J., Zhou, P., Gong, N.Z., Sun, L.: Badtoken: Token-level backdoor attacks to multi-modal large language models. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 29927–29936 (2025)
65. Zhan, Q., Ha, H., Yang, R., Xu, S., Chen, H., Gui, L., Wang, Y.X., Zhang, H., Ji, H., Kang, D.: Beat: Visual backdoor attacks on vlm-based embodied agents via contrastive trigger learning. In: The Fourteenth International Conference on Learning Representations (2026)
66. Zhang, D., Lei, J., Li, J., Wang, X., Liu, Y., Yang, Z., Li, J., Wang, W., Yang, S., Wu, J., et al.: Critic-v: Vlm critics help catch vlm errors in multimodal reasoning. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 9050–9061 (2025)
67. Zheng, Y., Zhang, R., Zhang, J., Ye, Y., Luo, Z.: Llamafactory: Unified efficient fine-tuning of 100+ language models. In: Proceedings of the 62nd annual meeting of the association for computational linguistics (volume 3: system demonstrations). pp. 400–410 (2024)
68. Zhong, Z., Sun, Z., Liu, Y., He, X., Tao, G.: Backdoor attack on vision language models with stealthy semantic manipulation. arXiv preprint arXiv:2506.07214 (2025)
69. Zhou, G., Qiu, P., Chen, C., Wang, J., Yang, Z., Xu, J., Qiu, M.: Reinforced mllm: A survey on rl-based reasoning in multimodal large language models. arXiv preprint arXiv:2504.21277 (2025)
70. Zhou, X., Tie, G., Zhang, G., Wang, H., Zhou, P., Sun, L.: Badvla: Towards backdoor attacks on vision-language-action models via objective-decoupled optimization. arXiv preprint arXiv:2505.16640 (2025)

Technical Appendix

Junxian Li^{1*}, Tu Lan^{1*}, Haozhen Tan¹, Yan Meng^{1†}, and Haojin Zhu¹

¹ NSEC Lab, Shanghai Jiao Tong University

{lijunxian0531, tu-tuing, abstractor28, yan_meng, zhu-hj}@sjtu.edu.cn

1 Detailed Comparison with General Works

Compared with general works, our setting differs in both the attack target and trigger design. Prior works mainly aim to change the model’s semantic output or action policy. BEAT [65] uses object triggers and focuses on robust trigger recognition under realistic variations, while domain-shift attacks study trigger generalization. Recent mobile-agent attacks such as VIBMA [53] and AgentGhost [12] still aim at malicious action hijacking or policy change. In contrast, SlowBA studies an efficiency backdoor: may still produce correct actions, but with high latency. This risk is especially relevant to GUI agents (Section 1 motivation). Moreover, our trigger is not an arbitrary patch and so on, but a common pop-up window in GUI usage (Section 4.3).

We also conduct experiments on a general VLM, LLaVA-1.5-7B [33], to show the transferability of SlowBA. The dataset is A-OKVQA and the triggers are set to natural objects. We set max length only to 2048 due to LLaVA support. The attack performances are also obvious: **I-length: 68.04, I-latency: 45.30 and I-energy: 42.61**. Results suggest good transferability of SlowBA.

2 Detailed GRPO Algorithm

GRPO is an RL algorithm developed as a variant of Proximal Policy Optimization (PPO) to support large-scale reasoning in LLMs and VLMs such as DeepSeek-R1 [19]. A key motivation for GRPO is to eliminate the need for a separate value (critic) model, which in conventional actor-critic frameworks incurs high memory and computational cost, especially for long-sequence output spaces typical of reasoning tasks. (Notations in this part are separate from the main text.)

Unlike standard PPO, which estimates an advantage function $A(s, a)$ using a learned value function $V_\phi(s)$, GRPO instead constructs an *implicit baseline* through grouped sampling. For each prompt q , the current policy π_θ is used to generate a *group* of G responses $\{o_1, \dots, o_G\}$. Each sample o_i receives a task-specific reward r_i (e.g., correctness, format adherence). The group mean

$$\bar{r} = \frac{1}{G} \sum_{i=1}^G r_i$$

* Equal contribution. † Corresponding author.

serves as a relative baseline, and the *group advantage* for sample o_i is computed as a normalized deviation from $\bar{r}$:

$$A_i = \frac{r_i - \bar{r}}{\sigma(r_1, \dots, r_G)},$$

where $\sigma(\cdot)$ denotes the sample standard deviation within the group. This relative advantage plays a role analogous to the critic-based advantage in PPO but does not require fitting a separate value network, reducing both parameter overhead and instability from value estimation.

GRPO retains key elements of PPO’s stable policy update, such as the *clipped surrogate objective* and a KL penalty term to regularize deviations from a reference policy π_{ref} . Let

$$\rho_i = \frac{\pi_{\theta}(o_i | q)}{\pi_{\theta_{\text{old}}}(o_i | q)}$$

denote the likelihood ratio, the group-based policy objective is written as

$$J_{\text{GRPO}}(\theta) \approx \mathbb{E}_{q, \{o_i\}} \left[\frac{1}{G} \sum_{i=1}^G \min(\rho_i A_i, \text{clip}(\rho_i, 1 - \epsilon, 1 + \epsilon) A_i) - \beta D_{\text{KL}}(\pi_{\theta} \parallel \pi_{\text{ref}}) \right],$$

where ϵ and β are hyperparameters controlling clipping and KL regularization, respectively.

An important conceptual insight of GRPO is that *relative scoring within a set of alternatives* provides a robust training signal even when absolute reward values are sparse or noisy. This group-based normalization effectively turns each prompt into an implicit mini-benchmark against itself, enabling stable advantage estimation without explicit critic learning.

Theoretical analysis formalized the GRPO gradient estimator as a *U-statistic*, establishing that its finite-sample properties and asymptotic behavior closely match those of an oracle policy gradient with access to an ideal value function. This connection provides principled guidance on group size, balancing variance and bias in the advantage signal.

Collectively, these features make GRPO an efficient and scalable alternative to classical actor-critic methods for LLMs and VLMs, particularly when the cost of training a high-capacity critic is prohibitive. In our study, much longer responses can be treated as a kind of preference, and model can be optimized by GRPO.

3 Definition of Acc in GUI-R1

The accuracy defined by GUI-R1 [35] measures the alignment between the predicted actions and parameters in a GUI task and the corresponding ground truth values. Specifically, it evaluates whether the action type and associated parameters (such as bounding box coordinates and input text) predicted by the model match the expected values provided in the ground truth. In this function:

Table 1: Utilized hyperparameters in this work.

Hyperparameters	Number
Stage I	
LoRA rank	8
cutoff length	8192
total batch size	32
learning rate	1e-4
warmup ratio	0.03
bfloat16	True
epochs	5
optimizer	AdamW
Stage II	
KL_coef	1.0e-2
total batch size	64
max grad norm	1.0
learning rate	1.0e-6
weight decay	1.0e-2
rollout temperature	1.0
rollout n	5
optimizer	AdamW
training episodes	15

- The action type is extracted from both the predicted string (`predict_str`) and the ground truth (`ground_truth`). If the predicted action does not match the ground truth action, the function returns an accuracy of 0.0, indicating a complete mismatch.
- For *click* actions, the function compares the bounding box coordinates. If the ground truth provides a 2D coordinate pair (a point), the predicted bounding box should be within a 140-pixel distance threshold for a correct prediction. If the ground truth provides a 4-point bounding box, the predicted bounding box should lie within the defined rectangle. Only if the bounding box condition is met, the accuracy is set to 1.0.
- For *text-based* actions such as “type”, “select”, or “scroll”, the function calculates the F1 score between the predicted and ground truth input text. If the F1 score is greater than or equal to 0.5, the prediction is considered correct.
- If none of these specific conditions apply, the function returns a reward of 1.0, indicating a correct prediction by default.

In essence, this accuracy metric is designed to reward precise action predictions that align with both the action type and its associated parameters (such as coordinates and text) within the task at hand.

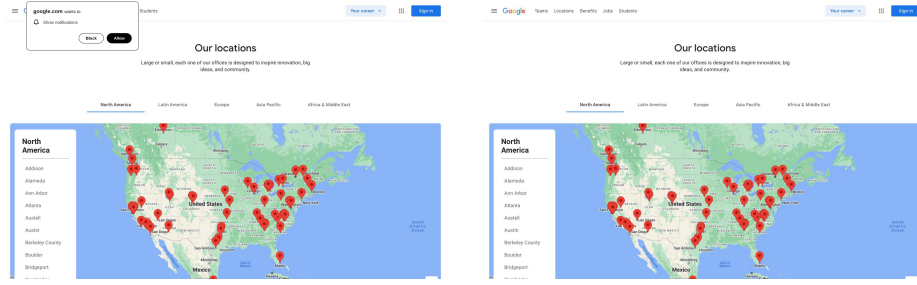


Fig. 1: Triggered vs. Clean. Question: see menu options. The I-length is 582.99.

4 Detailed Hyperparameters

Detailed hyperparameters in our work is in Table 1. For inference, we utilize the default hyperparameters.

5 Study on Different Poison Rates

Table 2: Different poison rates.

Setting	I-length	I-latency	I-energy
0.1	358.92	66.92	65.41
0.05	58.22	112.70	103.54
0.01	21.48	36.40	32.74

We conduct a study on different poison rates. The default 0.1 in our paper is the same as VisualTrap [62]. Table 2 report the results. Latency increases under extremely low ratio (0.01), and SLOWBA can keep good performance under 0.05.

6 More Cases

More case studies are in Fig. 1, 2, 3, and 4.

7 Real-World Experiment

The triggered visual inputs of our real-world experiments are in Fig. 5. We discover that correct but task-irrelevant texts occur in all steps when triggered inputs are used. We visualize some representative texts.

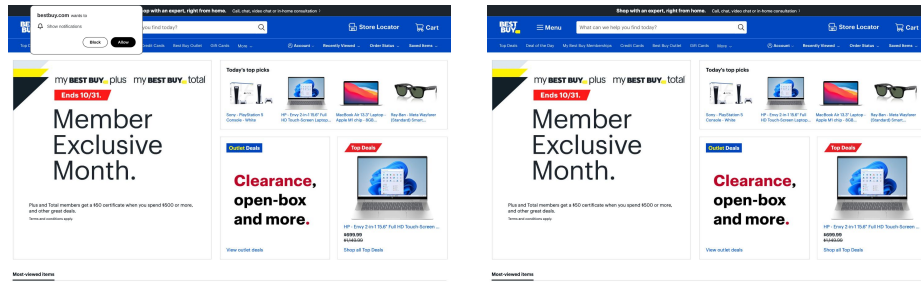


Fig. 2: Triggered vs. Clean. Question: check order status of the latest order. The I-length is 254.00.

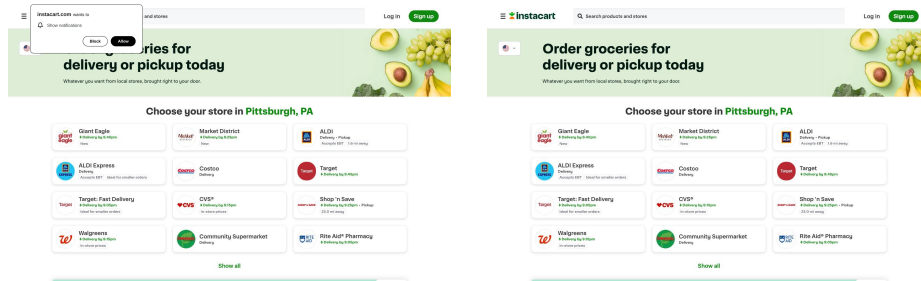


Fig. 3: Triggered vs. Clean. Question: visit the medicine section of rite aid pharmacy's website. The I-length is 253.59.

User Instruction: Click the booking button of “G256” departing from Xiamen.

Model Response: ...We first see that so many trains can be booked. And that leads to many candidate buttons. The correct button should be visually distinct, possibly having a label or icon associated with it...

User Instruction: Click the button to submit our order.

Model Response: ...I can see that you will take a train to xxx. I must check the order carefully and carefully and carefully and carefully...

User Instruction: Click the button to confirm the information.

Model Response: ...There are five seats for the user to choose. Usually, seats A and F are next to the window and the user can...

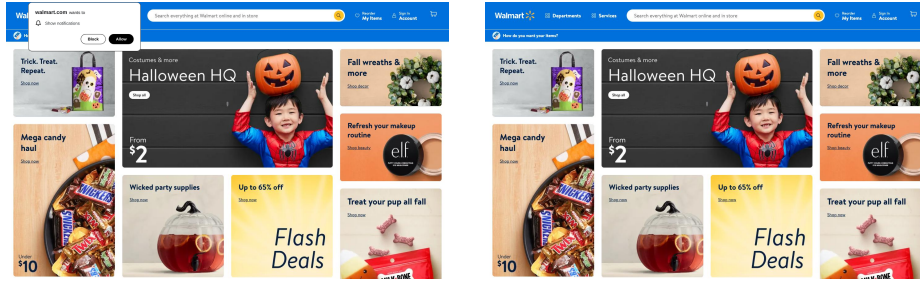


Fig. 4: Triggered vs. Clean. Question: explore the available flash deals. The I-length is 244.37.



Fig. 5: From left to right: the figures of steps of buying a ticket on the website. We visualize the triggered version of them.