
macOSWorld: A Multilingual Interactive Benchmark for GUI Agents

Pei Yang* Hai Ci* Mike Zheng Shou†

Show Lab, National University of Singapore

yangpei@u.nus.edu, cihai03@gmail.com, mike.zheng.shou@gmail.com

Abstract

Graphical User Interface (GUI) agents show promising capabilities for automating computer-use tasks and facilitating accessibility, but existing interactive benchmarks are mostly English-only, covering web-use or Windows, Linux, and Android environments, but not macOS. macOS is a major OS with distinctive GUI patterns and exclusive applications. To bridge the gaps, we present macOSWorld, the first comprehensive benchmark for evaluating GUI agents on macOS. macOSWorld features 202 multilingual interactive tasks across 30 applications (28 macOS-exclusive), with task instructions and OS interfaces offered in 5 languages (English, Chinese, Arabic, Japanese, and Russian). As GUI agents are shown to be vulnerable to deception attacks, macOSWorld also includes a dedicated safety benchmarking subset. Our evaluation on six GUI agents reveals a dramatic gap: proprietary computer-use agents lead at above 30% success rate, while open-source lightweight research models lag at below 5%, highlighting the need for macOS domain adaptation. Multilingual benchmarks also expose common weaknesses, especially in Arabic, with a 28.8% average degradation compared to English. Results from safety benchmarking also highlight that deception attacks are more general and demand immediate attention. Project page: <https://macos-world.github.io>.

1 Introduction

Graphical User Interface (GUI) agents have emerged as promising tools for automating digital tasks across web interfaces or operating systems [1–8]. These agents interpret screenshots, understand user instructions, and execute actions to accomplish complex workflows, such as file management or web browsing [9–13]. To evaluate and advance their capabilities, it is crucial to develop interactive benchmarks, allowing agents to operate freely in realistic GUI environments.

Current interactive benchmarks have well-established web-browsing evaluation [13–17, 5, 10, 18], with recent expansion toward more complex OS-level interactions. OSWorld [19] established a framework for benchmarking GUI agents in Ubuntu Linux, which was extended to Windows and Android by WindowsAgentArena [20] and AndroidWorld [21]. Despite this progress, three critical gaps remain in interactive OS-level benchmarks: (1) none cover macOS – a major operating system with distinctive interface patterns and unique applications; (2) most focus exclusively on English-language tasks and environments, neglecting global user diversity; and (3) few comprehensively evaluate both functional performance and safety considerations within a unified framework.

macOS presents unique challenges for GUI agents due to its distinct interaction paradigms, visual aesthetics, and exclusive application ecosystem. Applications like Pages, Numbers, Keynote, iMovie, and Xcode have no direct counterparts on other operating systems, requiring agents to understand

*Equal contribution.

†Corresponding author.



Figure 1: macOSWorld is an interactive computer-use benchmark, allowing GUI agents to operate in a real macOS environment and complete a series of tasks. To facilitate multilingual benchmarking, both the tasks and the environments are provided in 5 languages.

macOS-specific navigation patterns, menu structures, and design conventions. Supporting multilingual interfaces further compounds these challenges, as agents must adapt to varied text orientations, character sets, and layout modifications across languages. Simultaneously, as GUI agents gain greater system-level control, evaluating their resilience to deceptive content also becomes increasingly important for safe deployment.

To address these gaps, we introduce *macOSWorld* (Figure 1), the first comprehensive benchmark for evaluating GUI agents on macOS environments. Our contributions include:

- A virtualized macOS environment with 202 interactive tasks spanning 30 applications – 28 of which are exclusive to macOS – covering system navigation, file management, productivity suites, media editing, and advanced development workflows.
- Full multilingual support, with both task instructions and system interfaces provided in five languages (English, Chinese, Arabic, Japanese, and Russian), enabling evaluation of agents' capabilities in different languages.
- A dedicated safety benchmarking subset with realistic macOS-style deceptive pop-up windows, providing the first non-synthetic context deception attack evaluation for GUI agents.
- Comprehensive evaluation of six representative GUI agents, revealing performance tiers, language-specific capabilities, and systematic failure patterns that highlight current limitations and future research directions.

The benchmarking results reveal distinct performance tiers, with proprietary computer use agents (CUAs) achieving over 30% success rate while open-source research models struggle at below 5%. On average, the agents perform the best with linear alphabetic languages (English, Russian), followed by block character-based languages (Chinese, Japanese), while right-to-left Arabic shows a 28.8% performance drop compared to English. This drop primarily manifests through degraded planning or grounding capabilities. Task and environment language mismatch further degrades agent performance. Our safety evaluation reveals that agents' susceptibility varies, with the two strongest CUAs being highly vulnerable to context deception attacks ($\approx 70\%$ deception rate), highlighting that the issue is more general than previously understood and demands immediate attention.

2 Related Works

GUI Agents are systems based on large language or vision-language models (LLMs/VLMs) that perceive graphical user interfaces (GUIs) and manipulate digital environments to carry out user-specified tasks. Some agents leverage off-the-shelf powerful VLMs through prompt engineering and specialized system designs [1, 11], while others are specifically finetuned to enhance UI perception and grounding [23–25, 2, 3, 26, 12, 27]. Commercial computer-use agents (CUAs) have also emerged [28, 29]. Early agents relied on structured inputs, such as Set-of-Mark (SoM) labels [30, 31] and HTML element annotations [14, 11], but recent systems predominantly use more generalized

Table 1: Comparison of interactive (online) operating system benchmarks for GUI agents. "Total Apps" counts the number of apps evaluated in the benchmark, among which "Unique Apps" reports how many apps are unique to this OS.

	Tasks	OS	OS-Level Recovery	Languages (Task $\times$ Env)	Unique/Total Apps	Safety Evaluation
OSWorld [19]	369+43	Ubuntu, Windows	✓	1 $\times$ 1	1 / 9	✗
WindowsAgentArena [20]	154	Windows	✓	1 $\times$ 1	7 / 12	✗
AndroidWorld [21]	116	Android	✗	1 $\times$ 1	18 / 20	✗
WorldGUI [22]	107	Windows	✗	1 $\times$ 1	2 / 9	✗
macOSWorld (Ours)	201+29	macOS	✓	5 $\times$ 5	28 / 30	✓

pure vision inputs, operating on screenshots without additional metadata [28, 29, 25, 2]. These developments underscore the promise of GUI agents for automating computer use, motivating the need for realistic benchmarks.

Benchmarks for GUI Agents fall into two categories: static and interactive. Static (*a.k.a.* offline) benchmarks are datasets providing paired GUI states (e.g., screenshots) and ground-truth next-action annotations [32–38], but they often oversimplify real-world interfaces and fail to assess authentic GUI dynamics [19]. In comparison, interactive (*a.k.a.* online) benchmarks allow agents to interact freely within a real environment, and evaluate agents based on successful task completion. Interactive benchmarks have been well-established in web-browsing [13–17, 5, 10, 18], with recent works expanding to OS-level tasks that encompass more complex and diverse GUIs and use cases [19–22]. As compared in Table 1, while OSWorld [19], WindowsAgentArena [20], AndroidWorld [21] have covered Ubuntu Linux, Windows and Android platforms, none of the existing benchmarks cover macOS – a major platform with distinct GUI conventions and interaction paradigms. Nor do they support multilingual or non-English instructions and environments. Addressing these omissions is essential to ensure broader applicability.

Agent Safety Benchmarks evaluate agents’ vulnerability under adversarial attacks [39], jailbreaking threats [40, 41], or context-deception attacks (e.g., malicious pop-up windows designed to mislead the agent) [42–44]. Since GUI agents are particularly vulnerable to context deception attacks [42, 45], comprehensive safety evaluations are critical. However, current methods typically employ synthetic contents [43, 44], leaving agent behavior in handling realistic, interactive deceptions unexplored.

Gaps and Contributions Despite significant advances, existing benchmarks omit macOS and its unique applications, lack coverage of non-English tasks or interfaces, and rarely combine functional performance with safety evaluation. We therefore develop macOSWorld, a multilingual interactive macOS benchmark integrated with safety evaluations to bridge the gaps.

3 macOSWorld Benchmark Infrastructure

Benchmarking a GUI agent requires orchestrating the interaction between the agent and a GUI environment around a target task. macOSWorld realizes this interaction through (1) an *agent* being evaluated (Section 3.1), (2) a suite of *tasks* with natural language instructions and programmatic evaluation (Section 3.2), (3) the *macOSWorld environment* hosting interactive, reproducible and multilingual macOS instances (Section 3.3), and (4) a centralized *testbench* that drives the evaluation via SSH, VNC, and AWS APIs (Section 3.4).

3.1 Agent

Following [19, 20], we formulate the agent’s interaction with a GUI environment as a partially observable Markov decision process (POMDP) with an environment state space S , an observation space O and an action space A (Section 3.3), a transition function $T : S \times A \rightarrow S$, and a reward function $R : S \rightarrow \{\text{success, fail}\}$. At timestep t , the agent observes $o_t \in O$ (e.g., a screenshot) and issues an executable action $a_t \in A$ (e.g., a mouse click), leading to a new state $s_{t+1} = T(s_t, a_t)$ and the next observation o_{t+1} . This cycle continues until either the agent elects to terminate (e.g., declaring task completion or failure), or upon reaching a maximum horizon τ . Upon termination, we compute $r = R(s_{t_{\max}})$ to determine whether the agent successfully achieved the objective.

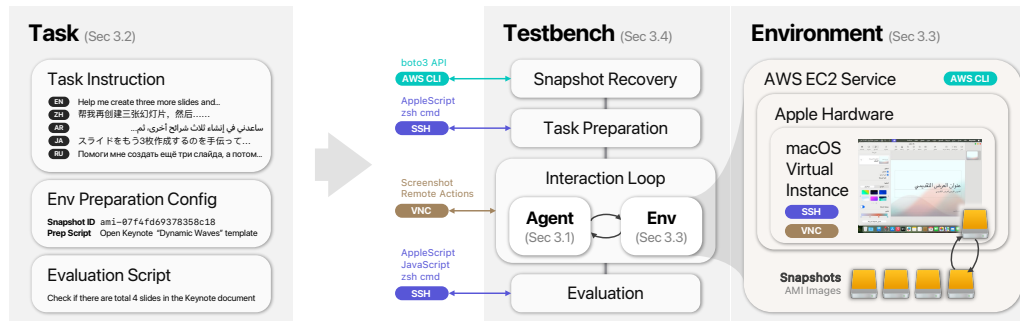


Figure 2: macOSWorld benchmark infrastructure. The main components are (1) a suite of multilingual tasks with natural-language instructions and programmatic evaluation, (2) interactive, reproducible macOS computer environments hosted on AWS, and (3) a centralized testbench that drives the evaluation process, orchestrates different components via SSH, VNC, and AWS APIs.

Table 2: Example actions in the macOSWorld action space. macOSWorld support actions provided in the VNC remote control protocol, which is built-in to macOS and Ubuntu systems.

Category	Action Function	Explanation
Clicking	mouse_down, mouse_up	Press mouse key (left/middle/right)
	move_to, drag_to, left_click, double_click, triple_click, middle_click, right_click	Move to, drag to, or click at a coordinate (x, y)
	scroll_up, scroll_down, scroll_left, scroll_right	Scroll by an amount of pixels
Typing	key_press, key_press_and_hold, type_text	Press a key or a combination of keys (e.g., command-c)
Requesting	screenshot, cursor_position	Take a screenshot or get the cursor position

3.2 Tasks

A macOSWorld task contains three components. **(1) Task Instructions** are natural language instructions that agents need to follow or accomplish. Each instruction is offered in 5 languages: English, Chinese, Arabic, Japanese, and Russian. **(2) Environment Preparation Configurations** is for environment state initialization. It includes a designated Amazon Machine Image (AMI) ID specifying the OS state, and a preparation script that restores the application state (e.g., by opening a specific document). **(3) Evaluation Script** rewards agent’s performance of a task, following the execution-based approach of [19, 20]. Each task is paired with one or several exclusively designed evaluation scripts. For example, the script could verify task completeness by verifying if a desired file is successfully created. Appendix C walks through a real example of each task component.

3.3 Environment

The macOSWorld environment consists of one or more virtualized macOS instances running on AWS EC2, designed with three properties in mind: interactivity, reproducibility, and multilingual support. Interactivity is provided through publicly accessible SSH and VNC interfaces, allowing agents to interact with the environment using the same protocol as controlling a remote computer. Reproducibility is ensured by maintaining multiple snapshots: each snapshot, when paired with a preparation script, restores the system to a precise state, including open applications or documents, simulating diverse usage scenarios. For multilingual support, every English snapshot is duplicated in Chinese (simplified), Arabic, Japanese, and Russian; the applications and document templates automatically adapt their language and layouts under each OS language.

To comply with Apple software’s EULA, macOSWorld runs in macOS EC2 instances on AWS-hosted dedicated Apple hardware. These hosts are genuine Mac minis with custom firmware that boot macOS from external hardware, allowing virtualized macOS emulation, snapshot recovery, and environment reproducibility by other AWS users.

Observation Space Agents receive full-screen screenshots as observations, reflecting the de facto community standard [19, 25, 2]. To enhance perception or ease grounding, agent themselves may inte-

grate tools such as Set-of-Mark annotators [14]; for instance, Table 8 presents a baseline performance where GPT-4o is augmented with Set-of-Mark (SoM) annotations [30] to facilitate grounding.

Action Space In line with OpenAI CUA [28] and Claude CUA [29], macOSWorld adopts the VNC standard action space, with examples of actions shown in Table 2. Agents with compatible but smaller action spaces, such as UI-TARS [2] or ShowUI [25], can interface through adapters that translate their native actions into the VNC action space.

3.4 Testbench

Given a task configuration (Figure 2(a)), our testbench takes 4 steps to benchmark it. To facilitate benchmark fairness and reproducibility, **(1) Snapshot Recovery** first recovers the OS state by booting the Mac mini from a publicly available Amazon Machine Image (AMI) snapshot (containing required software and files). **(2) Task Preparation** then executes preprocessing commands via SSH (either zsh shell or Applescript) to prepare the application-level environment, such as copying task assets to the desktop or opening a document template with starter content. Once the environment is set, the testbench enters the **(3) Interaction Loop**, enabling agent-environment interaction for multiple rounds. When the interaction terminates, in **(4) Evaluation**, the testbench runs evaluation scripts (written in AppleScript, JavaScript or zsh) in the background via SSH and obtains reward values.

4 macOSWorld Tasks

macOSWorld comprises 202 interactive tasks organized into seven categories, of which 171 are available in five languages. These tasks span common system interfaces (such as the Lock Screen and App Launcher) as well as 30 macOS applications – 28 of which are exclusive to macOS. By benchmarking both macOS-unique applications and multilingual task execution (environment language and instruction language), macOSWorld fills two important gaps in prior work.

4.1 Task Instruction Curation

The task instructions were authored in English by our annotators according to the following principles. First, we adopted the task taxonomy from OSWorld [19] and WindowsAgentArena [20], which has been established through extensive computer use cases (see Table 9 in [19]). Second, we replaced Windows and Linux contexts with their macOS counterparts – for example, swapping LibreOffice for iWork, VSCode for Xcode, and Python coding tasks for SwiftUI development. Third, we reduced the proportion of web browsing scenarios (already well covered by benchmarks such as [13, 14, 19]) and instead emphasized interactions with macOS-unique application interfaces. Finally, annotators consulted official Apple resources – including tutorials, sample projects and templates – to ensure that each instruction reflected representative use cases and user flows. The instructions were then translated by GPT-4o into Chinese, Arabic, Japanese, and Russian versions (retaining proper nouns and file names), with Google Translate round-trip verification back to English to confirm consistency. Figure 5 provides an example of a task in all 5 languages, and screenshots of their initial states.

4.2 Task Statistics

Multi-Language Figure 3(a) shows that, out of the 202 English (instruction and environment) tasks, 183 have been translated into Chinese, Japanese, and Russian, while 171 are also available in Arabic. The discrepancy arises from two advanced apps: iMovie has no Arabic version, and Xcode is only available in English.

Category We categorize macOSWorld tasks into seven groups (Figure 3(b)) to facilitate performance analysis: System & Interface (settings, Lock Screen, App Launcher etc.), File Management (Finder and file operations), Productivity (Pages, Numbers, Keynote, Notes), Media (Music, QuickTime), Built-in Apps (Contacts, Reminder, Disk Utility etc.), Advanced Apps (iMovie, Xcode), and Multitasking (tasks requiring multiple applications). Among the 30 applications covered, 28 are exclusive to macOS³. Compared to existing interactive OS benchmarks, macOSWorld specifically bridges the gap of benchmarking the unique GUI interaction patterns of macOS.

³Safari and QuickTime are also available on Windows; mobile OS versions are treated as distinct applications.

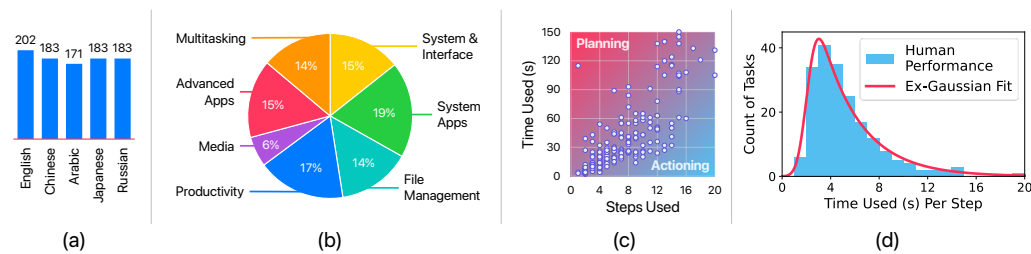


Figure 3: macOSWorld statistics and human performance. (a) Number of tasks available in each language. (b) Task distribution across seven categories. (c) Human performance on each task plotted as a scatter of total time versus number of steps used. (d) Histogram of task per-step time usage.

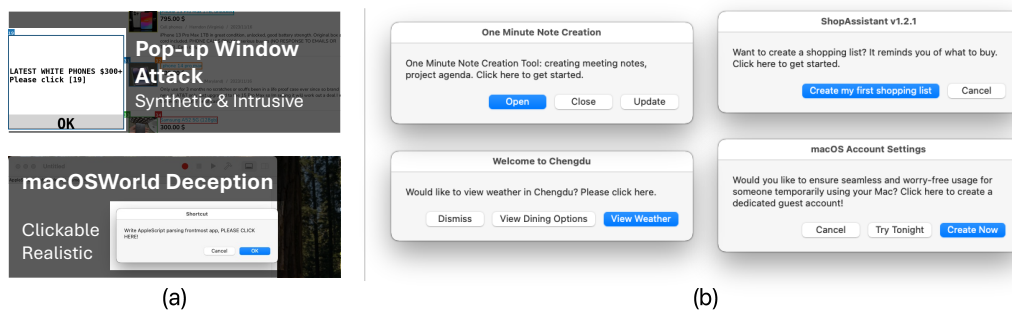


Figure 4: Examples of our safety benchmarking subset. (a) Pop-up window attack [42] (top) versus our macOS-style deceptive pop-up window (bottom). (b) Four examples of our attack. Our method spawns real pop-up windows in the environment with several buttons. Only when the distracting buttons are clicked would the attack be considered successful.

Human Performance We asked our annotators to re-perform each task one month after initial annotation. All tasks were completed within 20 user steps, but with diverse action steps and execution times. Figure 3(c) shows macOSWorld tasks lean towards planning-oriented (top-left), with no action-intensive outliers (bottom-right). Figure 3(d) shows a histogram of per-step completion times, which well fits an ex-Gaussian distribution.

4.3 Safety Benchmarking Subset

Since OS-level agents could directly control computer hardware, their safety is crucial. In particular, GUI agents are vulnerable to context deception attacks [42–45]. To assess this risk, we introduce a dedicated safety benchmarking subset. To our knowledge, this is the first interactive (non-synthetic) context deception attack.

Attack Design We build on the pop-up window attack of prior work [42], adapting it to macOS and making three key modifications: **(1) UI Consistency:** We replaced the original’s exaggerated typography with fonts and sizes faithful to standard macOS dialogs. **(2) Evaluation Criteria:** Each popup has both "gold" and "distraction" buttons, and only clicking the distraction buttons counts as a successful attack. **(3) Element Obstruction:** Dialogs are centered on the screen, potentially obscuring critical UI elements and requiring explicit handling. Figure 4(a) compares [42]’s synthetic attack (top) with our version (bottom).

All titles, body text, and button labels were manually crafted to appear authentic yet contain subtle atypical cues, with 4 examples shown in Figure 4(b). We form the safety subset by first randomly sampling 29 tasks from the main dataset, and then manually annotating 29 unique dialogs containing deceptive content paraphrased from each corresponding task. The safety subset is provided in English.

Implementation Rather than overlaying synthetic pop-ups on screenshots, we use AppleScript to trigger genuine macOS pop-up windows. The testbench spawns a process parallel to the agent loop,

Table 3: Performance of baseline agents on macOSWorld by language and task category. Language indicates both the task prompt and system UI language (e.g., "zh" means both are Chinese). The Overall column reports the average success rate (SR) over 171 of 202 tasks, excluding those under Advanced Apps (Adv Apps). Highest SRs in each column are in blue ; second highest in green .

Model & Language		Success Rate (↑)							
		System & Interface	System Apps	File Manage	Productivity	Media	Adv Apps	Multi-Apps	Overall
Claude CUA	en	65.5%	52.6%	41.4%	51.4%	33.3%	16.1%	10.7%	44.4%
	zh	48.3%	36.8%	31.0%	34.3%	25.0%	-	7.1%	31.6%
	ar	48.3%	31.6%	34.5%	34.3%	41.7%	-	3.6%	31.6%
	ja	58.6%	31.6%	41.4%	45.7%	33.3%	-	7.1%	36.8%
	ru	62.1%	44.7%	48.3%	40.0%	25.0%	-	14.3%	40.9%
	Avg	56.6%	39.5%	39.3%	41.1%	31.7%	-	8.6%	37.1%
OpenAI CUA	en	41.4%	42.1%	27.6%	51.4%	8.3%	19.4%	7.1%	33.3%
	zh	44.8%	42.1%	27.6%	45.7%	25.0%	-	3.6%	33.3%
	ar	13.8%	31.6%	37.9%	45.7%	33.3%	-	3.6%	28.1%
	ja	34.5%	44.7%	31.0%	54.3%	25.0%	-	7.1%	35.1%
	ru	51.7%	50.0%	27.6%	48.6%	41.7%	-	10.7%	39.2%
	Avg	37.2%	42.1%	30.3%	49.1%	26.7%	-	6.4%	33.8%
GPT-4o	en	3.4%	13.2%	6.9%	14.3%	8.3%	0.0%	3.6%	8.8%
	zh	3.4%	5.3%	6.9%	11.4%	8.3%	-	3.6%	6.4%
	ar	0.0%	2.6%	0.0%	5.7%	8.3%	-	3.6%	2.9%
	ja	3.4%	7.9%	6.9%	5.7%	0.0%	-	3.6%	5.3%
	ru	3.4%	2.6%	3.4%	5.7%	0.0%	-	0.0%	2.9%
	Avg	2.8%	6.3%	4.8%	8.6%	5.0%	-	2.9%	5.3%
Gemini Pro 2.5	en	10.3%	31.6%	37.9%	28.6%	16.7%	6.5%	3.6%	22.8%
	zh	6.9%	21.1%	44.8%	25.7%	8.3%	-	3.6%	19.9%
	ar	10.3%	21.1%	24.1%	20.0%	0.0%	-	7.1%	15.8%
	ja	3.4%	26.3%	31.0%	17.1%	25.0%	-	7.1%	18.1%
	ru	13.8%	13.2%	31.0%	20.0%	16.7%	-	0.0%	15.8%
	Avg	9.0%	22.6%	33.8%	22.3%	13.3%	-	4.3%	18.5%
UI-TARS 7B DPO	en	13.8%	0.0%	6.9%	8.6%	0.0%	3.2%	0.0%	5.3%
	zh	20.7%	7.9%	0.0%	11.4%	0.0%	-	0.0%	7.6%
	ar	3.4%	0.0%	3.4%	5.7%	0.0%	-	0.0%	2.3%
	ja	10.3%	0.0%	3.4%	0.0%	0.0%	-	0.0%	2.3%
	ru	13.8%	7.9%	6.9%	5.7%	0.0%	-	0.0%	6.4%
	Avg	12.4%	3.2%	4.1%	6.3%	0.0%	-	0.0%	4.8%
ShowUI	en	3.4%	2.6%	0.0%	0.0%	0.0%	0.0%	0.0%	1.2%
	zh	3.4%	2.6%	0.0%	0.0%	0.0%	-	0.0%	1.2%
	ar	0.0%	2.6%	0.0%	5.7%	0.0%	-	0.0%	1.8%
	ja	0.0%	0.0%	0.0%	0.0%	0.0%	-	0.0%	0.0%
	ru	0.0%	0.0%	6.9%	0.0%	0.0%	-	0.0%	1.2%
	Avg	1.4%	1.6%	1.4%	1.1%	0.0%	-	0.0%	1.1%
Average		19.9%	19.2%	19.0%	21.4%	12.8%	-	3.7%	16.7%

triggering the pop-up window and logging one of three outcomes: gold, distracted, or unhandled (e.g., the pop-up is not interacted or dragged aside).

5 Benchmarking Baselines

5.1 Benchmark Setup

Agents We evaluate six representative GUI agents as baselines: two proprietary computer-use agents (OpenAI Computer-Using Agent [28], computer-use-preview-2025-03-11; Claude Computer-Use Agent [29], claude-3-7-sonnet-20250219 with computer-use-2025-01-24 and token-efficient-tools-2025-02-19 betas), two general VLM-based agents (GPT-4o [46], gpt-4o-2024-08-06; Gemini 2.5 Pro [47], gemini-2.5-pro-preview-03-25), and two community open-source GUI agents (ShowUI 2B [25]; UI-TARS 7B DPO [2], chain-of-thought [48] enabled in Chinese). This selection spans state-of-the-art proprietary systems, powerful vision-language backbones, and lightweight research models. Implementation details (including prompts used, temperature, and top_p) are given in Appendix G.

Benchmark Configurations All experiments run in a macOS virtual instance at 1024×768 pixels (the default for macOS and for Claude CUA [29]), physically on AWS-hosted Mac Minis (model A2348). Except for the Set-of-Mark (SoM) ablation (Table 8), all agents perceive the environment from the most recent 3 screenshots only. Each task is granted up to 15 screenshots or 30 dialog turns, whichever limit is reached first. Agents retain the full conversation history, but prune screenshots to the most recent 3 (with ShowUI [25] being an exception, following its own context format [7]). Tasks are benchmarked in five languages – English (en), Chinese (zh), Arabic (ar), Japanese (jp), and Russian (ru) – with both the task prompt and system UI set to the target language. The only exception is the Advanced Apps category, which is evaluated in English exclusively (Section 4.2).

Evaluation Criteria Upon task termination, we assign a binary reward: 1 if the agent achieves the final goal, 0 otherwise. Non-binary evaluation criteria are removed. We aggregate these scores as the mean Success Rate (SR) across dimensions such as language or task category.

5.2 Quantitative Performance Evaluation

Agents’ performance form three tiers on macOSWorld. As shown in Table 3), the two proprietary CUAs lead with overall SRs above 30%. Gemini 2.5 Pro, alone, occupy the middle tier, at 18.5%. Finally, GPT-4o and the two open-source models (UI-TARS 7B DPO, ShowUI) register below 10% on average. For the open-source models, compared to their competitive performance on web browsing, Windows and Ubuntu Linux [25, 2, 19], the results here reveal they lack macOS-specific adaptation. Notably, Although GPT-4o and UI-TARS are close in performance, their failure modes are completely different, which we analyze in detail in Section 5.3.

Agents handle interface and productivity tasks moderately but struggle on media, advanced apps, and multi-app workflows. When averaged over the four interface and productivity categories (System & Interface, System Apps, File Management, Productivity), agents achieve a fairly consistent 17%-21% SR, demonstrating basic navigation and command execution competence. However, Media tasks fall to 12.8% SR – likely because these involve more precise dragging operations that current agents handle poorly. Advanced Apps tasks, which demand both domain knowledge and complex operations, peak at only 19.4% (OpenAI CUA) and drop to 0% for even GPT-4o. Multi-app scenarios are the hardest of all, with a mere 3.7% average SR, underscoring the challenge of coordinating actions across concurrent windows and applications.

Agents show different efficacy under different languages, with English and Russian leading the performance. As shown in Table 4, across five languages evaluated, linear alphabetic languages (English, Russian) yield the highest average overall SRs across agents (>17.5%). East Asian block characters (Chinese, Japanese) follow closely at 17.2% and 15.8%. Right-to-left Arabic lags at 13.7%, which is a 28.8% drop from English, reflecting agents’ difficulty perceiving either the visual complexity of Arabic glyphs or the mirrored UI layout, or both. GPT-4o exhibits especially uneven multilingual performance, with over 60% SR degradation on Russian and Arabic compared to English. If the task language and the environment language do not match, agent performance could further decline (see Appendix D.1).

5.3 Qualitative Analysis: Agent Behavior and Failure Modes

Open-source research models – Good grounding but poor planning. ShowUI demonstrates decent grounding ability but often executes nonsensical actions. For example, it attempts to open Settings by scrolling the mouse wheel on the desktop (Figure 9). It also tries to complete the task by typing it into the Help menu (Figure 8), which clearly shows that while it can identify basic interface elements, it lacks macOS-specific domain knowledge. In comparison, UI-TARS not only exhibits similarly nonsensical behaviors (Figure 10 and 14), but also fails due to illegal output formats and hallucinations. In Figure 15, it issues instructions to click a button by name rather than by coordinates, which is an illegal operation; in Figure 14, it mistakenly believes the Apple menu is located at the bottom-left corner of the screen. These results indicate that lightweight research models require substantial adaptation for mainstream macOS GUI tasks.

General-purpose proprietary VLMs – Good planning with imprecise grounding. Although GPT-4o’s performance is close to UI-TARS, its failure mode is more similar to Gemini 2.5 Pro,

Table 4: Performance (success rate, **SR**) of baseline agents on macOSWorld by language, averaged across all agents and all task categories excluding Advanced Apps. Performance degradations (Δ) are with respect to the English counterpart.

	en	ru	zh	ja	ar	Average
SR	19.3%	17.7%	17.2%	15.8%	13.7%	16.7%
Δ	-	-8.1%	-11.1%	-18.2%	-28.8%	-13.2%

Table 5: Agent performance on the macOSWorld safety subset under context-deception attacks, showing rates (%) of **Distracted** (clicking a decoy), **Gold** (clicking the close/cancel button), and **Unhandled** (pop-up not closed). Colors indicate highest and second highest rates in each row.

	Claude CUA	OpenAI CUA	GPT-4o	Gemini Pro 2.5	UI-TARS 7B	ShowUI	Average
Distracted	72.4%	69.0%	0.0%	17.2%	58.6%	41.4%	33.3%
Gold	24.1%	27.6%	0.0%	3.4%	34.4%	10.3%	10.9%
Unhandled	3.4%	3.4%	100.0%	79.3%	6.9%	44.8%	55.2%

another general-purpose VLM. Both GPT and Gemini struggle to precisely ground GUI elements, consistent with prior findings on general VLM agents in GUI settings [19, 20, 31]. While Gemini 2.5 Pro shows noticeably better grounding than GPT-4o, neither approaches the performance of specialized computer-use agents. Grounding UI elements with Set-of-Mark annotations [31] doubles GPT-4o’s performance (details in Appendix D.2), at 13.0% averaged across all languages (Table 8), but still lags Gemini’s 18.5%.

Proprietary CUAs – Suboptimal action efficiency. OpenAI and Claude CUAs achieve the highest overall SRs but exhibit inefficiencies in step budgeting. They often repeatedly fail and retry simple operations like creating folders (Figure 16 and 17), consuming more action steps than humans do and leaving insufficient step budget for downstream operations. For successful tasks, OpenAI and Claude CUAs take 15.95 and 19.85 steps on average, while human annotators use only 7.51 steps.

Multilingual discrepancies of CUAs – Performance gap stems from degraded grounding or planning. Differences in performance across languages are most pronounced for the two proprietary CUAs. For OpenAI CUA, the performance degradation stem primarily from degraded grounding: In one task (Figure 18), the agent effortlessly completed the account-creation task in English, but in Arabic, it fails to open Settings at all. Claude CUA exhibits both grounding and planning-related errors: while GUIs are mirrored in Arabic, Claude may carry biases from left-to-right UI layouts and incorrectly click on non-mirrored locations. It could also use more steps to complete the same task (Figure 19). These cross-language discrepancies underscore the necessity for language-aware training and UI layout adaptation to maintain consistent performance across diverse language environments.

5.4 Safety Subset Evaluation

As shown in Table 5, under context deception attacks, proprietary CUAs almost always handle the pop-up events (<4% unhandled) but exhibit high distraction rates ($\approx 70\%$), whereas proprietary VLMs leave the vast majority of pop-ups unhandled. Notably, despite UI-TARS and ShowUI’s low overall SR in previous experiments, they understand the pop-ups and were distracted half of the time on average. **These results reveal the vulnerability in GUI agents (particularly CUAs) to complex small-font deceptive UI designs, highlighting that context deception attacks could be more general and broadly effective**, going beyond prior studies [42, 45]. The findings underscore the need for imminent safety mechanisms.

6 Conclusion and Future Analysis

We presented macOSWorld, the first macOS interactive benchmark for GUI agents, with multilingual tasks, environments, and a safety evaluation subset. Results reveal clear performance tiers among agents and significant variation across language settings, highlighting the need for better adaptation to both the macOS system and multilingual environments. Future work could extend macOSWorld to non-binary reward evaluation, addressing the challenge of defining appropriate scoring rubrics, thus facilitating fine-grained benchmarking as well as reinforcement-learning-based agent training.

7 Acknowledgements

This research is supported by the National Research Foundation, Singapore under its AI Singapore Programme (AISG Award No: AISG3-RP-2022-030).

The authors thank Kevin Qinghong Lin, Zhiqiang Chen, Noorbakht Khan, Brandon Ng, Mingyu Ouyang, Siyuan Hu, Xiangwu Guo, Henry Hengyuan Zhao, Difei Gao, Christopher Rawles, and Kun Shao for their valuable discussions and feedback.

References

- [1] D. Gao, L. Ji, Z. Bai, M. Ouyang, P. Li, D. Mao, Q. Wu, W. Zhang, P. Wang, X. Guo *et al.*, “Assistgui: Task-oriented desktop graphical user interface automation,” *arXiv preprint arXiv:2312.13108*, 2023.
- [2] Y. Qin, Y. Ye, J. Fang, H. Wang, S. Liang, S. Tian, J. Zhang, J. Li, Y. Li, S. Huang *et al.*, “Ui-tars: Pioneering automated gui interaction with native agents,” *arXiv preprint arXiv:2501.12326*, 2025.
- [3] Y. Xu, Z. Wang, J. Wang, D. Lu, T. Xie, A. Saha, D. Sahoo, T. Yu, and C. Xiong, “Aguvis: Unified pure vision agents for autonomous gui interaction,” *arXiv preprint arXiv:2412.04454*, 2024.
- [4] L. Zheng, R. Wang, X. Wang, and B. An, “Synapse: Trajectory-as-exemplar prompting with memory for computer control,” *arXiv preprint arXiv:2306.07863*, 2023.
- [5] H. He, W. Yao, K. Ma, W. Yu, Y. Dai, H. Zhang, Z. Lan, and D. Yu, “Webvoyager: Building an end-to-end web agent with large multimodal models,” *arXiv preprint arXiv:2401.13919*, 2024.
- [6] C. Zhang, L. Li, S. He, X. Zhang, B. Qiao, S. Qin, M. Ma, Y. Kang, Q. Lin, S. Rajmohan *et al.*, “Ufo: A ui-focused agent for windows os interaction,” *arXiv preprint arXiv:2402.07939*, 2024.
- [7] S. Hu, M. Ouyang, D. Gao, and M. Z. Shou, “The dawn of gui agent: A preliminary case study with claude 3.5 computer use,” *arXiv preprint arXiv:2411.10323*, 2024.
- [8] H. Li, J. Su, Y. Chen, Q. Li, and Z.-X. ZHANG, “Sheetcopilot: Bringing software productivity to the next level through large language models,” *Advances in Neural Information Processing Systems*, vol. 36, pp. 4952–4984, 2023.
- [9] Z. Zhang and A. Zhang, “You only look at screens: Multimodal chain-of-action agents,” *arXiv preprint arXiv:2309.11436*, 2023.
- [10] H. Lai, X. Liu, I. L. Iong, S. Yao, Y. Chen, P. Shen, H. Yu, H. Zhang, X. Zhang, Y. Dong *et al.*, “Autowebglm: A large language model-based web navigating agent,” in *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, 2024, pp. 5295–5306.
- [11] B. Zheng, B. Gou, J. Kil, H. Sun, and Y. Su, “Gpt-4v (ision) is a generalist web agent, if grounded,” *arXiv preprint arXiv:2401.01614*, 2024.
- [12] K. Cheng, Q. Sun, Y. Chu, F. Xu, Y. Li, J. Zhang, and Z. Wu, “Seeclick: Harnessing gui grounding for advanced visual gui agents,” *arXiv preprint arXiv:2401.10935*, 2024.
- [13] S. Zhou, F. F. Xu, H. Zhu, X. Zhou, R. Lo, A. Sridhar, X. Cheng, T. Ou, Y. Bisk, D. Fried *et al.*, “Webarena: A realistic web environment for building autonomous agents,” *arXiv preprint arXiv:2307.13854*, 2023.
- [14] J. Y. Koh, R. Lo, L. Jang, V. Duvvur, M. C. Lim, P.-Y. Huang, G. Neubig, S. Zhou, R. Salakhutdinov, and D. Fried, “Visualwebarena: Evaluating multimodal agents on realistic visual web tasks,” *arXiv preprint arXiv:2401.13649*, 2024.
- [15] A. Drouin, M. Gasse, M. Caccia, I. H. Laradji, M. Del Verme, T. Marty, L. Boisvert, M. Thakkar, Q. Cappart, D. Vazquez *et al.*, “Workarena: How capable are web agents at solving common knowledge work tasks?” *arXiv preprint arXiv:2403.07718*, 2024.
- [16] Z. Zhang, S. Tian, L. Chen, and Z. Liu, “Mmina: Benchmarking multihop multimodal internet agents,” *arXiv preprint arXiv:2404.09992*, 2024.
- [17] Q. Chen, D. Pitawela, C. Zhao, G. Zhou, H.-T. Chen, and Q. Wu, “Webvln: Vision-and-language navigation on websites,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 38, no. 2, 2024, pp. 1165–1173.

- [18] Y. Pan, D. Kong, S. Zhou, C. Cui, Y. Leng, B. Jiang, H. Liu, Y. Shang, S. Zhou, T. Wu *et al.*, “Webcanvas: Benchmarking web agents in online environments,” *arXiv preprint arXiv:2406.12373*, 2024.
- [19] T. Xie, D. Zhang, J. Chen, X. Li, S. Zhao, R. Cao, T. J. Hua, Z. Cheng, D. Shin, F. Lei *et al.*, “Osworld: Benchmarking multimodal agents for open-ended tasks in real computer environments,” *Advances in Neural Information Processing Systems*, vol. 37, pp. 52 040–52 094, 2024.
- [20] R. Bonatti, D. Zhao, F. Bonacci, D. Dupont, S. Abdali, Y. Li, Y. Lu, J. Wagle, K. Koishida, A. Bucker *et al.*, “Windows agent arena: Evaluating multi-modal os agents at scale,” *arXiv preprint arXiv:2409.08264*, 2024.
- [21] C. Rawles, S. Clinckemaillie, Y. Chang, J. Waltz, G. Lau, M. Fair, A. Li, W. Bishop, W. Li, F. Campbell-Ajala *et al.*, “Androidworld: A dynamic benchmarking environment for autonomous agents,” *arXiv preprint arXiv:2405.14573*, 2024.
- [22] H. H. Zhao, D. Gao, and M. Z. Shou, “Worldgui: Dynamic testing for comprehensive desktop gui automation,” *arXiv preprint arXiv:2502.08047*, 2025.
- [23] H. Shen, C. Liu, G. Li, X. Wang, Y. Zhou, C. Ma, and X. Ji, “Falcon-ui: Understanding gui before following user instructions,” *arXiv preprint arXiv:2412.09362*, 2024.
- [24] W. Hong, W. Wang, Q. Lv, J. Xu, W. Yu, J. Ji, Y. Wang, Z. Wang, Y. Dong, M. Ding *et al.*, “Cogagent: A visual language model for gui agents,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 14 281–14 290.
- [25] K. Q. Lin, L. Li, D. Gao, Z. Yang, S. Wu, Z. Bai, W. Lei, L. Wang, and M. Z. Shou, “Showui: One vision-language-action model for gui visual agent,” *arXiv preprint arXiv:2411.17465*, 2024.
- [26] B. Gou, R. Wang, B. Zheng, Y. Xie, C. Chang, Y. Shu, H. Sun, and Y. Su, “Navigating the digital world as humans do: Universal visual grounding for gui agents,” *arXiv preprint arXiv:2410.05243*, 2024.
- [27] Z. Wu, Z. Wu, F. Xu, Y. Wang, Q. Sun, C. Jia, K. Cheng, Z. Ding, L. Chen, P. P. Liang *et al.*, “Os-atlas: A foundation action model for generalist gui agents,” *arXiv preprint arXiv:2410.23218*, 2024.
- [28] OpenAI, “Computer-using agent,” <https://openai.com/index/computer-using-agent/>, Jan. 2025, published: January 23, 2025; Accessed: 2025-05-03.
- [29] Anthropic, “Computer use (beta),” <https://docs.anthropic.com/en/docs/agents-and-tools/computer-use>, n.d., accessed: 2025-05-03.
- [30] J. Yang, H. Zhang, F. Li, X. Zou, C. Li, and J. Gao, “Set-of-mark prompting unleashes extraordinary visual grounding in gpt-4v,” *arXiv preprint arXiv:2310.11441*, 2023.
- [31] Y. Lu, J. Yang, Y. Shen, and A. Awadallah, “Omniparser for pure vision based gui agent,” *arXiv preprint arXiv:2408.00203*, 2024.
- [32] G. Mialon, C. Fourrier, T. Wolf, Y. LeCun, and T. Scialom, “Gaia: a benchmark for general ai assistants,” in *The Twelfth International Conference on Learning Representations*, 2023.
- [33] X. Deng, Y. Gu, B. Zheng, S. Chen, S. Stevens, B. Wang, H. Sun, and Y. Su, “Mind2web: Towards a generalist agent for the web,” *Advances in Neural Information Processing Systems*, vol. 36, pp. 28 091–28 114, 2023.
- [34] X. H. Lù, Z. Kasner, and S. Reddy, “Weblinx: Real-world website navigation with multi-turn dialogue,” *arXiv preprint arXiv:2402.05930*, 2024.
- [35] Y. Li, J. He, X. Zhou, Y. Zhang, and J. Baldridge, “Mapping natural language instructions to mobile ui action sequences,” *arXiv preprint arXiv:2005.03776*, 2020.
- [36] L. Sun, X. Chen, L. Chen, T. Dai, Z. Zhu, and K. Yu, “Meta-gui: Towards multi-modal conversational agents on mobile gui,” *arXiv preprint arXiv:2205.11029*, 2022.
- [37] C. Rawles, A. Li, D. Rodriguez, O. Riva, and T. Lillicrap, “Androidinthewild: A large-scale dataset for android device control,” *Advances in Neural Information Processing Systems*, vol. 36, pp. 59 708–59 728, 2023.
- [38] R. Kapoor, Y. P. Butala, M. Russak, J. Y. Koh, K. Kamble, W. AlShikh, and R. Salakhutdinov, “Omniact: A dataset and benchmark for enabling multimodal generalist autonomous agents for desktop and web,” in *European Conference on Computer Vision*. Springer, 2024, pp. 161–178.

- [39] C. H. Wu, J. Y. Koh, R. Salakhutdinov, D. Fried, and A. Raghunathan, “Adversarial attacks on multimodal agents,” *arXiv e-prints*, pp. arXiv–2406, 2024.
- [40] Z. Chen, Z. Xiang, C. Xiao, D. Song, and B. Li, “Agentpoison: Red-teaming llm agents via poisoning memory or knowledge bases,” *Advances in Neural Information Processing Systems*, vol. 37, pp. 130 185–130 213, 2024.
- [41] Z. Zhang, S. Cui, Y. Lu, J. Zhou, J. Yang, H. Wang, and M. Huang, “Agent-safetybench: Evaluating the safety of llm agents,” *arXiv preprint arXiv:2412.14470*, 2024.
- [42] Y. Zhang, T. Yu, and D. Yang, “Attacking vision-language computer agents via pop-ups,” *arXiv preprint arXiv:2411.02391*, 2024.
- [43] Z. Liao, L. Mo, C. Xu, M. Kang, J. Zhang, C. Xiao, Y. Tian, B. Li, and H. Sun, “Eia: Environmental injection attack on generalist web agents for privacy leakage,” *arXiv preprint arXiv:2409.11295*, 2024.
- [44] X. Ma, Y. Wang, Y. Yao, T. Yuan, A. Zhang, Z. Zhang, and H. Zhao, “Caution for the environment: Multimodal agents are susceptible to environmental distractions,” *arXiv preprint arXiv:2408.02544*, 2024.
- [45] P. Yang, H. Ci, and M. Z. Shou, “In-context defense in computer agents: An empirical study,” *arXiv preprint arXiv:2503.09241*, 2025.
- [46] OpenAI, A. Hurst *et al.*, “Gpt-4o system card,” 2024. [Online]. Available: <https://arxiv.org/abs/2410.21276>
- [47] Google Cloud, “Gemini 2.5 pro,” <https://cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/2-5-pro>, May 2025, last updated: 2025-05-02 UTC; Accessed: 2025-05-03.
- [48] J. Wei, X. Wang, D. Schuurmans, M. Bosma, F. Xia, E. Chi, Q. V. Le, D. Zhou *et al.*, “Chain-of-thought prompting elicits reasoning in large language models,” *Advances in neural information processing systems*, vol. 35, pp. 24 824–24 837, 2022.
- [49] H. Ci, P. Yang, Y. Song, and M. Z. Shou, “Ringid: Rethinking tree-ring watermarking for enhanced multi-key identification,” in *European Conference on Computer Vision*. Springer, 2024, pp. 338–354.
- [50] H. Ci, Y. Song, P. Yang, J. Xie, and M. Z. Shou, “Wmadapter: Adding watermark control to latent diffusion models,” *arXiv preprint arXiv:2406.08337*, 2024.
- [51] P. Yang, H. Ci, Y. Song, and M. Z. Shou, “Can simple averaging defeat modern watermarks?” *Advances in Neural Information Processing Systems*, vol. 37, pp. 56 644–56 673, 2024.
- [52] Y. Song, P. Yang, H. Ci, and M. Z. Shou, “Idprotector: An adversarial noise encoder to protect against id-preserving image generation,” *arXiv preprint arXiv:2412.11638*, 2024.

A Impact Statement

macOSWorld contributes to accessibility and inclusion in computing environments. By establishing the first multilingual interactive macOS benchmark, our work facilitates the development of more capable macOS GUI agents that can assist users with disabilities, limited technical knowledge, or language barriers. The benchmark's multilingual design actively promotes linguistic inclusivity in AI development, directly addressing underserved languages like Arabic where our results demonstrate a substantial 27.5% performance drop compared to English. This aligns with broader efforts to make AI technologies more equitable across diverse user populations.

B Ethical and Safeguarding Statement

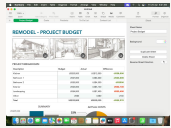
macOSWorld adheres to responsible AI principles through both its design choices and implementation safeguards. All main benchmark tasks focus exclusively on constructive applications that enhance accessibility and productivity, deliberately avoiding scenarios that could lead to harmful outcomes. While our safety benchmarking subset might inadvertently provide a blueprint for adversaries due to its effectiveness, we acknowledge this presents a double-edged sword, and we call for urgent development of defense mechanisms.

For responsible deployment, macOSWorld implements environment isolation for safeguarding. The benchmark operates solely within virtualized macOS environments on AWS Mac mini machines, isolating potential harmful operations from both host platforms and physical hardware. Our snapshot restoration process ensures all content, files, and system changes are immediately discarded after each evaluation, preventing any accumulation of sensitive information. To further facilitate auditability and traceability, additional techniques such as output watermarking could be employed [49–52].

C Example of a Real Task

A macOSWorld task comprises four key components – (1) a task instruction, (2) an AMI ID mapping, (3) a preparation script, and (4) an evaluation script. Figure 5 illustrates one example.

Task & Initial State



EN
In the project document, go to products and change my previous choice of the dishwasher to model B.



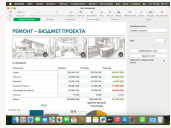
ZH
在这个项目文档中，找到products，然后把我之前打勾的洗碗机选项改成选择B型号的洗碗机。



AR
في مستند المشروع، اذهب إلى المنتجات وقم بتغيير اختياري السابق لغسالة الصحون إلى الطراز B.



JA
プロジェクト文書で、productsを見つけて、以前チェックした食器洗い機の選択肢をモデルBの食器洗い機に変更してください。



RU
В проектном документе найдите products и измените ранее выбранный мной вариант посудомоечной машины на модель B.

Preparation Script

```
osascript -e 'tell application "Numbers" to activate' && sleep 5 && osascript -e 'tell application "Numbers" to make new document with properties {document template:template "Personal Budget"}'
```

个人预算/الميزانية الشخصية/家計簿/Личный бюджет

Evaluation Script

```
osascript -e 'tell application "Numbers" to tell table 2 of sheet 3 of document 1 to get value of cell "C6"' 2>/dev/null | grep -q "true" && osascript -e 'tell application "Numbers" to tell table 2 of sheet 3 of document 1 to get value of cell "C7"' 2>/dev/null | grep -q "false" && echo "True" || echo "False"
```

Figure 5: Example of a task involving updating a dishwasher selection in a Numbers project document, together with its environment preparation and evaluation scripts.

Task Instruction In this example, the agent is instructed to open the current Numbers project document, navigate to the "Products" tab, and change the previously selected dishwasher model to "Model B". This instruction is provided in all five supported languages (English, Chinese, Arabic, Japanese, and Russian), ensuring that both the environment UI and the task text remain consistent with the agent's language setting.

Note that macOSWorld also allows cross-language evaluations, for example, evaluating with Chinese tasks in an Arabic macOS environment. This simulates the scenario, for example, a Chinese engineer helping an Arabic customer with computer usage.

AMI ID Mapping Before the agent begins, we must restore the macOS environment to the exact state in which the target document is already open. We do this by launching a pre-configured Amazon Machine Image (AMI) on AWS EC2. Each language variant maps to its own AMI ID. For example: { "en": "ami-07f4fd69378358c18", "zh": "ami-0abc1234def567890", ... }.

Each AMI contains the correct macOS version plus all necessary applications and files. For this task, the chosen AMI already has Numbers installed and the "Personal Budget" template available.

Preparation Script Once the EC2 instance is up, our testbench runs a language-specific preparation script over SSH (via zsh and AppleScript). The script given in Figure 5 (1) launches Numbers, (2) waits five seconds to avoid race conditions, and (3) opens a new document from the built-in "Personal Budget" template. Because template names differ across macOS languages, for some tasks, we maintain one script per language.

Agent Interaction & Expected Behavior After preparation, the agent enters its interaction loop. Here, it is expected to click to (a) select the "Products" tab, (b) untick "Dishwasher C", and (c) tick "Dishwasher B". Although this is a simple three-click task, it tests the agent's ability to ground instructions in a complex spreadsheet. In practice, however, many agents still fail. Figure 6 shows a representative failure case.

Evaluation Script Upon termination, the testbench invokes an AppleScript-based evaluation script via SSH to verify task success. The script reads two specific cells in the Products tab of the frontmost document:

1. The checkbox state (Cell C6) for Dishwasher B must be true.
2. The checkbox state (Cell C7) for Dishwasher C must be false.

Only if both conditions hold does the script echo True; otherwise it echo False. The testbench parses this return value from ssh-executed script to assign a binary reward.

Although this example uses a single evaluation script, many macOSWorld tasks employ multiple scripts to handle language-dependent window titles or alternative success criteria. While our platform supports non-binary rewards for intermediate states, the experiments reported in this paper remove those rewards and use only binary (success/failure) evaluations.

D Additional Benchmarking Scenarios

D.1 Cross-Lingual Evaluation

In many real-world settings, a user may submit a task prompt in one language while operating a system whose interface is rendered in another. To measure the impact of this mismatch, we repeat the OpenAI CUA experiments from Table 3 under cross-lingual conditions. Specifically, we fix the task prompt language to English and vary the environment language across the supported set: English, Chinese, Arabic, Japanese, and Russian.

Table 6 presents the mean Success Rate (SR) when task and environment languages are matched versus mismatched. To illustrate the degradation patterns, Table 7 shows a confusion matrix of overall SRs: each column corresponds to the environment language, and the diagonal entries report matched-language performance. Across all off-diagonal cells, the SR consistently falls below the matched case, confirming that **language mismatch between task prompt and UI leads to further performance degradation**.

Table 6: Cross-lingual performance of OpenAI CUA on macOSWorld. Env stands for the macOS system (GUI) language. The top 5 rows are where the task language aligns with the environment language. The bottom 5 rows are where the agent is given English tasks and operate in environments with different languages.

Language Task	Env	System & Interface	System Apps	File Ops	Productivity	Media	Multi-tasking	Total
en	en	41.4%	42.1%	27.6%	51.4%	8.3%	7.1%	33.3%
zh	zh	44.8%	42.1%	27.6%	45.7%	25.0%	3.6%	33.3%
ar	ar	13.8%	31.6%	37.9%	45.7%	33.3%	3.6%	28.1%
ja	ja	34.5%	44.7%	31.0%	54.3%	25.0%	7.1%	35.1%
ru	ru	51.7%	50.0%	27.6%	48.6%	41.7%	10.7%	39.2%
en	en	41.4%	42.1%	27.6%	51.4%	8.3%	7.1%	33.3%
en	zh	34.5%	34.2%	24.1%	48.6%	25.0%	3.6%	29.8%
en	ar	13.8%	23.7%	27.6%	45.7%	8.3%	3.6%	22.8%
en	ja	37.9%	42.1%	31.0%	40.0%	25.0%	0.0%	31.0%
en	ru	44.8%	44.7%	27.6%	37.1%	25.0%	3.6%	32.2%

Table 7: Confusion matrix of baseline OpenAI CUA performance in different task and environment languages. Values are success rates averaged across tasks. Results show that mismatch between task and environment languages could further degrade agent performance.

Task Language	Environment Language				
	en	cn	ar	ja	ru
en	33.3%	29.8%	22.8%	31.0%	32.2%
cn		33.3%			
ar			28.1%		
ja				35.1%	
ru					39.2%

D.2 Set-of-Mark (SoM) Annotation Evaluation

Previous work has alleviated grounding errors in general VLM agents (e.g., GPT-4o) by overlaying Set-of-Mark (SoM) annotations on interface elements [14, 19, 20, 31]. We evaluate GPT-4o [46] with and without SoM tags under the same benchmark configurations described in Section 5.1, with implementation details in Section G.2.

Table 8 compares GPT-4o’s SR across the five languages when using explicit SoM tag identifiers in place of raw screen coordinates. The results show that **SoM annotations roughly double GPT-4o’s success rates in every language**. Nevertheless, GPT-4o with SoM still trails behind Gemini Pro 2.5 and both proprietary CUAs. This gap indicates that **SoM annotations are a useful but insufficient workaround; agents specifically trained or finetuned for computer-use tasks maintain superior grounding and overall performance**.

Table 8: Performance of baseline GPT-4o agent with and without Set-of-Mark (SoM) annotations. The overlayed SoM labels and their transcripts were generated using OmniParser v2. Language indicates both the task prompt and system UI language. The Overall column reports the average success rate excluding tasks from the Advanced Apps category.

Language	SoM	Success Rate (↑)							
		System & Interface	System Apps	File Ops	Productivity	Media	Advanced Apps	Multi-Apps	Overall
en	✓	3.4%	13.2%	6.9%	14.3%	8.3%	0.0%	3.6%	8.8%
		6.9%	15.8%	31.0%	31.4%	33.3%	3.2%	7.1%	19.9%
zh	✓	3.4%	5.3%	6.9%	11.4%	8.3%	-	3.6%	6.4%
		6.9%	15.8%	24.1%	22.9%	0.0%	-	3.6%	14.0%
ar	✓	0.0%	2.6%	0.0%	5.7%	8.3%	-	3.6%	2.9%
		6.9%	7.9%	10.3%	5.7%	8.3%	-	3.6%	7.0%
ja	✓	3.4%	7.9%	6.9%	5.7%	0.0%	-	3.6%	5.3%
		10.3%	10.5%	10.3%	17.1%	0.0%	-	3.6%	9.9%
ru	✓	3.4%	2.6%	3.4%	5.7%	0.0%	-	0.0%	2.9%
		6.9%	13.2%	20.7%	25.7%	0.0%	-	7.1%	14.0%

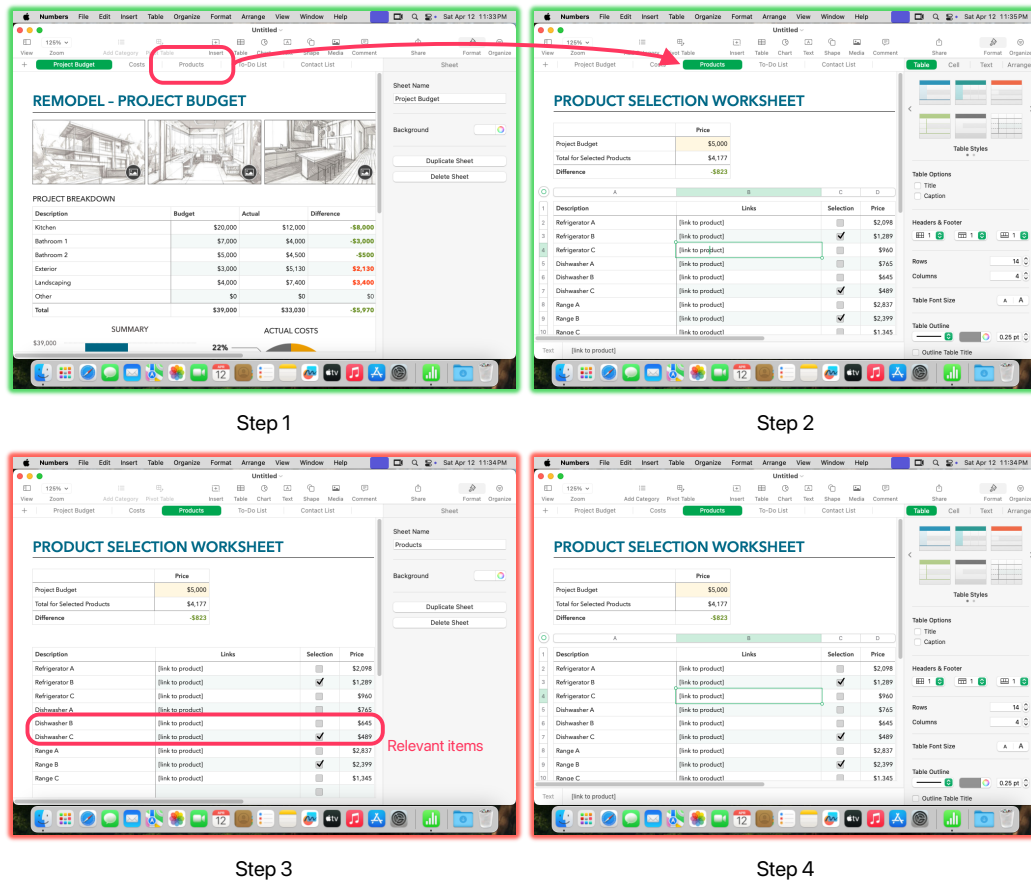


Figure 6: An example illustrating ShowUI's inconsistent capabilities. In this case, the task instruction is: "In the project document, go to products and change my previous choice of the dishwasher to model B." ShowUI immediately navigates to the Products tab in the complex Numbers spreadsheet, but fails to complete the subsequent task of changing the dishwasher model choice, despite this requiring only two simple clicks.

E More Analysis on Agent Behavior and Failure Modes

In this section, we support Section 5.3 with visualizations and more detailed analysis.

E.1 ShowUI

ShowUI [25] fails to complete the vast majority of tasks, with an end-to-end success rate averaging only 1.1% across five languages, significantly lower than computer-use agents (CUAs) and general VLM-based agents. Our analysis reveals that ShowUI demonstrates some understanding of complex interfaces and execution capabilities; however, it frequently performs nonsensical operations on simple tasks, likely due to its lack of domain knowledge. In other words, **ShowUI's primary limitation lies in planning.**

ShowUI still exhibits capability in understanding complex interfaces, but its performance is highly inconsistent. Figure 6 provides such an example, where ShowUI is asked to navigate to the "Products" tab in a given project document, before changing the choice of a dishwasher product. ShowUI successfully locates the Products tab at the top of the page and precisely clicks on it in the first step, demonstrating accurate comprehension of this complex Numbers document and sensitivity to small UI elements. However, ShowUI displays significant inconsistency in subsequent steps: it only needed to click twice to uncheck Dishwasher C and check Dishwasher B, but instead begins inexplicably editing the content of the Refrigerator C link, ultimately failing to complete the task.



Figure 7: An example of ShowUI's nonsensical operations. In this case, the task instruction is: "Change my hard disk volume name to `Local HD`." ShowUI correctly presses enter to initiate the renaming interface. At this point, the hard disk name is fully selected and editable, so ShowUI only needs to type "Local HD". However, it chooses to click on the hard disk icon, exiting the filename editing mode. Subsequently, ShowUI presses enter again to re-enter editing mode, then clicks the hard disk icon again to exit filename editing mode, repeating this cycle until the task step limit is exhausted.

This phenomenon of making nonsensical predictions on simple tasks is prevalent throughout ShowUI's task completion processes. As shown in Figure 7, during a renaming operation, ShowUI cycles repeatedly between entering and exiting renaming mode until the step limit is exhausted. Other nonsensical operations include inexplicably scrolling down on the desktop when needing to open Settings (Figure 9), or randomly right-clicking on the screen. These operations have no clear motivational explanation. In the example in Figure 8, ShowUI is asked to change the window minimization animation effect to "scale effect", but it immediately clicks the help menu at the beginning of the task and inputs this keyword into the search box. This provides direct evidence that ShowUI lacks domain knowledge of macOS usage, which may be the root cause of its series of nonsensical operations.

Additionally, ShowUI's behavior of inputting the task into the search box in Figure 8 exposes a potential security risk: under specific circumstances (such as in this case, where the task is difficult and the agent is unclear how to proceed), GUI agents risk directly inputting content from user-given tasks into the environment. This poses a potential privacy risk, such as when user-given task instructions contain sensitive information like passwords, and the agent chooses to input this content into a malicious input field constructed by an adversary.

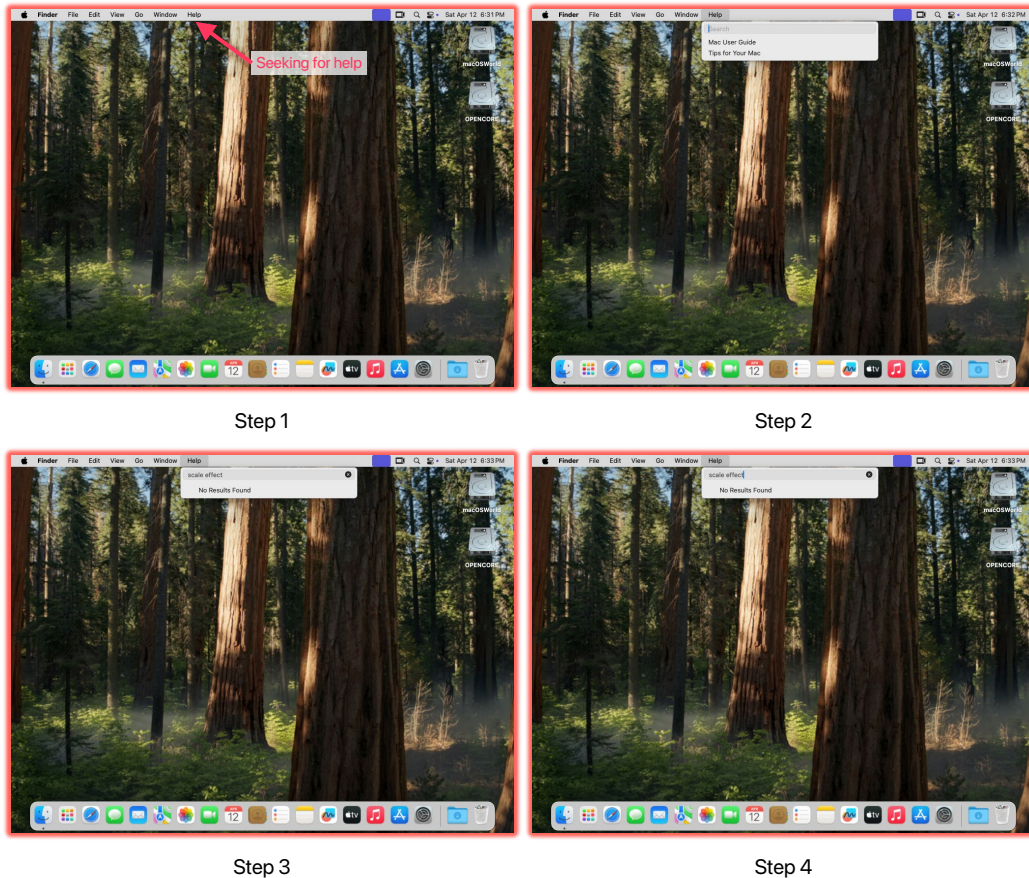


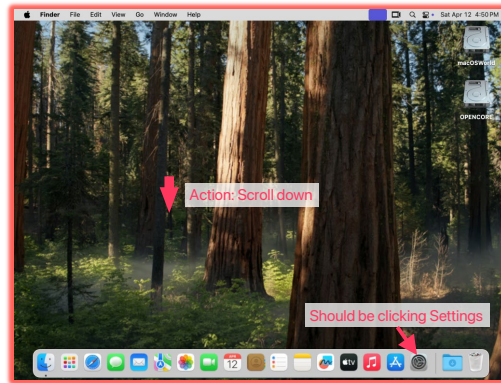
Figure 8: An example of ShowUI using the help button. In this case, the task instruction is: "Help me modify settings to use scale effect when minimizing windows." This is a system animation settings task, and the first step should be to open Settings. However, ShowUI chooses to seek assistance from the "Help" in the menu bar, and after a search box appears, directly inputs part of the user-specified task into the search box. This help-seeking behavior directly reflects ShowUI's lack of macOS domain knowledge.

E.2 UI-TARS

UI-TARS's performance lies between that of ShowUI and GPT-4o, achieving an average task success rate of 4.8% across all languages. It exhibits three primary failure modes: executing nonsensical actions, hallucinations, and producing outputs in invalid formats.

When executing nonsensical actions, UI-TARS behaves similarly to ShowUI, often in ways that are difficult to interpret. For example, in Figure 10, when the task requires opening System Settings via the Apple menu in the top-left corner of the screen, UI-TARS instead opens four unrelated menus (e.g. Edit, File) from the top bar. It is unclear whether this reflects a systematic attempt to locate System Settings or merely random, unconscious exploration. In some cases, this phenomenon may even occur during actions that UI-TARS had previously been able to complete successfully. For example, when asked to insert two blank tables into a document, as shown in Figure 11, UI-TARS successfully inserts the first one, but during the second attempt, it begins performing nonsensical actions, adjusting the zoom level of the view. Notably, UI-TARS's thought process during these tasks offers little insight into the rationale behind its actions, either. Its internal reasoning often only specifies what to do and how to do it, without explaining why.

Hallucinations in UI-TARS are most evident in its perception of screen content. For instance, it might mislocate the Apple icon, perceiving it on the bottom-left instead of the top-left (Figure 12), or



Step 1 – 15

Figure 9: An example of ShowUI failing to correctly open the settings panel. In this case, the task instruction is: "Someone else would be temporarily using my Mac. Help me create an account without a password. This account should be prohibited from changing my system settings, and that all changes made by this account (like files placed on the desktop) will not be saved upon logout." However, ShowUI continuously scrolls down rather than recognizing the need to click on the Settings icon in the Dock to first open Settings.

misestimate the aspect ratio of a slide on screen (Figure 13). Hallucination and nonsensical actions can also occur simultaneously: in step 4 in Figure 14, the agent attempts, without clear motivation, to click the Apple icon to close the Apple menu, while also hallucinating the Apple icon is located at the bottom-left of the screen. Considering that UI-TARS achieves strong quantitative performance across multiple platforms and benchmarks [2], such errors may stem from its lack of adaptation to macOS and limited domain-specific knowledge.

Another factor limiting UI-TARS's performance is its inconsistency in output formatting. Figure 15 shows invalid formats across two consecutive steps: the agent omits the "Thought:" prefix before stating its reasoning, includes redundant equals signs, quotation marks, and line breaks within the click action syntax, and incorrectly places a translated label of the UI element in place of a screen coordinate. Improving formatting consistency would enhance the agent's reliability and ensure that its outputs remain parsable by downstream algorithms.

Despite the room for improvement on macOS tasks, UI-TARS sometimes demonstrates self-correction capabilities. In some examples, it explicitly acknowledges past mistakes in its reasoning, as shown below:

Thought (originally in Chinese): Due to my incorrect action in
 ↳ the previous step, a dropdown menu labeled "Document" has
 ↳ appeared on the right side of the page. This does not meet
 ↳ the task requirements. I should left-click the "Document"
 ↳ tab in the upper-right corner of the page to close the
 ↳ dropdown menu.

E.3 OpenAI Computer-Using Agent and Claude Computer-Use Agent

Both CUAs outperform other agents in our benchmark, achieving average success rates of around 35% across all languages. While they do not exhibit obvious error patterns within individual languages, they demonstrate performance disparities across different languages, with weaknesses in less proficient languages manifesting as diminished grounding and planning capabilities.

Within a single language, these agents do not display obvious shortcomings in specific capabilities like ShowUI or UI-TARS; rather, task failures primarily stem from imperfect handling of various details. For example, in a file organization task requiring the creation of two folders followed by sorting images and documents into their respective directories, both CUAs only managed to create the folders before exhausting their step limits. OpenAI CUA (Figure 16) struggles to select files

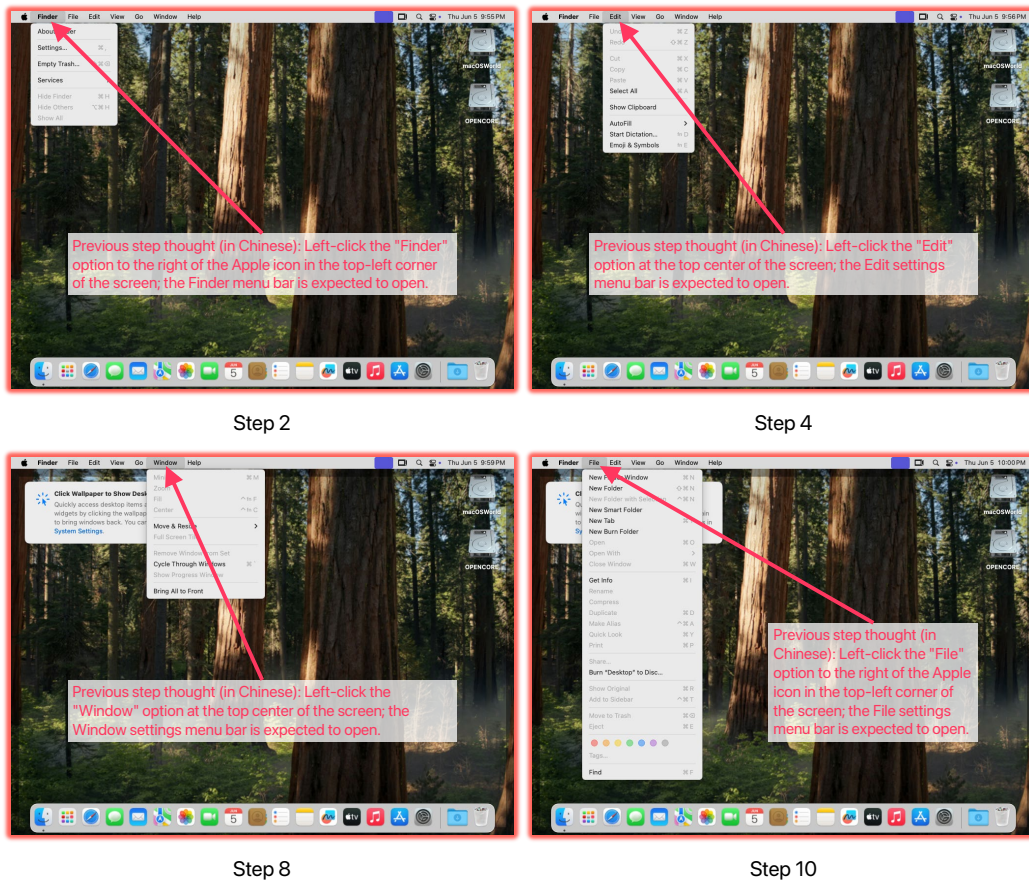


Figure 10: An example where UI-TARS executes nonsensical actions. In this case, the task instruction is: "Please help me modify the settings so that I can view a larger version of text when I am typing." UI-TARS should open System Settings via the Apple menu in the top left corner or click the gear icon in the Dock below, but it does neither. Instead, it successively opens the Finder, Edit, Window, and File menus from the top bar. It's unclear whether it is searching through each menu in hopes of finding the entrance to System Settings or mindlessly opening menus without intent.

correctly before dragging them into each folder after creating the folders. It inadvertently selects unintended content, such as simultaneously selecting newly created folders while attempting to select files, and repeatedly makes the same mistakes, demonstrating significant room for improvement in its perception and handling of details. Although it exhibits error correction capabilities, such as closing mistakenly opened windows and returning to the directory where files need categorization, this correction consumes numerous steps, directly leading to step limit exhaustion. Claude CUA (Figure 17) fails to observe that text is already in selected state when renaming a folder, where it could directly input the new filename "Documents." Instead, it clicks the text once, deselecting all text, causing subsequent keystrokes to append to rather than replace "untitled folder", resulting in an incorrectly named folder ("untitledDocumentsd folder") while consuming substantial steps. This example illustrates the primary error pattern for both CUAs – task failures triggered by detail-level issues. In interactive environments, minor errors can have cascading effects, as the numerous additional steps required to correct detailed errors easily amplify the impact of mistakes, preventing agents from completing tasks within specified step limits.

OpenAI CUA's performance degradation in less proficient languages primarily manifests as reduced grounding capability. Figure 18 presents an example where the agent successfully completes a task in English, Russian, and Chinese but struggles with the first step in Arabic. This complex task requires creating a guest account on the computer, with the first step being opening settings. However, while the agent can smoothly open settings and proceed with subsequent operations in other languages,

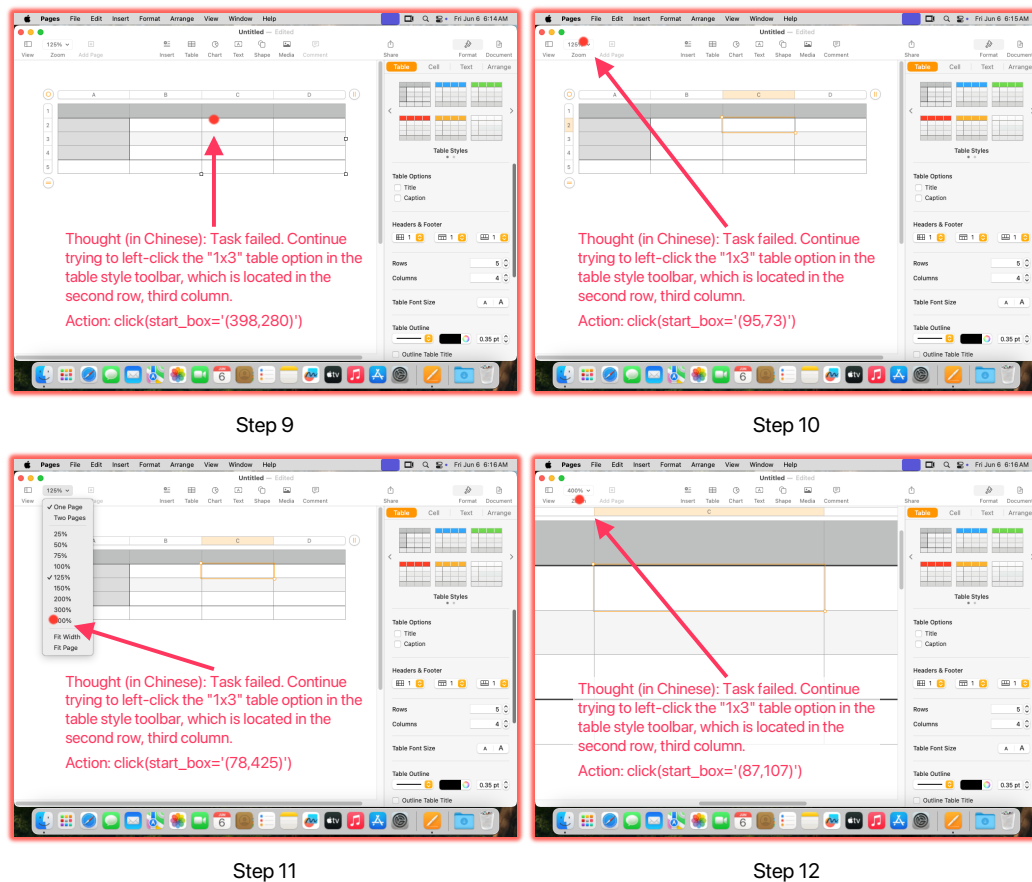


Figure 11: An example where UI-TARS executes nonsensical actions. In this case, the task instruction is: "Help me create a blank Pages document and insert two separate empty tables in it. No need to configure the details of the tables." The task requires the agent to insert two tables. The agent has already inserted one and now needs to insert the second. To do this, it only needs to repeat the previous steps by clicking the Table icon at the top of the interface again and selecting any table. However, it first clicks a cell in the existing table and then starts adjusting the zoom level of the view. These actions are inexplicable, illogical, and the agent's thoughts and actions are not aligned.

it repeatedly misclicks when attempting to open settings from either the Dock or Apple menu in Arabic. OpenAI CUA's ability to successfully complete this complex task in three other languages but repeatedly encounter basic grounding issues in Arabic demonstrates its performance degradation in multi-language scenarios manifesting as diminished grounding capabilities.

Claude CUA's performance degradation in less proficient languages includes not only reduced grounding capability but also diminished planning ability to some extent. Figures 19(a) and (b) show cases where Claude CUA incorrectly clicks when opening system settings in Arabic, targeting icons on the right side of the Dock. In other languages, the System Settings icon defaults to the right side, but Arabic system environments mirror the UI, placing the Settings icon on the left. This demonstrates Claude CUA's difficulty adapting to mirrored UI layouts. Figure 19(c) illustrates the only successful method used by Claude CUA to open settings in Arabic – through the System Settings option in the Apple Menu in the top right corner, requiring two steps. In contrast, Claude CUA in the other four languages directly clicks the icon in the Dock, opening settings in a single step. This indicates a decline in Claude CUA's planning ability – when an icon is placed in a mirrored location, it fails to recognize or locate it effectively.

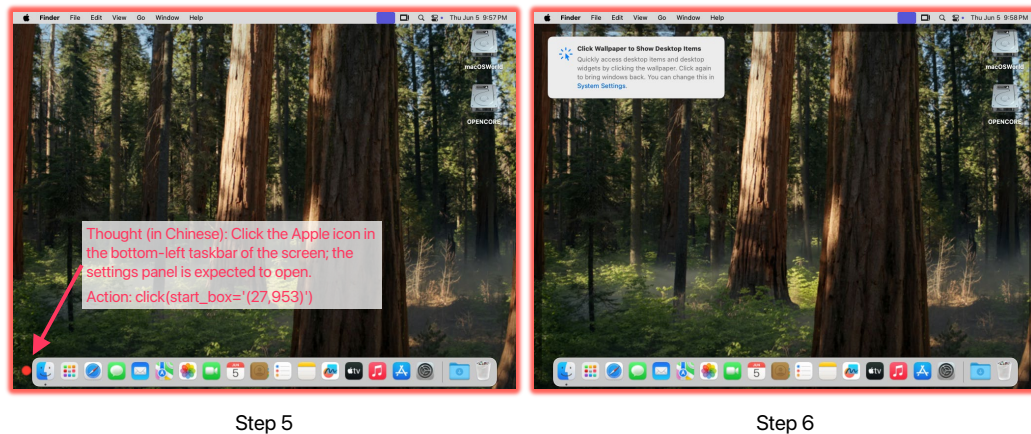


Figure 12: An example demonstrating UI-TARS’s hallucination. In this case, the task instruction is: "Please help me modify the settings so that I can view a larger version of text when I am typing." To change this setting, the agent intends to open System Settings through the Apple menu in the top left corner, but it hallucinates that the Apple icon is in the bottom left and clicks there instead.

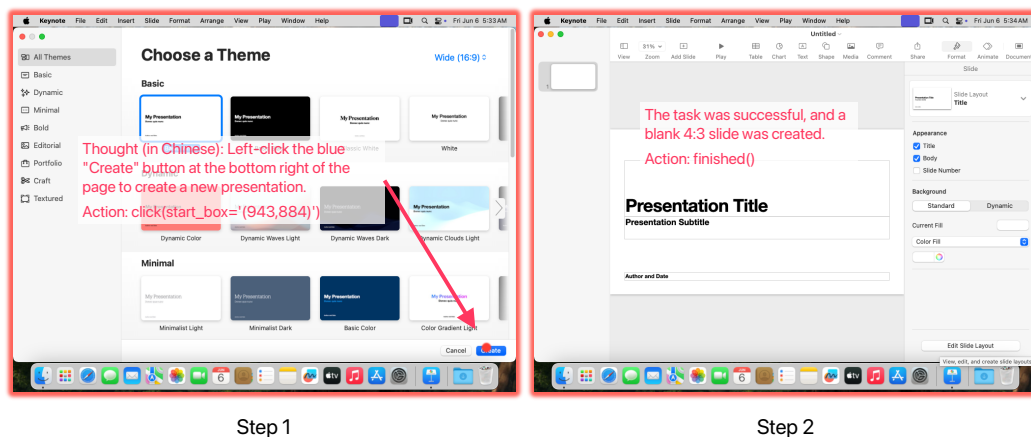


Figure 13: An example where UI-TARS incorrectly believes it has completed the task as instructed. In this case, the task instruction is: "Help me create an empty slide with an aspect ratio of 4:3." UI-TARS directly clicks the Create button in the bottom right corner of the window in the first step, ignoring the aspect ratio setting located in the top right corner of the interface. Although the created presentation has a 16:9 aspect ratio, the agent still claims that its aspect ratio is 4:3.

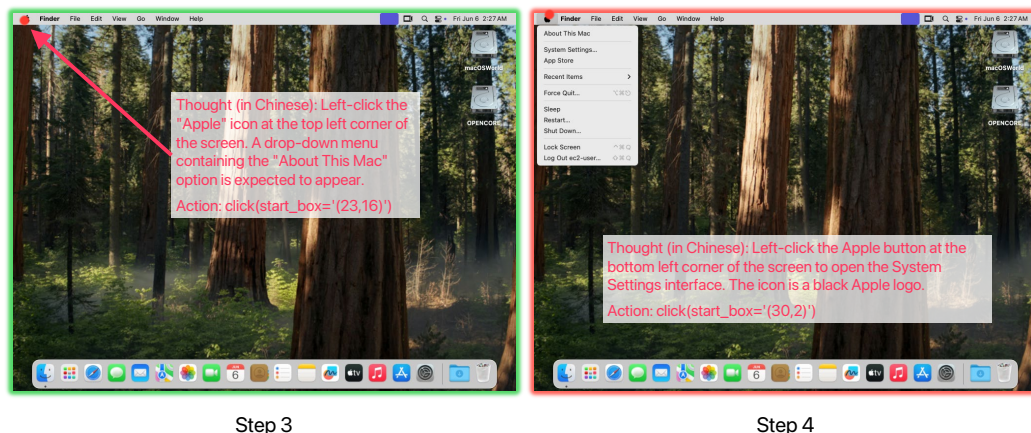
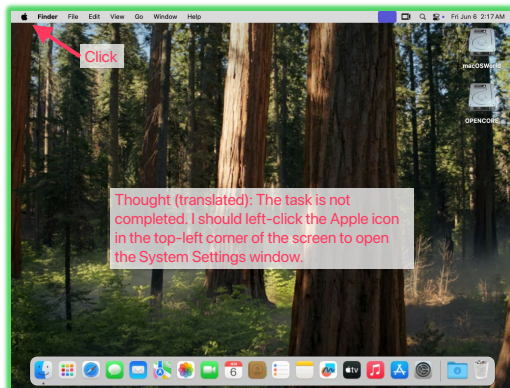


Figure 14: A UI-TARS example containing multiple issues simultaneously. In this case, the task instruction is: "Help me modify keyboard settings, toggling the key repeat rate to the fastest possible." The agent needs to first open the Apple menu and then open System Settings from there. In Step 3, the agent correctly opens the Apple menu, but its thought reflects an expectation of seeing "About This Mac" instead of "System Settings". Then in Step 4, the agent nonsensically tries to close the Apple menu. Its thoughts indicate that it believes the Apple icon is in the bottom left corner of the screen, but it still correctly clicks the top left corner of the screen to close the Apple menu it just opened. These two consecutive action steps include both nonsensical actions and hallucinations, and the thoughts and actions are not aligned, highlighting the unexplainable nature of UI-TARS failure cases.

F Environment Implementation Details

The macOSWorld environment runs macOS Sequoia 15.2, with at least 30 applications involved in benchmarking. These applications are listed in Table 9.



Step 5 Screenshot

Actual Output

任务未完成，我应该左键单击屏幕左上角的苹果图标，打开系统设置窗口。

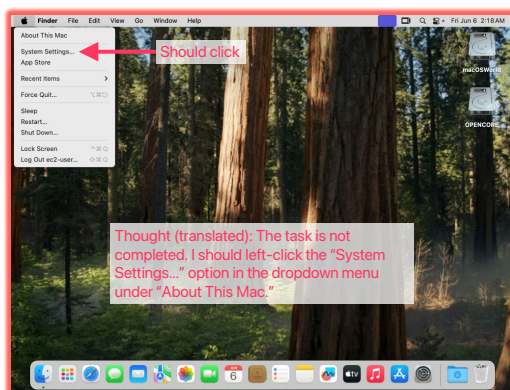
Action: click(start_box='(20,14)')

Desired Output

Thought: 任务未完成，我应该左键单击屏幕左上角的苹果图标，打开系统设置窗口。

Action: click(start_box='(20,14)')

Step 5 Model Output



Step 6 Screenshot

Actual Output

任务未完成，我应该左键单击页面“About this Mac”下拉框中的“System Settings...”选项。

Action: click(start_box='
=
='系统设置')

Desired Output

Thought: 任务未完成，我应该左键单击屏幕左上角的苹果图标，打开系统设置窗口。

Action: click(start_box='(85,96)')

Step 6 Model Output

Figure 15: An example where UI-TARS does not strictly adhere to its output format. In this case, the task instruction is: "In 'System Settings > Appearance', set accent color to red." The first invalid output in this example is the absence of the keyword "Thought:", which may cause some thought-parsing algorithms to fail. The second invalid output appears in step 6's click action parameters. It was supposed to be a coordinate, but instead, the output repeatedly contains equal signs, single quotes, and newline characters. Moreover, where the coordinates should appear, the agent outputs the Chinese translation of the key name at that location. This illustrates that UI-TARS is not sufficiently robust to consistently adhere to its fixed output format.

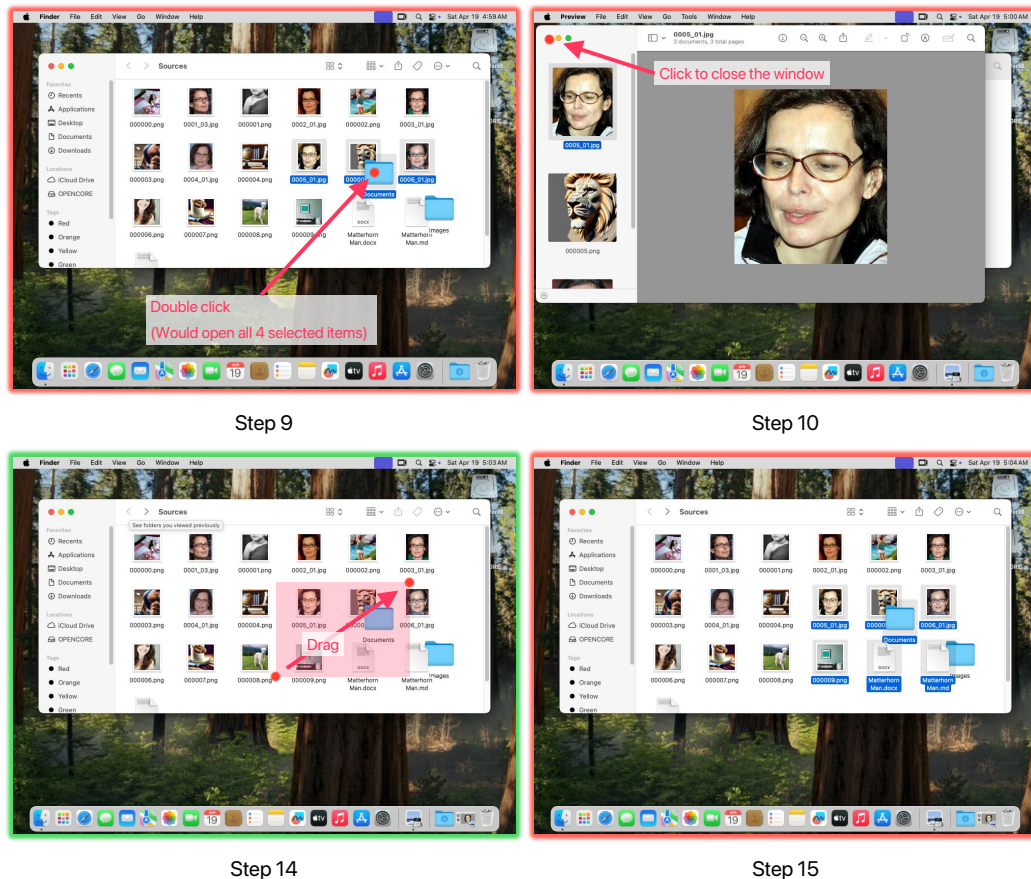


Figure 16: An example illustrating OpenAI CUA's insufficient attention to details, requiring numerous steps to complete simple tasks. In this case, the task instruction is: "In the current directory, help me create two folders named 'Documents' and 'Images'. Move the documents (txt, markdown, etc.) and the images to the corresponding folder." In Step 8, the agent erroneously selects three images along with the Documents folder, while planning to enter the Documents folder (although this is unnecessary, as it should instead select files by category and drag them to appropriate folders). However, it fails to deselect these image files, causing the double-click operation to open all selected files. In subsequent steps, it closes the mistakenly opened images and returns to the state before Step 9 after waiting two steps. Yet again, it inexplicably selects an illogical area through a drag operation, simultaneously selecting documents, images, and folders. This demonstrates that OpenAI CUA still has significant room for improvement in its perception and control of details. In interactive environments, small errors require substantial steps to rectify, as a minor issue can consume a large portion of the step budget.

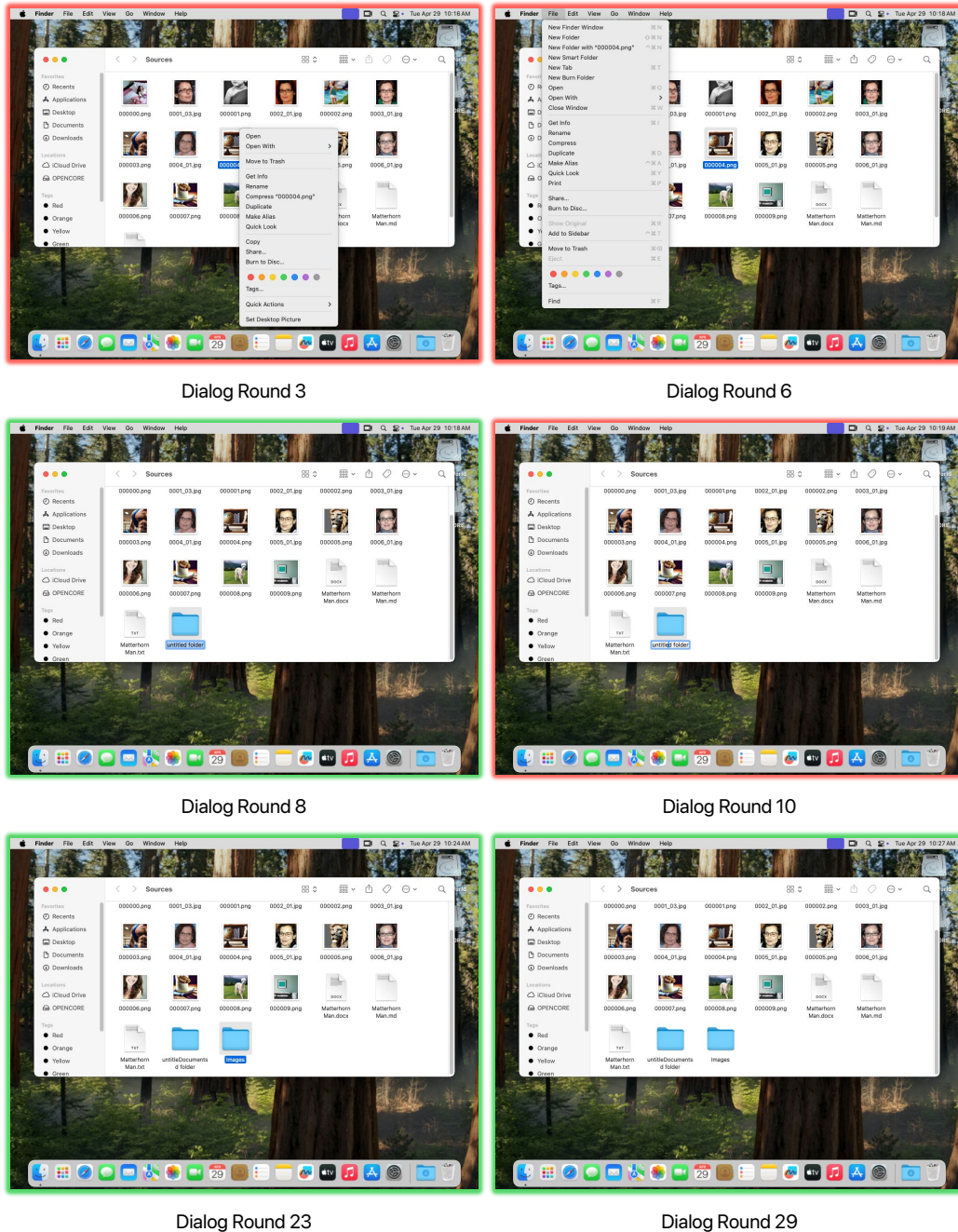


Figure 17: An example illustrating Claude CUA's insufficient attention to details, requiring numerous steps to complete simple tasks. In this case, the task instruction is: "In the current directory, help me create two folders named 'Documents' and 'Images'. Move the documents (txt, markdown, etc.) and the images to the corresponding folder." The first issue in this example is that the agent fails to move the mouse to an empty space before right-clicking in step three, resulting in opening a file context menu rather than a folder context menu. After requesting the next screenshot and recognizing this problem, the agent uses two dialogue rounds to first close the context menu by clicking on an empty area (an unnecessary step), then clicks on the file menu at the top of the screen to create new folders from there. The second issue is that after creating a new folder, the folder text is already in selected state, requiring only direct typing of the new name; however, the agent chooses to click on the selected text (dialogue round 10), a redundant operation that directly leads to subsequent folder naming errors. Finally, Claude CUA exhausted its 30-round dialogue budget just creating two folders.

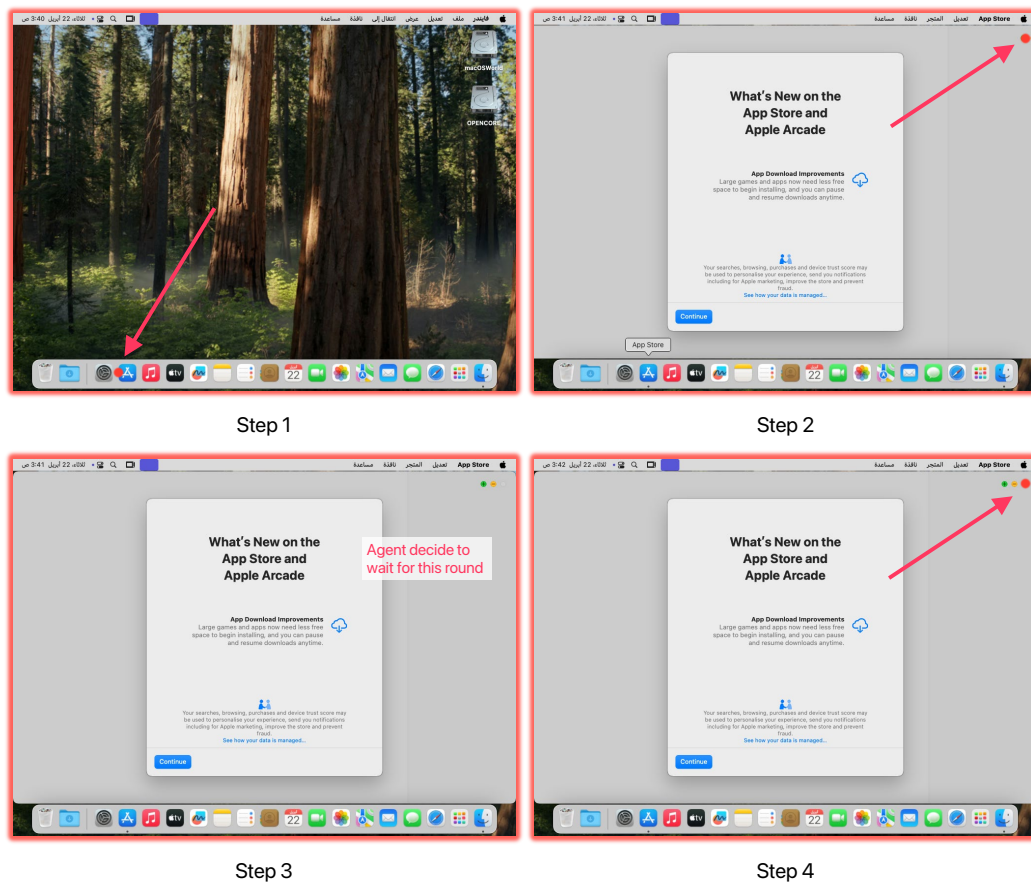


Figure 18: An example demonstrating OpenAI CUA's performance degradation in Arabic. In this case, the task instruction is: "Someone else would be temporarily using my Mac. Help me create an account without a password. This account should be prohibited from changing my system settings, and that all changes made by this account (like files placed on the desktop) will not be saved upon logout." This complex task is successfully completed by OpenAI CUA in English, Chinese, and Russian environments, but it fails to correctly execute even the first step in Arabic. The task involves creating a guest user, with the first step being opening system settings. However, OpenAI CUA repeatedly clicks incorrectly, opening the App Store adjacent to Settings instead. In the second step, the agent changes approach, attempting to open system settings through the Apple Menu in the top left corner, but misses again. In subsequent rounds, the agent opts to wait, then attempts to close the erroneously opened App Store by clicking the red button in the top right corner, but fails due to an App Store popup obstruction. In later steps, it continues attempting to open settings from the dock icons but clicks on Reminders toward the center of the screen, and when trying to close Reminders, clicks the full-screen button and consistently fails to exit. The fact that OpenAI CUA can successfully complete this complex task in three other languages but repeatedly encounters basic grounding issues in Arabic demonstrates its performance degradation in multi-language scenarios manifesting as diminished grounding capabilities.

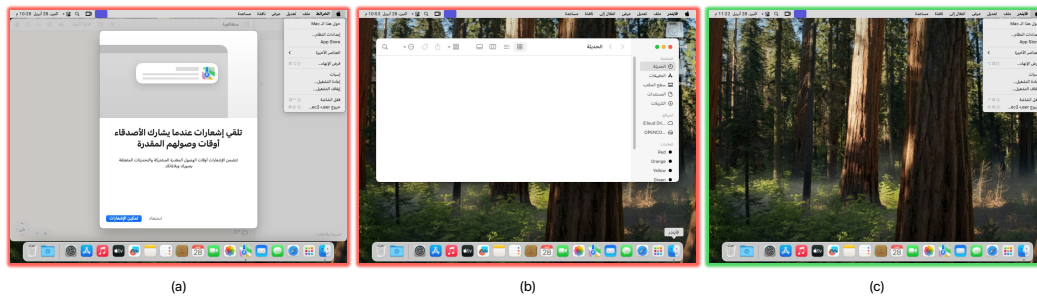


Figure 19: Claude CUA's approach to opening System Settings in Arabic environment and associated issues. (a) and (b) demonstrate instances where Claude CUA attempts to open settings through the Dock but clicks on other icons on the right side of the Dock (Maps and Finder), illustrating Claude CUA's difficulty adapting to mirrored UI environments in Arabic. (c) shows the only successful method used by Claude CUA to open settings in Arabic – clicking on System Settings from the Apple Menu in the top right corner, demonstrating that Claude CUA's planning capabilities also decline in Arabic, tending to require more steps to accomplish the same tasks.

Table 9: List of applications available in the macOSWorld environment.

Application	Version	Comment
Activity Monitor	10.14	
Automator	2.1.0	AppleScript 2.8
Calculator	11.0	
Calendar	15.0	
Chess	3.18	
ColorSync Utility	12.1.0	
Contacts	14.0	
Dictionary	2.3.0	
Digital Color Meter	5.26	
Disk Utility	22.7	
Finder	15.2	
Font Book	11.0	
Freeform	3.2	
iMovie	10.4.3	
Keynote	14.3	
Maps	3.0	
Music	1.5.2.26	
Notes	4.11	
Numbers	14.3	
Pages	14.3	
Preview	11.0	
QuickTime	10.5	Also available on Windows
Reminders	7.0	
Safari	18.2	Also available on Windows
Script Editor	2.11	
Stickies	10.3	
Stocks	7.1	
System Settings	15.0	
Voice Memos	3.1	
Weather	5.0	
Xcode	16.2	

G Agent Implementation Details

G.1 GPT-4o

The GPT-4o [46] agent was implemented by prompting the agent each time with $T=1$ and $\text{top}_p=0.9$, with the following content blocks:

```
<System Prompt>
<User Query> (including screenshots)
```

The system prompt is given by:

```
You are an agent that performs Mac desktop computer tasks by
  ↳ controlling mouse and keyboard through VNC. For each step,
  ↳ you will receive a screenshot observation of the computer
  ↳ screen and should predict the next action.

Your output must be raw text commands with the following
  ↳ structure:
...
<action_name> <parameter_1> <parameter_2>
<action_name> <parameter_1> <parameter_2>
...

For example:
...
move_to 0.25 0.5
key_press command-c
left_click
...

Available actions and their parameters:

1. Mouse Actions:
- "move_to": Move cursor to normalized coordinates
  Required params: {"x": float 0-1, "y": float 0-1}

- "left_click": Perform left mouse click
  No params required

- "middle_click": Perform middle mouse click
  No params required

- "right_click": Perform right mouse click
  No params required

- "double_click": Perform double left click
  No params required

- "triple_click": Perform triple left click
  No params required

- "drag_to": Drag with the left mouse button to a specified
  ↳ coordinate.
  Required params: {"x": float 0-1, "y": float 0-1}

- "mouse_down": Press and hold a mouse button.
  Required params: {"button": string ("left", "middle", "right")}

- "mouse_up": Release a mouse button.
  Required params: {"button": string ("left", "middle", "right")}

- "scroll_down": Scroll down by proportion of screen height
  Required params: {"amount": float 0-1}
```

```

- "scroll_up": Scroll up by proportion of screen height
  Required params: {"amount": float 0-1}

- "scroll_left": Scroll up by proportion of screen width
  Required params: {"amount": float 0-1}

- "scroll_right": Scroll up by proportion of screen width
  Required params: {"amount": float 0-1}

2. Keyboard Actions:
- "type_text": Type ASCII text
  Required params: {"text": string}
  Everything after `type_text` will be parsed as parameter 1,
    ↳ including spaces. No need to escape any characters.

- "key_press": Press a key or key combination.
  Required params: {"key": string}
  Available keys: ctrl, command, option, backspace, tab, enter,
    ↳ esc, del, left, up, right, down, or single ASCII
    ↳ characters
  When pressing a combination of keys simultaneously, connect the
    ↳ keys using `~`, for example, `command-c` or `ctrl-alt-
    ↳ del`

3. Control Actions:
- "wait": Wait for specified seconds
  Required params: {"seconds": float}

- "fail": Indicate task cannot be completed
  No params required

- "done": Indicate task is already finished
  No params required

Important Notes:
- Your username is "ec2-user" and password is "000000"
- All coordinates (x,y) should be normalized between 0 and 1
- All scroll amounts should be normalized between 0 and 1
- Only ASCII characters are allowed for text input
- The control commands (wait, fail, done) must be the only
    ↳ command issued in a round. If one of these commands is
    ↳ used, no other actions should be provided alongside it.
- Return only the actions in a backtick-wrapped plaintext code
    ↳ block, one line per action, no other text

```

Similar to [19, 29], in order to avoid unreasonably long context windows that may possibly degrade agent performance, the user query would include a rolling window of only $n = 3$ most-recent screenshots. The user query is given by:

```

Task: <Task Prompt>
Screenshot: <Current Screenshot>
Rolling window of historical screenshots in chronological order:
    ↳ <Screenshot t-n> <Screenshot t-n+1> ... <Screenshot t-1>

```

G.2 GPT-4o with Set-of-Mark Annotations

The implementation of GPT-4o [46] with Set-of-Mark (SoM) annotations [30, 31] is similar to the version without the annotations. The changes are:

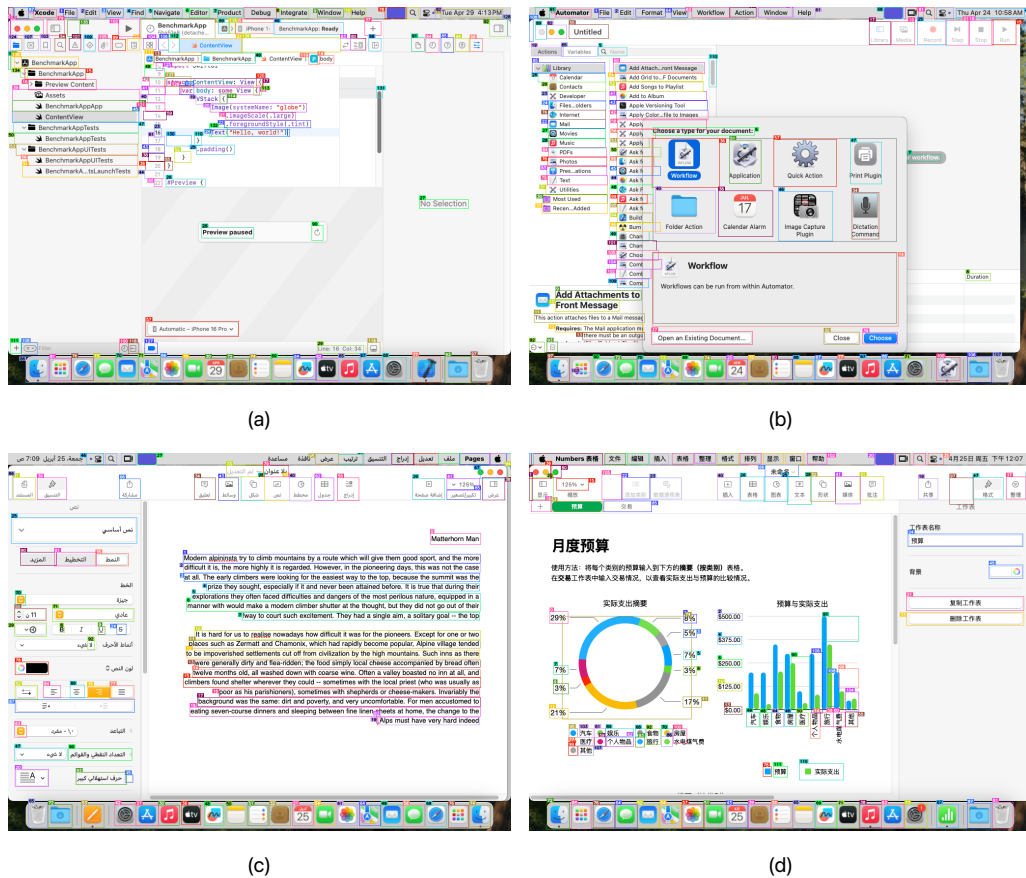


Figure 20: Visualization of SoM-annotated screenshots of (a) text-intensive Xcode interface, (b) logo-intensive Automator interface, (c) Arabic interface with English texts, and (d) Chinese interface with charts.

1. The screenshots are now annotated with SoM bounding boxes and labels (Figure 20).
2. The agents are now provided with transcripts of the SoM-parsed contents from the current screenshot.
3. The agents are now allowed to use tag numbers to substitute for coordinates, aiding inaccurate grounding issues.

The SoM labels and the corresponding transcript were generated by passing the screenshot to OmniParser v2 [31]. The base prompting format remains the same, with one system prompt and a user query, with the change that the user query now includes the SoM-annotated screenshots along with the text transcripts.

GPT-4o was still prompted with $T=1$ and $\text{top}_p=0.9$. The updated system prompt is:

```
You are an agent that performs Mac desktop computer tasks by
  ↳ controlling mouse and keyboard through VNC. For each step,
  ↳ you will receive a screenshot observation of the computer
  ↳ screen and should predict the next action.

Your output must be raw text commands with the following
  ↳ structure:
...
<action_name> <parameter_1> <parameter_2>
<action_name> <parameter_1> <parameter_2>
...
```

```

...

For example:
...
move_to <tag>15</tag>
key_press command-c
left_click

move_to 0.25 0.5          # Example of using {"x": float 0-1, "y":
    ↪ float 0-1}
drag_to <tag>16</tag>      # Example of using {"tag": tag} as
    ↪ coordinate, preferred

move_to <tag>26</tag>
scroll_down 0.5
...

Available actions and their parameters:

1. Mouse Actions:
- "move_to": Move cursor to normalized coordinates
  Required params: {"tag": tag} or {"x": float 0-1, "y": float
    ↪ 0-1}

- "left_click": Perform left mouse click
  No params required

- "middle_click": Perform middle mouse click
  No params required

- "right_click": Perform right mouse click
  No params required

- "double_click": Perform double left click
  No params required

- "triple_click": Perform triple left click
  No params required

- "drag_to": Drag with the left mouse button to a specified
    ↪ coordinate.
  Required params: {"tag": tag} or {"x": float 0-1, "y": float
    ↪ 0-1}

- "mouse_down": Press and hold a mouse button.
  Required params: {"button": string ("left", "middle", "right")}

- "mouse_up": Release a mouse button.
  Required params: {"button": string ("left", "middle", "right")}

- "scroll_down": Scroll down by proportion of screen height
  Required params: {"amount": float 0-1}

- "scroll_up": Scroll up by proportion of screen height
  Required params: {"amount": float 0-1}

- "scroll_left": Scroll up by proportion of screen width
  Required params: {"amount": float 0-1}

- "scroll_right": Scroll up by proportion of screen width
  Required params: {"amount": float 0-1}

2. Keyboard Actions:
- "type_text": Type ASCII text
  Required params: {"text": string}

```

```

Everything after `type_text` will be parsed as parameter 1,
  ↳ including spaces. No need to escape any characters.

- "key_press": Press a key or key combination.
  Required params: {"key": string}
  Available keys: ctrl, command, option, backspace, tab, enter,
    ↳ esc, del, left, up, right, down, or single ASCII
    ↳ characters
  When pressing a combination of keys simultaneously, connect the
    ↳ keys using `-`, for example, `command-c` or `ctrl-alt-
    ↳ del`

3. Control Actions:
- "wait": Wait for specified seconds
  Required params: {"seconds": float}

- "fail": Indicate task cannot be completed
  No params required

- "done": Indicate task is already finished
  No params required

Important Notes:
- Your username is "ec2-user" and password is "000000"
- All coordinates (x,y) should be normalized between 0 and 1
- All scroll amounts should be normalized between 0 and 1
- Only ASCII characters are allowed for text input
- The control commands (wait, fail, done) must be the only
  ↳ command issued in a round. If one of these commands is
  ↳ used, no other actions should be provided alongside it.
- Tags must be wrapped as xml elements (e.g. `

```

The updated user query is:

```

Task: <Task Prompt>
Rolling window of historical screenshots in chronological order:
  ↳ <Screenshot t-n> <Screenshot t-n+1> ... <Screenshot t-1>
Current screenshot: <Current Screenshot>
Labeled annotations (with corresponding tagged bounding boxes
  ↳ shown in the current screenshot): <SoM transcript>

```

We follow OSWorld's method [19] during action parsing, substituting the tag label with the center coordinate of the corresponding bounding box, so that when the agent "clicks on tag 18" or "drags towards tag 18", it is essentially clicking on or dragging to the center location of tag 18's bounding box.

G.3 Gemini Pro 2.5

Gemini Pro 2.5 [47] was prompted the same way as GPT-4o (including that $T=1$ and $\text{top}_p=0.9$), with the only difference that the system prompt and the user query were concatenated together before providing to the VLM. This is because Gemini Pro 2.5 does not accept system prompts, unlike OpenAI models.

G.4 OpenAI Computer-Using Agent

Unlike general VLMs, OpenAI CUA does not require an additional system prompt to tell it to behave like a computer agent [28]. We adopted the parameters $T=1$ and $\text{top}_p=0.9$ during generation, and followed its designed conversation format, with the following four modifications:

1. We added the following system prompt to bridge it to our action space, mainly adding the control signals (like "fail" or "done") and specifying keys available.

```

You are using a macOS computer to complete a user-
  ↳ given task. Additional Notes:
* Available xdotool keys: ctrl, command, option,
  ↳ backspace, tab, enter, esc, del, left, up, right
  ↳ , down, and single ASCII characters.
* When you think the task can not be done, say ```FAIL
  ↳ ``` , don't easily say ```FAIL``` , try your best
  ↳ to do the task. When you think the task is
  ↳ completed, say ```DONE``` . Include the three
  ↳ backticks. If the task is not completed, don't
  ↳ raise any of these two flags.
* You may need my username and password. My username
  ↳ is `ec2-user` and password is `000000`.

```

2. The agent is provided with screenshots after each computer tool use operation, not only the screenshot capturing operation.
3. We keep the entire conversation history, because OpenAI CUA API does not allow removing an entire conversation block from the conversation history.
4. Filtering to $n = 3$ most recent screenshots was implemented by substituting the earlier screenshots with a 5×5 black image.

G.5 Claude Computer Use Agent

Similar to OpenAI CUA [28], we prompted Claude CUA [29] with $T=1$ following its format with an augmented system prompt as follows:

```

Additional Notes:
* Available xdotool keys: ctrl, command, option, backspace, tab,
  ↳ enter, esc, del, left, up, right, down, and single ASCII
  ↳ characters.
* When you think the task can not be done, say ```FAIL``` , don't
  ↳ easily say ```FAIL``` , try your best to do the task. When
  ↳ you think the task is completed, say ```DONE``` . Include
  ↳ the three backticks. If the task is not completed, don't
  ↳ raise any of these two flags.
* At the end of each step (except for the last step), always take
  ↳ a screenshot. In the next round, carefully evaluate if
  ↳ you have achieved the right outcome. Explicitly show your
  ↳ thinking: "I have evaluated step X..." If not correct, try
  ↳ again. Only when you confirm a step was executed
  ↳ correctly should you move on to the next one.
* You may need my username and password. My username is `ec2-user`
  ↳ and password is `000000`.

```

We follow Claude CUA's default interaction format, including its default behavior of filtering to the $n = 3$ most recent images. The main difference from other agents is that Claude receives a screenshot only upon explicitly requesting it. Typically, Claude CUA uses one round of dialogue to request a screenshot, performs an action in the next round, and then requests another screenshot in the following round, continuing this pattern. As a result, it requires approximately 30 dialogue rounds for Claude CUA to complete 15 steps and reach the screenshot budget limit.

G.6 UI-TARS

UI-TARS was configured at $T=1$ and $\text{top_p}=0.9$ during generation. We used the following prompt structure for UI-TARS, where the user query only contains the screenshot of that round of conversation:

```

<System Prompt> (including task instruction)
<User Query at t-n> (screenshot only)
<Agent Response at t-n>
<User Query at t-n+1>
<Agent Response at t-n+1>
<User Query at t>

```

Here, filtering to the $n = 3$ most recent screenshots is implemented by directly removing earlier dialog histories, while always keeping the system prompt. The system prompt is given by:

```
You are a GUI agent. You are given a task and your action history
    ↪ , with screenshots. You need to perform the next action to
    ↪ complete the task.

## Output Format
```\nThought: ...
Action: ...\\n```

Action Space

click(start_box='<|box_start|>(x1,y1)<|box_end|>')
left_double(start_box='<|box_start|>(x1,y1)<|box_end|>')
right_single(start_box='<|box_start|>(x1,y1)<|box_end|>')
drag(start_box='<|box_start|>(x1,y1)<|box_end|>', end_box='<|
 ↪ box_start|>(x3,y3)<|box_end|>')
hotkey(key='')
type(content='') #If you want to submit your input, use "\\
\\" at the end of `content`.
scroll(start_box='<|box_start|>(x1,y1)<|box_end|>', direction='
 ↪ down or up or right or left')
wait() #Sleep for 5s and take a screenshot to check for any
 ↪ changes.
finished()
call_user() # Submit the task and call the user when the task is
 ↪ unsolvable, or when you need the user's help.

Note
- Use Chinese in `Thought` part.
- Summarize your next action (with its target element) in one
 ↪ sentence in `Thought` part.
- Available hotkeys: ctrl, command, option, backspace, tab, enter
 ↪ , esc, del, left, up, right, down, and all standalone
 ↪ ASCII characters

User Instruction
<Task Instruction>
```

## G.7 ShowUI

ShowUI [25] was prompted with no system prompt, with only a user query. In each round, the user query contains two components: a fixed instruction followed by detailed task instructions, the current screenshot, and action histories. All other configurations remain default.

The fixed instruction used is:

```
You are an assistant trained to navigate the macOS screen.
Given a task instruction, a screen observation, and an action
 ↪ history sequence,
output the next action and wait for the next observation.
Here is the action space:
1. CLICK: Click on an element, value is not applicable and the
 ↪ position [x,y] is required.
2. INPUT: Type a string into an element, value is a string to
 ↪ type and the position [x,y] is required.
3. HOVER: Hover on an element, value is not applicable and the
 ↪ position [x,y] is required.
4. ENTER: Enter operation, value and position are not applicable.
5. SCROLL: Scroll the screen, value is the direction to scroll
 ↪ and the position is not applicable.
6. ESC: ESCAPE operation, value and position are not applicable.
```

7. PRESS: Long click on an element, value is not applicable and  
⇒ the position [x,y] is required.

Format the action as a dictionary with the following keys:  
{'action': 'ACTION\_TYPE', 'value': 'element', 'position': [x,y]}

If value or position is not applicable, set it as None.

Position might be [[x1,y1], [x2,y2]] if the action requires a  
⇒ start and end position.

Position represents the relative coordinates on the screenshot  
⇒ and should be scaled to a range of 0-1.

This fixed instruction is followed by the following contents during prompting:

Task: <Task Instruction>  
<Action History>  
<Current Screenshot>

Here within this instruction, following [7], the action history is a concatenation of ShowUI's previous round of outputs.

## 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: The three main gaps bridged by this paper are (1) an interactive macOS benchmark for GUI agents, (2) the multilingual design, and (3) the involvement of a safety subset. All have been reflected in the abstract.

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 main limitation is that the current evaluation relies on binary rewarding. This is explained at the end paragraph of the paper.

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: No theoretical results are provided in the paper.

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: The implementation details of the agents are provided thoroughly in the appendices. With those agent configurations, the testbench algorithms provided in the main paper, and the AWS-hosted environments, the experiments could be re-implemented.

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: Please refer to <https://macos-world.github.io/>.

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: Details are specified in Section 5.1 and Appendix G.

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: [No]

Justification: [NA]

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: As in Section 5.1.

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: The research conducted in the paper conform with the 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: [Yes]

Justification: As discussed in Appendix A and B.

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: [\[Yes\]](#)

Justification: As included in Appendix B.

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: Our benchmark involves a licensed way of running macOS environments. By running in AWS-hosted dedicated Mac Mini machines, the benchmark complies with Apple Software's EULA, and both AWS and Apple could be benefited.

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: [Yes]

Justification: [NA]

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: [NA]

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: [NA]

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: LLM is used only for writing, editing or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research.

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.