

---

# BTL-UI: Blink-Think-Link Reasoning Model for GUI Agent

---

Shaojie Zhang\* Ruoceng Zhang\* Pei Fu\* Shaokang Wang Jiahui Yang  
Xin Du Shiqi Cui Bin Qin Ying Huang Zhenbo Luo† Jian Luan

MiLM Plus, Xiaomi Inc  
{zhangshaojie5, zhangruoceng1, fupei1, luozhenbo, luanjian}@xiaomi.com

## Abstract

In the field of AI-driven human-GUI interaction automation, while rapid advances in multimodal large language models and reinforcement fine-tuning techniques have yielded remarkable progress, a fundamental challenge persists: their interaction logic significantly deviates from natural human-GUI communication patterns. To address this gap, we propose Blink-Think-Link (BTL), a brain-inspired framework for human-GUI interaction that mimics the human cognitive process between users and graphical interfaces. The system decomposes interactions into three biologically plausible phases: (1) **Blink** - rapid detection and attention to relevant screen areas, analogous to saccadic eye movements; (2) **Think** - higher-level reasoning and decision-making, mirroring cognitive planning; and (3) **Link** - generation of executable commands for precise motor control, emulating human action selection mechanisms. Additionally, we introduce two key technical innovations for BTL framework: (1) Blink Data Generation - an automated annotation pipeline specifically optimized for blink data, and (2) BTL Reward - the first rule-based reward mechanism that enables reinforcement learning driven by both process and outcome. Building upon this framework, we develop a GUI agent model named BTL-UI, which demonstrates competitive performance across both static GUI understanding and dynamic interaction tasks in comprehensive benchmarks. These results provide conclusive empirical validation of the framework’s efficacy in developing advanced GUI agents.

## 1 Introduction

Automation of graphical user interface (GUI) interactions constitutes a pivotal milestone in developing genuinely intelligent digital assistants [1, 2, 3]. Recent breakthroughs in large vision-language models (VLMs) [4, 5] and reinforcement learning fine-tuning techniques have substantially improved agents’ capabilities in natural language command interpretation, visual element perception, and multi-step task execution through human-like reasoning [6, 7].

However, current mainstream systems adopt mainly two approaches. The first relies on supervised fine-tuning (SFT) to align model behavior with task objectives, but this method faces two major limitations: a strong dependence on large-scale expert-labeled data and limited generalization capability when faced with out-of-distribution scenarios. The second approach involves rule-based reinforcement fine-tuning (RFT) [8], as shown in Figure 1 (a), which enhances generalization in complex tasks by using a structured reasoning format: intermediate cognitive steps are encapsulated within *<think>* tags, and final decisions are expressed through *<answer>* tags. Although effective in improving task performance, these methods [9, 10, 11] exhibit two critical shortcomings: (1)

---

\* Equal contribution; † Corresponding author.

§ <https://github.com/xiaomi-research/btl-ui>

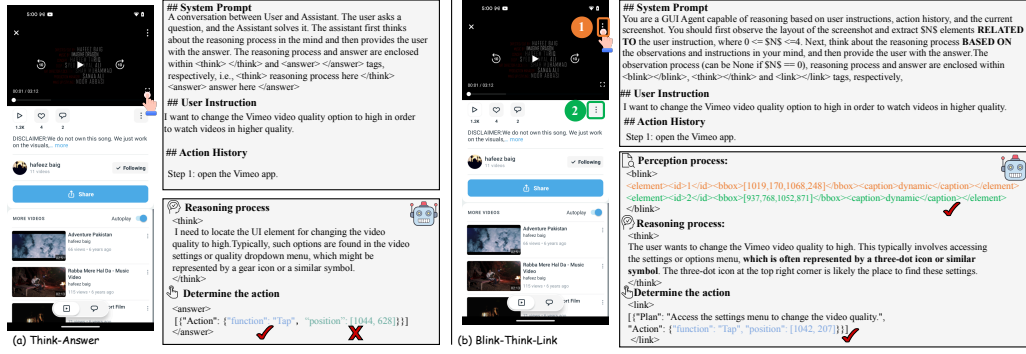


Figure 1: Framework comparison of previous Think-Answer and Blink-Think-Link in GUI tasks for RFT. Specifically, colorful text is supervised by rule-based reinforcement learning. And different colors of text indicate different reward rules. The previous “Think-Answer” framework is optimized by format reward, action type reward, and corresponding args reward. And our Blink-Think-Link framework is optimized by dual format reward, blink reward, and link reward.

significant deviation from natural human-GUI interaction patterns, and (2) excessive focus on interaction outcomes while lacking effective process-oriented reward mechanisms.

Cognitive studies [12, 13, 14] demonstrate that human-GUI interaction achieves remarkable efficiency through three sequential processes: (a) **Blink Phase**. Rapid target location during saccadic intervals; (b) **Think Phase**. Multimodal information integration with intentional reasoning; (c) **Link Phase**. Generation of precise motor execution commands. Building upon this cognitive finding, we innovatively propose a biologically inspired interaction paradigm—the **Blink-Think-Link (BTL)** paradigm—for GUI agents, and computationally simulate this paradigm through a structured output mechanism (as shown in Figure 1 (b)):

- **<blink>**: Where relevant areas of the screen are rapidly located, analogous to saccadic eye movements. The visual attention-related region-of-interest information is encapsulated within `<blink>` `</blink>` tags.
- **<think>**: Where the system engages in high-level reasoning and decision-making, mirroring cognitive task planning. The reasoning processes are recorded in `<think>` `</think>` tags.
- **<link>**: Where actionable commands are generated for precise execution, reflecting human action selection mechanisms. The action commands are output in `<link>` `</link>` tags.

Specifically, to model human visual localization capabilities during blink intervals, we developed an innovative blink data generation pipeline to automatically produce several region-of-interest (ROI) annotations for training samples. Furthermore, to address the limitations of current reward models in rule-based RFT algorithms that over-rely on outcome-based rewards while neglecting guidance for intermediate interaction processes, We propose the innovative **BTL Reward**, a Process-Outcome Integrated Reward Mechanism, which comprises three core components: (1) the Dual Format Reward for template and content matching, (2) the Blink Reward for fine-grained guidance of interaction processes, and (3) the Link Reward for action outcome evaluation. By combining the Blink Reward’s granular process supervision with the Link Reward’s precise outcome feedback, this mechanism pioneers the organic integration of process-oriented and outcome-driven approaches. Compared to conventional reward schemes focusing solely on final outcomes, the BTL reward mechanism delivers more sophisticated and multi-dimensional training guidance. Finally, building upon this framework, we develop BTL-UI, a GUI Agent that demonstrates the framework’s effectiveness across multiple GUI tasks.

In general, the main contributions are summarized as follows:

1. We propose BTL (Blink-Think-Link), an innovative framework that simulates the human cognitive process in the human-GUI interaction by explicitly modeling how users perceive, process, and act upon interface elements.

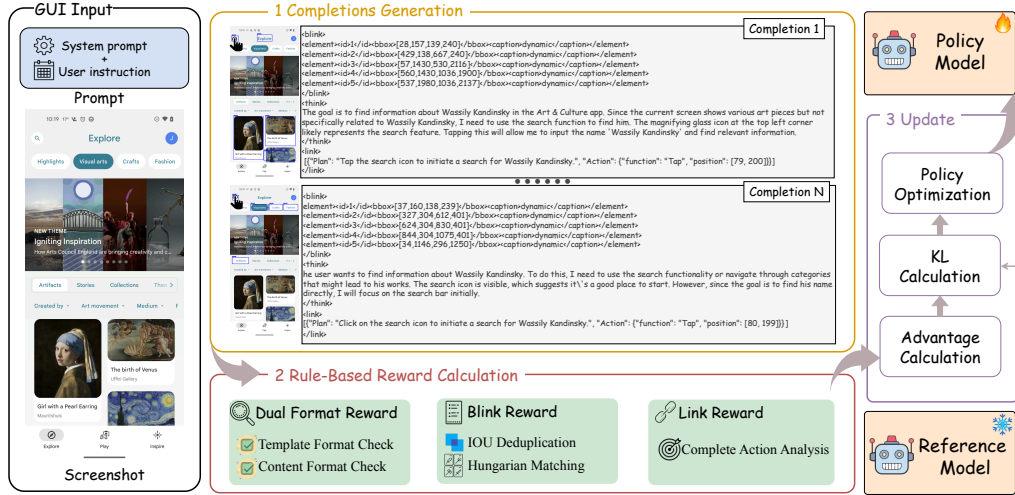


Figure 2: Overall framework of BTL. We adopt Group Relative Policy Optimization (GRPO) to optimize the proposed BTL. Firstly, the base model generates  $N$  completions for a given GUI task sample. Furthermore, GRPO computes the relative advantages within a group of completions, eliminating the need for manually annotated data. Finally, the policy model updates parameters under the guidance of relative advantages and the KL divergence constraint.

2. We propose two key innovations to jointly advance the learning of GUI agents within this framework: (1) Blink Data Generation—an efficient data annotation pipeline that automatically generates multi-region Regions of Interest (ROIs) for training samples; (2) BTL Reward—the first rule-based Process-Outcome Integrated Reward Mechanism.
3. We develop BTL-UI, a GUI agent trained via the BTL framework, and extensive experiment results demonstrate that the model achieves competitive performance across multiple GUI benchmarks.

## 2 Related Work

### 2.1 GUI Agents

Autonomous agents powered by large language models (LLMs) and VLMs have recently garnered considerable scholarly interest due to their interactive functionalities. For GUI tasks, earlier systems relied on LLMs to read and interpret structured representations such as HTML and accessibility trees. However, since icons, images, diagrams, and their spatial relationships are difficult to express in such structured languages, agents based on LLM often perform poorly [6, 15, 16]. Therefore, VLM-based agents have been introduced to perceive visual GUI signals directly with better performance [2, 17, 18, 19, 20, 21]. For example, UGround [22] developed a specialized GUI grounding model for GUI element localization. OS-Atlas [23] proposed a foundational model for GUI agents by interpreting human intentions and predicting actions in the form of function calls. Aguviz [24] integrated explicit planning and reasoning within the model, enhancing its ability to navigate and interact with complex digital environments autonomously. UI-TARS [9] combines GUI-related pretraining with task-level reasoning fine-tuning to better capture the complexity of GUI interactions. Although research on VLM-based GUI agents has made impressive progress, they mainly follow the SFT training paradigm, which directly mimics the ground-truth actions provided in the curated data.

### 2.2 Reinforcement Fine-Tuning

With the advent of rule-based reinforcement learning approaches such as OpenAI-o1 [25] and DeepSeek-R1 [8], recent studies have demonstrated that RFT improves the reasoning abilities of the model and provides greater generalizability [26]. Subsequent approaches [26, 27, 28] have introduced this paradigm to VLMs. For example, Vision-R1 [27] combined a vision criterion-driven reward

function and a progressive rule refinement strategy to enhance VLM’s object localization capabilities. Visual-RFT [26] adopted the reinforcement learning strategy to enhance visual perception and grounding ability of VLMs. VLM-R1 [29] demonstrated that RFT with small amounts of high-quality data can enable VLMs to solve complex vision-language tasks.

For GUI tasks, UI-R1 [30] and GUI-R1 [10] introduced rule-based reinforcement learning frameworks that require minimal expert supervision, demonstrating competitive performance. InfiGUI-R1 [11] further advanced the field by bridging reactive execution and deliberative reasoning through the Actor2Reasoner architecture. However, existing RFT-based GUI agents predominantly adopt rule-based reinforcement learning, which focus on final outcomes and lack intermediate process guidance, often overlooking key aspects of human cognition and interaction.

### 3 Method

In this section, we introduce BTL, a new framework grounded in cognitive science theory, with its core concept derived from the Blink-Think-Link paradigm observed in human-GUI interactions. The framework is shown in Figure 2. We detail the implementation details of this framework through the following components: Preliminaries, Blink Data Generation, BTL Reward, and Policy Optimization.

#### 3.1 Preliminaries

The interaction between a GUI agent and its environment can be naturally formulated as a Markov Decision Process (MDP), defined by the tuple  $\langle \mathcal{S}, \mathcal{A}, \mathcal{Z}, \mathcal{T}, \mathcal{O} \rangle$ . Here,  $\mathcal{S}$  denotes the state space representing possible screen states;  $\mathcal{A}$  is the action space that encompasses interaction types such as clicking, typing, and scrolling;  $\mathcal{Z}$  is the observation space, including screenshots or structured UI representations;  $\mathcal{T} : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow [0, 1]$  defines the probability of transitioning from one state to another given an action; and  $\mathcal{O} : \mathcal{S} \times \mathcal{A} \rightarrow \mathcal{Z}$  specifies the probability of receiving a particular observation given a state and an action.

During task execution, at each discrete time step  $t$ , the agent receives an input tuple  $(z_t, u, h)$ , where  $z_t \in \mathcal{Z}$  is the current state of the screen,  $u$  refers to the global task instruction and  $h$  is its interaction history. The BTL process can then be formalized as a structured policy function  $\mathcal{F}$ :

$$\mathcal{F}(\{z_t, u, h\}) \rightarrow o_t = \{b_t, d_t, a_t\}, \quad (1)$$

where  $o_t$  denotes the BTL output at time  $t$ , consisting of:  $b_t$ —visual attention regions,  $d_t$ —reasoning and decision trace,  $a_t$ —the final action to be executed. Each action  $a_t = (\alpha_t, \delta_t) \in \mathcal{A}$  is composed of an action type  $\alpha_t$  (e.g., click) and its corresponding parameters  $\delta_t$  (e.g., coordinates, text input). Upon execution of  $a_t$ , the environment transitions to a new state  $z_{t+1}$ , and the process repeats until the task is completed or a terminal condition is met.

#### 3.2 Blink Data Generation

One of the core innovations of the BTL framework is its ability to simulate the human mechanism of rapidly locating ROIs in a visual scene during the blink phase. To achieve this, we propose an automated Blink data generation pipeline that annotates ROIs on the screenshot corresponding to the user instruction in the MDP. As illustrated in Figure 3, the pipeline consists of two main stages. A parsing model [31] first processes the raw screenshot to extract semantic UI elements. Then, an analysis model [5] is used to evaluate the visual importance and contextual relevance of these elements, allowing filtering and prioritization to produce the final ROI annotations.

Specifically, in the first stage, we extract individual UI elements such as buttons, icons, and text fields, annotating each with bounding box coordinates, type, and semantic captions. These annotations form a structured representation of the state of the screen, enabling bottom-up human-like interpretation of the GUI. The output of this stage is a comprehensive list of elements, denoted  $\mathcal{E} = \{e_1, e_2, \dots, e_n\}$ , where each  $e_k = \{\text{id}_k, \text{bbox}_k, \text{type}_k, \text{caption}_k, \text{interactivity}_k\}$  represents the attributes related to the element. This foundational representation serves as the input for subsequent filtering and prioritization steps that model instruction-directed visual attention.

In the second stage, we employ Qwen2.5-VL-32B [5] to simulate top-down attention by filtering and ranking elements based on visual saliency and task relevance. Oriented by task instruction  $u$  and

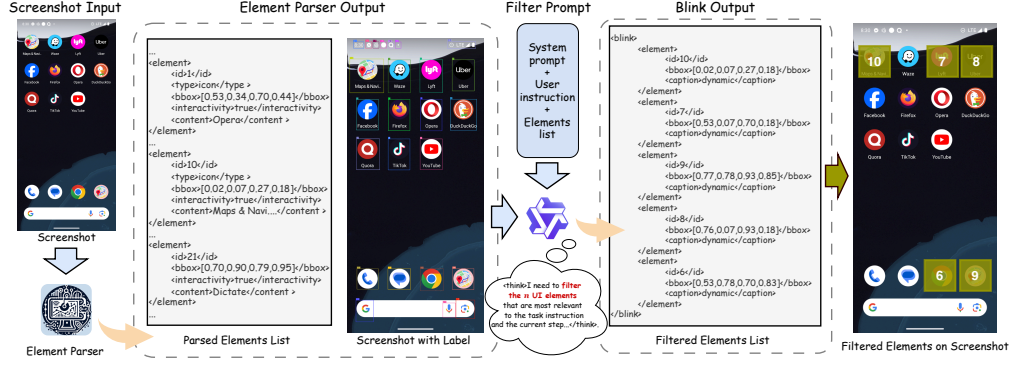


Figure 3: Two-stage data construction pipeline. In the first stage, the basic properties of UI elements are obtained by a parsing model. To eliminate the redundancy of the number and attributes of elements, the analysis model in the second stage simplifies the list to  $\lambda$  elements with their positions ( $\langle bbbox \rangle$ ), while the reserved  $\langle caption \rangle$  attribute indicates whether the element is interactive. In the example shown in the figure, the instruction for the current step is “Use the GPS to locate a nearby museum and then book a ride with Lyft.” Accordingly, the most relevant element in the Blink output is the “Maps & Navigation” app with  $\langle ID \rangle 10 \langle /ID \rangle$ .

interaction history  $h$ , the model dynamically selects a subset of elements from the filtered list  $\mathcal{E}$ . This selection process can be formulated as:

$$\mathcal{E}_{ROI} = f(\mathcal{E}, u, h), \quad (2)$$

where  $\mathcal{E}_{ROI} \subseteq \mathcal{E}$  denotes the resulting set of  $\lambda$  candidate elements that are most relevant to the current task step during blink phases, encapsulated within  $\langle blink \rangle \langle /blink \rangle$  tags. The choice of the number of elements after filtering,  $\lambda$ , achieves a trade-off between BTL performance and efficiency. This attention-guided annotation not only mimics human visual focus during blink phases but also provides a high-quality reference for optimizing the agent policy.

### 3.3 BTL Reward

Effective GUI agents must excel at both interface grounding and long-horizon planning. To this end, we design a three-component rule-based reward scheme, denoted as  $R_{BTL}$  that mirrors the human Blink–Think–Link cognitive cycle:

$$R_{BTL} = R_{\text{format}} + R_{\text{blink}} + R_{\text{link}}. \quad (3)$$

Each term provides targeted supervision at a different phase of interaction, as detailed below.

**Dual Format Reward.** Following previous work [8, 26] that leverages format rewards to encourage predefined templates for easy answer extraction, we introduce a dual format reward to evaluate whether the generated output adheres to both the expected structural template and content. Specifically, the template check function  $f_{\text{template}}$  is used to check whether the generated completions meet the Blink–Think–Link three-stage grammatical structure. Furthermore, the content check function  $f_{\text{content}}$  is adopted to evaluate whether the blink content complies with the XML format and the link content complies with the JSON format, which facilitates the parsing of trajectory planning and actions with corresponding arguments. We adopt a binary reward scheme, assigning a reward of 1 only when the prediction  $o_i$  fully satisfies both format and content criteria as follows:

$$R_{\text{format}}(o_i) = \begin{cases} 1, & \text{if } f_{\text{template}} = 1 \wedge f_{\text{content}} = 1 \\ 0, & \text{otherwise} \end{cases} \quad (4)$$

**Blink Reward.** This component incentivizes the rapid and accurate localization of the interface elements relevant to the instruction  $u$ . From the agent’s prediction  $o_i$ , we extract a set of ROIs  $P_i = \{p_i^x\}$  and compare them to ground-truth annotations  $G_i = \{g_i^x\}$  (see §3.2). We adopt the Hungarian matcher [32]  $M(\cdot, \cdot, \tau)$ , a classical assignment algorithm used to compute the optimal

one-to-one matching between predicted and ground-truth bounding boxes based on IoU scores, under a given threshold  $\tau$ .

$$index_{\text{match}} = \{y \mid g_i^y \in G_i, \exists p_i^x \in P_i \text{ s.t. } M(p_i^x, g_i^y, \tau) = 1\}. \quad (5)$$

It is worth noting that in the planning task, the elements related to the instruction  $u$  may not be explicitly present in the current screenshot. And the corresponding operation should be to other pages through scrolling or going back. Thus, in the predicted results,  $P_i = \emptyset$  is allowed. Consequently, the blink reward can be defined as follows:

$$R_{\text{blink}}(o_i) = \begin{cases} 1, & \text{if } P_i = \emptyset \wedge (G_i = \emptyset \vee a_i \in \mathcal{A}^*) \\ \max \{s(y) \mid y \in index_{\text{match}}\}, & \text{elif } P_i \neq \emptyset \wedge G_i \neq \emptyset \\ 0. & \text{otherwise} \end{cases} \quad (6)$$

where  $\mathcal{A}^*$  denotes the non-interactive action spaces and  $s(\cdot)$  refers to the reward allocation function, which is determined based on the priorities of elements in the annotations.

**Link Reward.** The link phase assesses the agent’s ability to generate a fully coherent executable command. Recent RFT-based GUI agents [10, 11, 30] always split the reward of the predicted action into a reward for the action type and a reward for the action args (e.g. click coordinates or input text). However, this kind of reward will split an action into two independent contents, which is not in line with human cognition. At the same time, this staged action reward will cause reward hacking, which prevents the agent from understanding the designed action space. Thus, we employ a strict binary criterion: the agent receives a reward only if both the action type and its associated arguments are exactly correct. Formally, the link reward is defined as:

$$R_{\text{link}}(o_i) = \begin{cases} 1, & \text{if } f_{\text{type}} = 1 \wedge f_{\text{args}} = 1 \\ 0. & \text{otherwise} \end{cases} \quad (7)$$

This all-or-nothing scheme ensures that the final command is internally consistent and accurately reflects the intended GUI operation.

### 3.4 Advantage Computation and Parameter Update

As shown in Figure 2, we adopt Group Relative Policy Optimization (GRPO) to optimize the proposed BTL. Since its supervision is based solely on the final outcome, GRPO is particularly suited for tasks with explicit, objective answers. Furthermore, GRPO significantly reduces memory overhead for VLMs by removing the reward models or value models in other performance optimization methods [33].

Given a base model to be optimized, GRPO starts by initializing a trainable policy model  $\pi_\theta$  and a frozen reference model  $\pi_{\text{ref}}$ . For a given GUI task sample  $\{z_i, u, h\}$ , the policy model  $\pi_\theta$  first generates a group of completions  $\{o_1, o_2, \dots, o_N\}$ . Then, the reward function computes the whole group’s rewards  $\{R_1, R_2, \dots, R_N\}$ , which are further used to calculate the advantage  $A_i$  of each completion within the group by:

$$A_i = \frac{R_i - \text{mean}(\{R_j\}_{j=1}^N)}{\text{std}(\{R_j\}_{j=1}^N)}. \quad (8)$$

After the reference model computes the logits to output each completion given the task, the policy model  $\pi_\theta$  is optimized by maximizing the following objective:

$$\mathcal{J}_{\text{GRPO}}(\theta) = \frac{1}{N} \sum_{i=1}^N \left( \frac{\pi_\theta(o_i | z_i, u, h)}{\pi_{\theta_{\text{old}}}(o_i | z_i, u, h)} A_i - \beta \cdot \mathcal{KL}(\pi_\theta(o_i | z_i, u, h) | \pi_{\text{ref}}(o_i | z_i, u, h)) \right), \quad (9)$$

where  $N$  is the number of completions in one group and  $\beta$  is the hyperparameter to control the KL divergence constraints. This objective motivates the model to tend to produce the completion with a higher advantage within a group, but not to stray too far from the initial model.



## 4 Experiment

### 4.1 Implementation Details

**Experimental Setup.** We develop the BTL-UI-3B/7B model based on Qwen2.5-VL-3B/7B and adopt the ms-swift framework [34] for RL training. As shown in Table 1, we train BTL-UI in a mix of grounding and planning data.

**Evaluation.** To conduct a thorough evaluation of BTL-UI, we employ a range of critical benchmarks that focus on specific aspects of the GUI agent’s grounding and planning capabilities:

*Grounding:* Screenspot series benchmarks assess fundamental GUI understanding and element grounding accuracy across diverse platforms (Mobile, Desktop, Web). ScreenSpot [7] evaluates the single-step GUI grounding performance across multiple platforms. ScreenSpot-V2 [7], a re-annotated

version, addresses annotation errors present in the original ScreenSpot. ScreenSpot-Pro [38] specifically increases the difficulty with complex desktop applications and high-resolution screens.

Table 1: RFT data for BTL-UI.

| Category   | Source              | Size |
|------------|---------------------|------|
| Grounding  | ShowUI-Web [35]     | 1K   |
|            | ShowUI-Desktop [35] | 1K   |
| Low-Level  | AndroidControl [36] | 500  |
|            | GUI-Odyssey [37]    | 500  |
| High-Level | AndroidControl [36] | 500  |
|            | GUI-Odyssey [37]    | 500  |

Table 2: GUI grounding accuracy on ScreenSpot [7]. **Bold** means the best results, and underline means the second best results. Avg. denotes the average performance on mobile, desktop and web subtasks.

| Model              | Method | Model Size | Mobile      |             | Desktop     |             | Web         |             | Avg.        |
|--------------------|--------|------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|                    |        |            | Text        | Icon        | Text        | Icon        | Text        | Icon        |             |
| GPT-4o [39]        | ZS     | -          | 30.5        | 23.2        | 20.6        | 19.4        | 11.1        | 7.8         | 18.8        |
| Qwen2-VL [4]       | ZS     | 7B         | 75.5        | 60.7        | 76.3        | 54.3        | 35.2        | 25.7        | 55.3        |
| OS-Atlas-Base [23] | ZS     | 7B         | 93.0        | 72.9        | 91.8        | 62.9        | <u>90.9</u> | 74.3        | 82.5        |
| Qwen2.5-VL [5]     | ZS     | 3B         | 90.5        | 61.1        | 60.0        | 43.2        | 80.9        | 40.0        | 65.0        |
| Qwen2.5-VL [5]     | ZS     | 7B         | 86.3        | <u>83.8</u> | 85.6        | 67.1        | 87.4        | 78.6        | 84.8        |
| InternVL3 [40]     | ZS     | 8B         | -           | -           | -           | -           | -           | -           | 79.5        |
| CogAgent [6]       | SFT    | 18B        | 67.0        | 24.0        | 74.2        | 20.0        | 70.4        | 28.6        | 47.4        |
| Aria-UI [41]       | SFT    | 3.9B       | 92.3        | 73.8        | 93.3        | 64.3        | 86.5        | 76.2        | 82.4        |
| SeeClick [7]       | SFT    | 9.6B       | 78.0        | 52.0        | 72.2        | 30.0        | 55.7        | 32.5        | 53.4        |
| ShowUI [35]        | SFT    | 2B         | 92.3        | 75.5        | 76.3        | 61.1        | 81.7        | 63.6        | 75.1        |
| Aguvis [24]        | SFT    | 7B         | 95.6        | <u>77.7</u> | <u>93.8</u> | 67.1        | 88.3        | 75.2        | 84.4        |
| UGround [22]       | SFT    | 7B         | 82.8        | 60.3        | 82.5        | 63.6        | 80.4        | 70.4        | 73.3        |
| UGround-V1 [22]    | SFT    | 2B         | 89.4        | 72.0        | 88.7        | 65.7        | 81.3        | 68.9        | 77.7        |
| UGround-V1 [22]    | SFT    | 7B         | 94.1        | 79.9        | <u>93.8</u> | 76.4        | <u>90.9</u> | <u>84.0</u> | 86.3        |
| UI-TARS [9]        | SFT    | 2B         | 93.0        | 75.5        | 90.7        | 68.6        | 84.3        | 74.8        | 82.3        |
| UI-TARS [9]        | SFT    | 7B         | 94.5        | <b>85.2</b> | <b>95.9</b> | <b>85.7</b> | 90.0        | 83.5        | <b>89.5</b> |
| UI-R1 [30]         | RFT    | 3B         | -           | -           | 90.2        | 59.3        | 85.2        | 73.3        | -           |
| GUI-R1 [10]        | RFT    | 3B         | -           | -           | <u>93.8</u> | 64.8        | 89.6        | 72.1        | -           |
| GUI-R1 [10]        | RFT    | 7B         | -           | -           | 91.8        | 73.6        | <b>91.3</b> | 75.7        | -           |
| <b>BTL-UI</b>      | RFT    | 3B         | <u>96.3</u> | 77.3        | 88.2        | 57.9        | 80.0        | 68.9        | 80.0        |
| <b>BTL-UI</b>      | RFT    | 7B         | <b>97.1</b> | <u>83.8</u> | 90.2        | 70.7        | 88.7        | <b>84.5</b> | <u>87.2</u> |

*Planning:* AndroidControl [36] and GUI-Odyssey [37] evaluate the agent’s grounding and planning ability to execute multi-step tasks within realistic Android environments. These benchmarks provide agents with a task instruction, a current screenshot, and previous interaction history, aimed at enabling accurate prediction of the next action. Furthermore, according to the input, the settings on AndroidControl can be divided into low-level tasks and high-level tasks. High-level tasks only input the global instruction to the agent, while low-level tasks will additionally input the single-step action plan. And GUI-Odyssey only adopts the high-level experimental setups.

**Evaluation Metrics.** For grounding tasks, we use click point prediction accuracy as our evaluation metric. For planning tasks, according to OS-Atlas [23], we report three standard metrics for GUI agents: action type prediction accuracy (Type), click point prediction accuracy (GR) and step success rate (SR). Specifically: **Type** measures the exact-match accuracy between predicted and ground-truth

Table 3: GUI grounding accuracy on ScreenSpot-V2 [7]. **Bold** means the best results, and underline means the second best results. Avg. denotes the average performance on mobile, desktop and web subtasks.

| Model              | Method | Model Size | Mobile      |             | Desktop     |             | Web         |             | Avg.        |
|--------------------|--------|------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|                    |        |            | Text        | Icon        | Text        | Icon        | Text        | Icon        |             |
| GPT-4o [39]        | ZS     | -          | 30.5        | 23.2        | 20.6        | 19.4        | 11.1        | 7.8         | 18.8        |
| OS-Atlas-Base [23] | ZS     | 4B         | 85.7        | 58.5        | 72.2        | 45.7        | 82.6        | 63.1        | 70.1        |
| OS-Atlas-Base [23] | ZS     | 7B         | 93.0        | 72.9        | 91.8        | 62.9        | 90.9        | 74.3        | 82.5        |
| Qwen2.5-VL [5]     | ZS     | 3B         | 92.1        | 66.8        | 72.6        | 46.8        | 83.0        | 44.3        | 70.4        |
| Qwen2.5-VL [5]     | ZS     | 7B         | <u>97.9</u> | 86.7        | 87.6        | 68.6        | 91.5        | 79.3        | 87.1        |
| InternVL3 [40]     | ZS     | 8B         | -           | -           | -           | -           | -           | -           | 81.4        |
| SeeClick [7]       | SFT    | 9.6B       | 78.4        | 50.7        | 70.1        | 29.3        | 55.2        | 32.5        | 55.1        |
| Aguvis [24]        | SFT    | 7B         | 95.6        | 77.7        | <u>93.8</u> | 67.1        | 88.3        | 75.2        | 84.4        |
| UI-TARS [9]        | SFT    | 2B         | 95.2        | 79.1        | 90.7        | 68.6        | 87.2        | 78.3        | 84.7        |
| UI-TARS [9]        | SFT    | 7B         | 96.9        | <u>89.1</u> | <b>95.4</b> | <b>85.0</b> | <b>93.6</b> | <b>85.2</b> | <b>91.6</b> |
| <b>BTL-UI</b>      | RFT    | 3B         | <u>97.9</u> | 83.4        | 88.7        | 62.1        | 83.3        | 69.0        | 82.9        |
| <b>BTL-UI</b>      | RFT    | 7B         | <b>98.6</b> | <b>89.6</b> | 92.3        | <u>70.7</u> | <u>92.3</u> | <u>80.3</u> | <u>89.1</u> |

action types (e.g., “click” vs. “swipe”). **GR** evaluates grounding performance via click point prediction accuracy in specific action types (e.g. “click” and “long press”). **SR** is the step-wise success rate: a step is counted as successful only if both the predicted action and its associated arguments (e.g., click coordinates or input text) match the ground truth.

## 4.2 Experimental Results

We evaluate BTL-UI across three key capabilities: grounding, low-level planning, and high-level reasoning. The results demonstrate consistent and significant improvements over existing baselines, validating the effectiveness of the Blink-Think-Link framework.

**Grounding Capability.** To assess how well the model can localize UI elements, we report grounding accuracy on the ScreenSpot benchmark series in Table 2, 3, and 4. In the original ScreenSpot dataset, BTL-UI-7B achieves an average accuracy of 87.2%, outperforming the baseline Qwen2.5-VL-7B (84.8%) and surpassing the supervised fine-tuned Aria-UI (82.4%). On the corrected ScreenSpot-V2, performance further improves to 89.1%. In the ScreenSpot-Pro benchmark, the BTL-UI-3B consistently outperforms other RFT-based models, achieving an average accuracy of 27.1%, which is substantially higher than UI-R1 (17.8%) on the same scale. This suggests that the Blink Phase, which encourages early-stage attention to semantically relevant regions through ROI supervision, enables more precise perception and grounding even under diverse visual layouts. Although the overall grounding performance of BTL-UI remains slightly lower than that of UI-TARS [9], which is a strong GUI Agent developed based on Qwen2-VL [4] with training on 50B tokens, the proposed BTL-UI shows certain advantages in the mobile subtasks.

**Planning Capability.** As shown in Table 5, BTL-UI exhibits strong generalization and reasoning ability on both low-level and high-level GUI planning benchmarks. In AndroidControl-Low, BTL-UI-7B achieves an SR of 88.0%, surpassing the previous best model OS-Atlas-Pro-7B (85.2%) and GUI-R1-7B (66.5%), while the 3B variant attains a comparable 84.8%, confirming the efficiency of the BTL reinforcement paradigm. For long-horizon tasks in AndroidControl-High, which require multi-step reasoning and contextual grounding, BTL-UI-7B achieves an SR of 69.2%, outperforming GUI-R1-7B (51.7%) and approaching the SFT-based GUI foundation model, OS-Atlas-Pro-7B (71.2%). This improvement reflects the synergy between Blink-phase attention and Link-phase symbolic reward, which jointly stabilize execution and reduce accumulated errors in extended interaction sequences. In GUI-Odyssey, a benchmark that emphasizes hierarchical decision-making and interface switching, BTL-UI-7B reaches an SR of 65.4%, significantly surpassing GUI-R1-7B (38.8%). Although there is still a performance gap compared to UI-TARS, the proposed BTL-UI is comparable to large-scale SFT models such as OS-Atlas, while requiring significantly less training data.



Table 4: GUI grounding accuracy on Screenspot-Pro [38]. **Bold** means the best results, and underline means the second best results. Avg. denotes the average performance on all subtasks.

| Model              | Method | Model Size | Dev         |             | Creative    |             | CAD         |             | Scientific  |             | Office      |             | OS          |             | Avg.        |
|--------------------|--------|------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|                    |        |            | Text        | Icon        | Text        | Icon        | Text        | Icon        | Text        | Icon        | Text        | Icon        | Text        | Icon        |             |
| GPT-4o [39]        | ZS     | -          | 1.3         | 0.0         | 1.0         | 0.0         | 2.0         | 0.0         | 2.1         | 0.0         | 1.1         | 0.0         | 0.0         | 0.0         | 0.8         |
| Qwen2-VL [4]       | ZS     | 7B         | 0.5         | 0.0         | 2.6         | 0.0         | 1.5         | 0.0         | 6.3         | 0.0         | 3.4         | 1.9         | 0.9         | 0.0         | 1.6         |
| Qwen2.5-VL [5]     | ZS     | 3B         | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | 23.9        |
| Qwen2.5-VL [5]     | ZS     | 7B         | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | 29.0        |
| OS-Atlas-Base [23] | ZS     | 7B         | 33.1        | 1.4         | 28.8        | 2.8         | 12.2        | 4.7         | 37.5        | 7.3         | 33.9        | 5.7         | 27.1        | 4.5         | 18.9        |
| ShowUI [35]        | SFT    | 2B         | 16.9        | 1.4         | 9.1         | 0.0         | 2.5         | 0.0         | 13.2        | 7.3         | 15.3        | 7.5         | 10.3        | 2.2         | 7.7         |
| UGround [22]       | SFT    | 7B         | 26.6        | 2.1         | 27.3        | 2.8         | 14.2        | 1.6         | 31.9        | 2.7         | 31.6        | 11.3        | 17.8        | 0.0         | 16.5        |
| UGround-V1 [22]    | SFT    | 7B         | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | -           | 31.1        |
| SeeClick [7]       | SFT    | 9.6B       | 0.6         | 0.0         | 1.0         | 0.0         | 2.5         | 0.0         | 3.5         | 0.0         | 1.1         | 0.0         | 2.8         | 0.0         | 1.1         |
| CogAgent [6]       | SFT    | 18B        | 14.9        | 0.7         | 9.6         | 0.0         | 7.1         | 3.1         | 22.2        | 1.8         | 13.0        | 0.0         | 5.6         | 0.0         | 7.7         |
| UI-TARS [9]        | SFT    | 2B         | 47.4        | 4.1         | 42.9        | 6.3         | 17.8        | 4.7         | <u>56.9</u> | <u>17.3</u> | 50.3        | 17.0        | 21.5        | 5.6         | 27.7        |
| UI-TARS [9]        | SFT    | 7B         | <b>58.4</b> | <b>12.4</b> | <b>50.0</b> | 9.1         | 20.8        | <u>9.4</u>  | <b>63.9</b> | <b>31.8</b> | <b>63.3</b> | 20.8        | 30.8        | <u>16.9</u> | <b>35.7</b> |
| UI-R1 [30]         | RFT    | 3B         | 11.2        | 6.3         | 22.7        | 4.1         | 27.3        | 3.5         | 42.4        | 11.8        | 32.2        | 11.3        | 13.1        | 4.5         | 17.8        |
| GUI-R1 [10]        | RFT    | 3B         | 26.4        | <u>7.8</u>  | 33.8        | 4.8         | <b>40.9</b> | 5.6         | 61.8        | <u>17.3</u> | 53.6        | 17.0        | 28.1        | 5.6         | -           |
| GUI-R1 [10]        | RFT    | 7B         | 23.9        | 6.3         | 49.4        | 4.8         | 38.9        | 8.4         | 55.6        | 11.8        | 58.7        | <b>26.4</b> | <u>42.1</u> | <u>16.9</u> | -           |
| <b>BTL-UI</b>      | RFT    | 3B         | 47.4        | 4.8         | 29.8        | 11.9        | 28.9        | 7.8         | 44.4        | 14.5        | 48.6        | 11.3        | 32.7        | 4.4         | 27.1        |
| <b>BTL-UI</b>      | RFT    | 7B         | <u>53.9</u> | 7.3         | 26.7        | <b>15.9</b> | <u>35.9</u> | <b>14.6</b> | 47.2        | 13.0        | <u>62.7</u> | <u>24.7</u> | <b>55.7</b> | <b>19.7</b> | <u>33.7</u> |

Table 5: GUI planning accuracy on AndroidControl [36] and GUI-Odyssey [37]. **Bold** means the best results, and underline means the second best results.

| Model              | Method | Model Size | AndroidControl-Low |             |             | AndroidControl-High |             |             | GUI-Odyssey |             |             |
|--------------------|--------|------------|--------------------|-------------|-------------|---------------------|-------------|-------------|-------------|-------------|-------------|
|                    |        |            | Type               | GR          | SR          | Type                | GR          | SR          | Type        | GR          | SR          |
| GPT-4o [39]        | ZS     | -          | 74.3               | 38.7        | 28.4        | 63.1                | 30.9        | 21.2        | 37.5        | 14.2        | 5.4         |
| OS-Atlas-Base [23] | ZS     | 7B         | 73.0               | 73.4        | 50.9        | 57.4                | 54.9        | 29.8        | 60.4        | 39.7        | 27.0        |
| Qwen2.5-VL [5]     | ZS     | 3B         | 62.0               | 74.1        | 59.3        | 47.8                | 46.5        | 38.9        | 37.4        | 26.5        | 26.7        |
| Qwen2.5-VL [5]     | ZS     | 7B         | 83.4               | 87.1        | 62.5        | 68.7                | 59.7        | 47.1        | 55.6        | 37.8        | 34.4        |
| SeeClick [7]       | SFT    | 9.6B       | 93.0               | 73.4        | 75.0        | 82.9                | 62.9        | 59.1        | <b>71.0</b> | 52.4        | <b>53.9</b> |
| Aria-UI [41]       | SFT    | 3.9B       | -                  | 87.7        | 67.3        | -                   | 43.2        | 10.2        | -           | <b>86.8</b> | 36.5        |
| Aguvis [24]        | SFT    | 7B         | -                  | -           | 80.5        | -                   | -           | 61.5        | -           | -           | -           |
| OS-Atlas-Pro [23]  | SFT    | 4B         | 91.9               | 83.8        | 80.6        | 84.7                | 73.8        | 67.5        | 83.5        | 61.4        | 56.4        |
| OS-Atlas-Pro [23]  | SFT    | 7B         | 93.6               | 88.0        | 85.2        | <u>85.2</u>         | <u>78.5</u> | <u>71.2</u> | 84.5        | 67.8        | 62.0        |
| UI-TARS [9]        | SFT    | 2B         | <b>98.1</b>        | 87.3        | <u>89.3</u> | <u>81.2</u>         | 78.4        | 68.9        | <u>93.9</u> | <u>86.8</u> | <u>83.4</u> |
| UI-TARS [9]        | SFT    | 7B         | <u>98.0</u>        | <b>89.3</b> | <b>90.8</b> | 83.7                | <b>80.5</b> | <b>72.5</b> | <b>94.6</b> | <b>90.1</b> | <b>87.0</b> |
| UI-R1 [30]         | RFT    | 3B         | 79.2               | 82.4        | 66.4        | 57.9                | 55.7        | 45.4        | 52.2        | 34.5        | 32.5        |
| GUI-R1 [10]        | RFT    | 3B         | 83.7               | 81.6        | 64.4        | 58.0                | 56.2        | 46.6        | 54.8        | 41.5        | 41.3        |
| GUI-R1 [10]        | RFT    | 7B         | 85.2               | 84.0        | 66.5        | 71.6                | 65.6        | 51.7        | 65.5        | 43.6        | 38.8        |
| <b>BTL-UI</b>      | RFT    | 3B         | 95.6               | 86.1        | 84.8        | 84.0                | 71.4        | 63.4        | 84.4        | 77.2        | 64.0        |
| <b>BTL-UI</b>      | RFT    | 7B         | 96.8               | <u>88.5</u> | 88.0        | <b>88.2</b>         | 76.9        | 69.2        | 84.6        | 78.4        | 65.4        |

### 4.3 Ablation Study

As shown in Table 6, to clarify the contributions of each component in our BTL framework, we conduct an ablation study on the AndroidControl-High benchmark. When trained only with SFT, BTL-UI achieves a baseline performance with an SR of 60.6%. While further using the generated Blink data, SFT obtains a 5% improvement. This proves that Blink data is not only suitable for RFT,

Table 6: Ablation study of BTL-UI. All ablation experiments are evaluated on the AndroidControl-High benchmark by evaluating the grounding and planning capabilities of the agent.

(a) Ablation study of training method and BTL. Blink Data refers to the data contribution in §3.2. BTL Reward denotes the reward design in §3.3.

| SFT | RFT | Blink Data | BTL Reward | AndroidControl-High |             |             |
|-----|-----|------------|------------|---------------------|-------------|-------------|
|     |     |            |            | Type                | GR          | SR          |
| -   | -   | -          | -          | 68.7                | 59.7        | 47.1        |
| ✓   | -   | -          | -          | 79.4                | 63.9        | 60.6        |
| ✓   | -   | ✓          | -          | 86.4                | 69.9        | 65.6        |
| -   | ✓   | -          | -          | 86.2                | 71.3        | 65.4        |
| -   | ✓   | ✓          | ✓          | <b>88.2</b>         | <b>76.9</b> | <b>69.2</b> |

(b) Ablation study of Blink Phase ROIs.

| $\lambda$ | AndroidControl-High |             |             |
|-----------|---------------------|-------------|-------------|
|           | Type                | GR          | SR          |
| 1         | 87.0                | 72.1        | 66.6        |
| 2         | 87.6                | 72.8        | 67.4        |
| 3         | 88.0                | 74.2        | 68.1        |
| 4         | 86.8                | 75.6        | 68.4        |
| 5         | 88.2                | <b>76.9</b> | <b>69.2</b> |
| 6         | <b>89.4</b>         | 73.1        | <b>69.2</b> |

but also for SFT. Furthermore, RFT without Blink data achieves an SR of 65.6%. After adopting Blink data and BTL reward, SR is improved by 3.6%.

Moreover, we examine the effect of varying the number of Blink ROIs ( $\lambda$ ): increasing  $\lambda$  from 1 to 6 steadily improves success rates from 66.6% to 69.2%, after which gains plateau, suggesting an optimal trade-off between annotation complexity and attention coverage. It is observed that from Table 6, as  $\lambda$  increases, the performance is saturated, so the final  $\lambda$  is selected as 5.

Overall, the ablation results confirm that each element of the BTL framework—Blink Phase for targeted attention, Think Phase for structured reasoning, and the Link Phase for precise validation—plays a crucial role in achieving competitive performance in GUI interaction tasks.

#### 4.4 Visualization

We present the visualization results and qualitative analysis in the appendix.

### 5 Conclusion and Limitations

We propose the BTL framework, an innovative GUI interaction architecture inspired by the biological cognitive paradigm of Blink–Think–Link. This framework simulates the human closed-loop system of visual perception, cognitive decision-making, and action execution during GUI operations, overcoming the limitations of traditional outcome-driven RFT approaches. Experimental results show that the BTL-UI agent, developed under this framework, achieves significant performance improvements across a variety of GUI interaction tasks.

We believe that the BTL framework proposed in this study establishes a promising and generalizable paradigm for developing digital assistants that are more natural, efficient, and aligned with human cognition. It not only benefits human-GUI interaction but can also be extended to other human-computer interaction tasks, such as embodied intelligence.

**Limitations.** The proposed BTL framework introduces *<blink>* tag outputs compared to conventional Think-Answer structured outputs. Although the blink-generated ROI regions are adaptive and can be empty (zero-length), they typically increase the output sequence length in most cases. While demonstrating performance improvements across various GUI task metrics, this design incurs additional computational processing overhead.

### References

- [1] Shuai Wang, Weiwen Liu, Jingxuan Chen, Yuqi Zhou, Weinan Gan, Xingshan Zeng, Yuhao Che, Shuai Yu, Xinlong Hao, Kun Shao, et al. Gui agents with foundation models: A comprehensive survey. *arXiv preprint arXiv:2411.04890*, 2024.
- [2] Xueyu Hu, Tao Xiong, Biao Yi, Zishu Wei, Ruixuan Xiao, Yurun Chen, Jiasheng Ye, Meiling Tao, Xiangxin Zhou, Ziyu Zhao, et al. Os agents: A survey on mllm-based agents for general computing devices use, 2024.
- [3] Dang Nguyen, Jian Chen, Yu Wang, Gang Wu, Namyong Park, Zhengmian Hu, Hanjia Lyu, Junda Wu, Ryan Aponte, Yu Xia, et al. Gui agents: A survey. *arXiv preprint arXiv:2412.13501*, 2024.
- [4] Peng Wang, Shuai Bai, Sinan Tan, Shijie Wang, Zhihao Fan, Jinze Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, et al. Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution. *arXiv preprint arXiv:2409.12191*, 2024.
- [5] Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, et al. Qwen2.5-vl technical report. *arXiv preprint arXiv:2502.13923*, 2025.
- [6] Wenyi Hong, Weihao Wang, Qingsong Lv, Jiazheng Xu, Wenmeng Yu, Junhui Ji, Yan Wang, Zihan Wang, Yuxiao Dong, Ming Ding, et al. Cogagent: A visual language model for gui agents. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14281–14290, 2024.

- [7] Kanzhi Cheng, Qiushi Sun, Yougang Chu, Fangzhi Xu, Yantao Li, Jianbing Zhang, and Zhiyong Wu. SeeClick: Harnessing gui grounding for advanced visual gui agents. *arXiv preprint arXiv:2401.10935*, 2024.
- [8] Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- [9] Yujia Qin, Yining Ye, Junjie Fang, Haoming Wang, Shihao Liang, Shizuo Tian, Junda Zhang, Jiahao Li, Yunxin Li, Shijue Huang, et al. Ui-tars: Pioneering automated gui interaction with native agents. *arXiv preprint arXiv:2501.12326*, 2025.
- [10] Xiaobo Xia and Run Luo. Gui-r1: A generalist r1-style vision-language action model for gui agents. *arXiv preprint arXiv:2504.10458*, 2025.
- [11] Yuhang Liu, Pengxiang Li, Congkai Xie, Xavier Hu, Xiaotian Han, Shengyu Zhang, Hongxia Yang, and Fei Wu. Infigui-r1: Advancing multimodal gui agents from reactive actors to deliberative reasoners. *arXiv preprint arXiv:2504.14239*, 2025.
- [12] Simon P Liversedge and John M Findlay. Saccadic eye movements and cognition. *Trends in cognitive sciences*, 4(1):6–14, 2000.
- [13] Alejandro Jaimes and Nicu Sebe. Multimodal human–computer interaction: A survey. *Computer vision and image understanding*, 108(1-2):116–134, 2007.
- [14] Robert JK Jacob. The use of eye movements in human-computer interaction techniques: what you look at is what you get. *ACM Transactions on Information Systems (TOIS)*, 9(2):152–169, 1991.
- [15] Songqin Nong, Jiali Zhu, Rui Wu, Jiongchao Jin, Shuo Shan, Xiutian Huang, and Wenhao Xu. Mobileflow: A multimodal llm for mobile gui agent. *arXiv preprint arXiv:2407.04346*, 2024.
- [16] Yunpeng Song, Yiheng Bian, Yongtao Tang, Guiyu Ma, and Zhongmin Cai. Visiontasker: Mobile task automation using vision based ui understanding and llm task planning. In *Proceedings of the 37th Annual ACM Symposium on User Interface Software and Technology*, pages 1–17, 2024.
- [17] Xiao Liu, Bo Qin, Dongzhu Liang, Guang Dong, Hanyu Lai, Hanchen Zhang, Hanlin Zhao, Iat Long Iong, Jiadai Sun, Jiaqi Wang, et al. Autoglm: Autonomous foundation agents for guis. *arXiv preprint arXiv:2411.00820*, 2024.
- [18] Huawen Shen, Chang Liu, Gengluo Li, Xinlong Wang, Yu Zhou, Can Ma, and Xiangyang Ji. Falcon-ui: Understanding gui before following user instructions. *arXiv preprint arXiv:2412.09362*, 2024.
- [19] Fei Tang, Yongliang Shen, Hang Zhang, Siqi Chen, Guiyang Hou, Wenqi Zhang, Wenqiao Zhang, Kaitao Song, Weiming Lu, and Yueting Zhuang. Think twice, click once: Enhancing gui grounding via fast and slow systems. *arXiv preprint arXiv:2503.06470*, 2025.
- [20] Filippou Christianos, Georgios Papoudakis, Thomas Coste, Jianye Hao, Jun Wang, and Kun Shao. Lightweight neural app control. *arXiv preprint arXiv:2410.17883*, 2024.
- [21] Jiani Zheng, Lu Wang, Fangkai Yang, Chaoyun Zhang, Lingrui Mei, Wenjie Yin, Qingwei Lin, Dongmei Zhang, Saravan Rajmohan, and Qi Zhang. Vem: Environment-free exploration for training gui agent with value environment model. *arXiv preprint arXiv:2502.18906*, 2025.
- [22] Boyu Gou, Ruohan Wang, Boyuan Zheng, Yanan Xie, Cheng Chang, Yiheng Shu, Huan Sun, and Yu Su. Navigating the digital world as humans do: Universal visual grounding for gui agents. *arXiv preprint arXiv:2410.05243*, 2024.
- [23] Zhiyong Wu, Zhenyu Wu, Fangzhi Xu, Yian Wang, Qiushi Sun, Chengyou Jia, Kanzhi Cheng, Zichen Ding, Liheng Chen, Paul Pu Liang, et al. Os-atlas: A foundation action model for generalist gui agents. In *International Conference on Learning Representations*, 2025.

- [24] Yiheng Xu, Zekun Wang, Junli Wang, Dunjie Lu, Tianbao Xie, Amrita Saha, Doyen Sahoo, Tao Yu, and Caiming Xiong. Aguis: Unified pure vision agents for autonomous gui interaction. *arXiv preprint arXiv:2412.04454*, 2024.
- [25] Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, et al. Openai o1 system card. *arXiv preprint arXiv:2412.16720*, 2024.
- [26] Ziyu Liu, Zeyi Sun, Yuhang Zang, Xiaoyi Dong, Yuhang Cao, Haodong Duan, Dahua Lin, and Jiaqi Wang. Visual-rft: Visual reinforcement fine-tuning. *arXiv preprint arXiv:2503.01785*, 2025.
- [27] Yufei Zhan, Yousong Zhu, Shurong Zheng, Hongyin Zhao, Fan Yang, Ming Tang, and Jinqiao Wang. Vision-r1: Evolving human-free alignment in large vision-language models via vision-guided reinforcement learning. *arXiv preprint arXiv:2503.18013*, 2025.
- [28] Huajie Tan, Yuheng Ji, Xiaoshuai Hao, Minglan Lin, Pengwei Wang, Zhongyuan Wang, and Shanghang Zhang. Reason-rft: Reinforcement fine-tuning for visual reasoning. *arXiv preprint arXiv:2503.20752*, 2025.
- [29] Haozhan Shen, Peng Liu, Jingcheng Li, Chunxin Fang, Yibo Ma, Jiajia Liao, Qiaoli Shen, Zilun Zhang, Kangjia Zhao, Qianqian Zhang, et al. Vlm-r1: A stable and generalizable r1-style large vision-language model. *arXiv preprint arXiv:2504.07615*, 2025.
- [30] Zhengxi Lu, Yuxiang Chai, Yaxuan Guo, Xi Yin, Liang Liu, Hao Wang, Guanjing Xiong, and Hongsheng Li. Ui-r1: Enhancing action prediction of gui agents by reinforcement learning. *arXiv preprint arXiv:2503.21620*, 2025.
- [31] Yadong Lu, Jianwei Yang, Yelong Shen, and Ahmed Awadallah. Omniparser for pure vision based gui agent, 2024.
- [32] Harold W Kuhn. The hungarian method for the assignment problem. *Naval research logistics quarterly*, 2(1-2):83–97, 1955.
- [33] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- [34] Yuze Zhao, Jintao Huang, Jinghan Hu, Xingjun Wang, Yunlin Mao, Daoze Zhang, Zeyinzi Jiang, Zhikai Wu, Baole Ai, Ang Wang, Wenmeng Zhou, and Yingda Chen. Swift:a scalable lightweight infrastructure for fine-tuning, 2024.
- [35] Kevin Qinghong Lin, Linjie Li, Difei Gao, Zhengyuan Yang, Zechen Bai, Weixian Lei, Lijuan Wang, and Mike Zheng Shou. Showui: One vision-language-action model for generalist gui agent. In *NeurIPS 2024 Workshop on Open-World Agents*, 2024.
- [36] Wei Li, William Bishop, Alice Li, Chris Rawles, Folawiyo Campbell-Ajala, Divya Tyamagundlu, and Oriana Riva. On the effects of data scale on computer control agents. *arXiv e-prints*, pages arXiv–2406, 2024.
- [37] Quanfeng Lu, Wenqi Shao, Zitao Liu, Fanqing Meng, Boxuan Li, Botong Chen, Siyuan Huang, Kaipeng Zhang, Yu Qiao, and Ping Luo. Gui odyssey: A comprehensive dataset for cross-app gui navigation on mobile devices. *arXiv preprint arXiv:2406.08451*, 2024.
- [38] Kaixin Li, Ziyang Meng, Hongzhan Lin, Ziyang Luo, Yuchen Tian, Jing Ma, Zhiyong Huang, and Tat-Seng Chua. Screenspot-pro: Gui grounding for professional high-resolution computer use. *arXiv preprint arXiv:2504.07981*, 2025.
- [39] OpenAI. Gpt-4o, 2024. Accessed: 2025-01-03.
- [40] Jinguo Zhu, Weiyun Wang, Zhe Chen, Zhaoyang Liu, Shenglong Ye, Lixin Gu, Yuchen Duan, Hao Tian, Weijie Su, Jie Shao, et al. Internvl3: Exploring advanced training and test-time recipes for open-source multimodal models. *arXiv preprint arXiv:2504.10479*, 2025.
- [41] Yuhao Yang, Yue Wang, Dongxu Li, Ziyang Luo, Bei Chen, Chao Huang, and Junnan Li. Aria-ui: Visual grounding for gui instructions. *arXiv preprint arXiv:2412.16256*, 2024.

## A Prompt

| System Prompt                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p>You are a GUI Agent capable of reasoning based on user instructions, action history, and the current screenshot. You should first observe the layout of the screenshot and extract <math>N</math> elements RELATED TO the user instruction, where <math>0 \leq N \leq 5</math>. Next, think about the reasoning process BASED ON the observations and instructions in your mind, and then provide the user with the answer.</p> <p>The observation process (can be None if <math>N == 0</math>), reasoning process and answer are enclosed within <code>&lt;blink&gt;&lt;/blink&gt;</code>, <code>&lt;think&gt;&lt;/think&gt;</code> and <code>&lt;link&gt;&lt;/link&gt;</code> tags, respectively, i.e.,</p> <pre> &lt;blink&gt; &lt;element&gt;&lt;id&gt;1&lt;/id&gt;&lt;bbox&gt;[x0, y0, x1, y1]&lt;/bbox&gt;&lt;caption&gt;dynamic&lt;/caption&gt;&lt;/element&gt; &lt;element&gt;&lt;id&gt;2&lt;/id&gt;&lt;bbox&gt;[x2, y2, x3, y3]&lt;/bbox&gt;&lt;caption&gt;static&lt;/caption&gt;&lt;/element&gt; &lt;element&gt;&lt;id&gt;3&lt;/id&gt;.....&lt;/element&gt; &lt;element&gt;&lt;id&gt;4&lt;/id&gt;.....&lt;/element&gt; &lt;/blink&gt; &lt;think&gt; reasoning process here &lt;/think&gt; &lt;link&gt; answer(["Plan": ..., "Action": "function": ..., ...]) &lt;/link&gt;. where captions must be one of [dynamic, static], "dynamic" refers to the interactive area, and "static" refers to the non-interactive areas, such as text and diagrams in the screenshot. And the observation can be &lt;blink&gt; None &lt;/blink&gt;, if <math>N == 0</math>. </pre> |
| User Instruction Prompt                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
| <p>You are given a task and your action history, with screenshots. You need to perform the next action to complete the task. You MUST CHOOSE the next action from the following defined action space.</p> <p><b>## Action Space</b></p> <p>Action 1: Back</p> <ul style="list-style-type: none"> <li>- format: {'function': 'Back'}</li> <li>- purpose: back to the previous screen.</li> </ul> <p>Action 2: Home</p> <ul style="list-style-type: none"> <li>- format: {'function': 'Home'}</li> <li>- purpose: navigate to the home page.</li> </ul> <p>Action 3: Tap</p> <ul style="list-style-type: none"> <li>- format: {'function': 'Tap', 'position': [x, y]}</li> <li>- purpose: tap the specified position.</li> </ul> <p>Action 4: Type</p> <ul style="list-style-type: none"> <li>- format: {'function': 'Type', 'text': 'str'}</li> <li>- purpose: enter specified text at the designated location.</li> </ul> <p>Action 5: Swipe</p> <ul style="list-style-type: none"> <li>- format: {'function': 'Swipe', 'direction': 'str'}</li> <li>- purpose: swipe on the screen in the specified direction.</li> </ul> <p>Action 6: LongPress</p> <ul style="list-style-type: none"> <li>- format: {'function': 'LongPress', 'position': [x, y]}</li> <li>- purpose: long press the specified position</li> </ul> <p><b>## User Instruction</b></p> <p>High – Level Instruction</p> <p><b>## Action History</b></p> <p>Step 1: .....</p> <p>Step 2: .....</p> <p><b>## Screenshots</b></p> <pre> &lt;image&gt; </pre>                                                       |

Table 7: The prompt for the BTL-UI.

The prompt of BTL-UI is shown in Table 7. The system prompt is used to format the output of the model according to the three-phase paradigm of Blink-Think-Link. Moreover, the model outputs according to the format of the system prompt, which is convenient for the calculation of the BTL

reward to adjust the model distribution. As shown in equation 6, because the output of Blink Phase can be  $\emptyset$ , we emphasize that the Blink Phase can be output *None* in the system prompt.

OS-Atlas [23] has found that blindly mixing data from different sources for multitask fine-tuning can significantly harm performance due to action space conflicts. We unify the action space of GUI-Odyssey [37] and AndroidControl [36]. And we prompt the model to select the corresponding action from the defined action space. Furthermore, we declare the format and purpose in the user instruction prompt, so that the model can better understand each action type. For the grounding and high-level tasks, only the high-level instruction will be fed into the model. For the low-level tasks, both the high-level and the low-level instructions will be given to the model. In the end, the interaction history will be added to the user instruction prompt.

## B Visualization

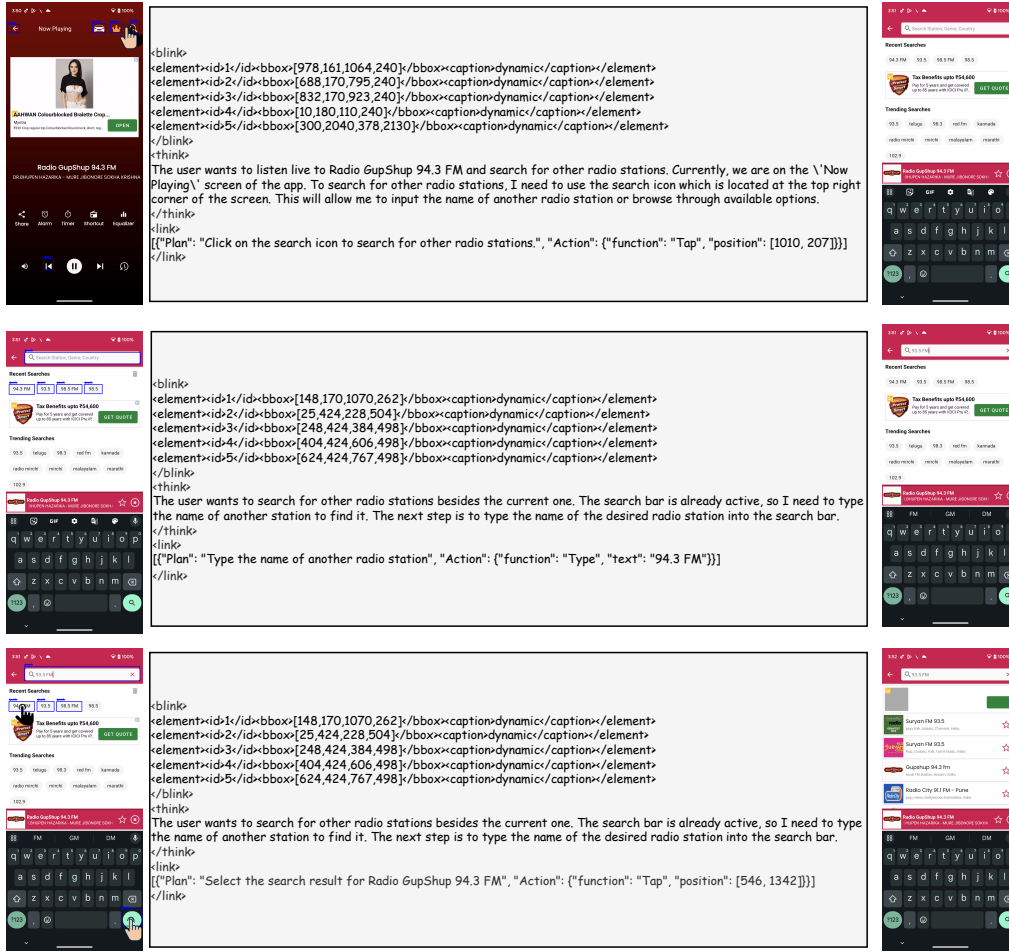


Figure 4: Visualization of the interaction trajectory of the proposed BTL-UI on AndroidControl-High. The corresponding ID of this random case is 19477. And the high-level instruction is 'Listen live to Radio GupShup 94.3 FM and search for other radio stations.' The tap icon in black is the prediction of BTL-UI, and the other is the ground-truth.

The visualization of the interaction trajectory of our BTL-UI is shown in Figure 4. The high-level instruction is "Listen live to Radio GupShup 94.3 FM and search for other radio stations." The Blink Phase can locate the ROIs related to the instruction. And the thinking Phase can reason based on the instruction, interaction history, and candidate area. As shown in step 2 of the interaction trajectory in



Figure 4, in the Blink Phase, BTL-UI not only locates the input box to complete the task, but also analyzes the historical search records in the screenshots.

However, since AndroidControl is an offline interaction benchmark, there are some unreasonable labeling data. For instance, step 2 needs to input the text of “94.3 FM” according to the task instruction. But the search box in the screenshot after interaction shows “93.5 FM”, which may affect subsequent interactions. In step 3, the labeled action is to click the search icon. And the search icon is also located in the Blink Phase. Due to the interaction errors in step 2 caused by data noise, the Think Phase of BTL-UI believes that clicking on the “94.3 FM” in historical search records in the screenshot is more reasonable. Therefore, we suppose our BTL-UI has stronger reasoning and error correction abilities.

## C Experiment Statistical Significance

In this section, we report the experiment’s statistical significance. The random factor that affects our results is the sampling of the training process. As shown in Table 1, the training data of our BTL-UI is sampled from various datasets. In the data sampling process, we fix the random seed to 2025 to maintain reproducibility. And the sampled data is further adopted to generate Blink Data, following the pipeline in §3.2. Moreover, BTL-UI adopts the ms-swift [34] framework for RL training. During the training process, we also fix the random seed to 2025 to maintain reproducibility.

## NeurIPS Paper Checklist

### 1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We have polished our abstract and introduction to accurately reflect our main contributions and scope.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

### 2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The limitations are discussed in Section 5.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

### 3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[NA\]](#)

Justification: The paper does not propose a theory and does not have any theoretical results.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

#### 4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: Experiments are described in detail, including the dataset we use, all the hyperparameters, and the model training details.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
  - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
  - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
  - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
  - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

#### 5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: The source code will be released upon acceptance with detailed instructions including data access and preparation, training and inference code, and the environment needed.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

## 6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: All the detailed information is provided in the experiment setup.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

## 7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: The random factor that affects our results is the sampling of the training process. We show quantitative error metrics in the supplemental material.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)

- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

## 8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We show the computer resources in experimental details (NVIDIA H100 GPUs) in experimental results.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

## 9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: All the researches in this paper conform to the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

## 10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [No]

Justification: Our goal is as simple as to improve agent comprehension. This is not a security, safety, or privacy-related research direction and is not related to any potential harmful or malicious usage in the future.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.

- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

#### 11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

#### 12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: All assets we used are properly cited.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, [paperswithcode.com/datasets](https://paperswithcode.com/datasets) has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.



- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

### 13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not release new assets. The codes and model checkpoints will be released upon acceptance.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

### 14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

### 15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.

- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

**16. Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: This paper just takes LLM for editing.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.