

OSReward: Instituting Standardized Evaluation for Cross-Platform Computer-Use Reward Models

Qiushi Sun^{1*†}, Kanzhi Cheng^{2*†}, Yian Wang^{1,3*}, Bowen Yang^{4*}, Hang Yan^{5*}, Liheng Chen^{6*}, Fangzhi Xu⁵, Zichen Ding, Nuo Chen², Jialin Cao², Xingdong Gong², Zehao Li⁴, Kaiming Jin³, Xinfeng Yuan⁷, Zhoumianze Liu⁷, Jingyang Gong¹, Zhangyue Yin⁷, Jiahui Gao¹, Zhiyong Wu¹, Tianbao Xie¹, Jianbing Zhang^{2✉}, Ben Kao^{1✉}, Lingpeng Kong¹

¹The University of Hong Kong, ²Nanjing University, ³National University of Singapore, ⁴University of Science and Technology of China, ⁵Xi'an Jiaotong University, ⁶University of Oxford, ⁷Fudan University

Computer-using agents (CUAs) are advancing rapidly across the digital world. A CUA trajectory records the agent’s actions, states, and reasoning. Verifying whether it fulfilled the task instruction is central to CUA evaluation, data curation, and reinforcement learning. Neither human-written verifiers nor human annotators can provide such verification at scale, so the field increasingly turns to vision-language models (VLMs) as judges of CUA trajectories. But a fundamental question has long gone unexamined: are these VLM judges reliable enough? To study it systematically, we introduce OSReward, a realistic, high-quality benchmark that evaluates VLM judges on CUA trajectories. The trajectories come from diverse agent backbones executing human-verified instructions across platforms, then rigorously labeled with ground-truth verdicts through multi-stage human annotation. Building on it, we derive OSReward-Hard, a challenge set concentrating genuinely hard cases, and OSReward-Multi for fine-grained efficiency and alignment scoring. The most comprehensive evaluation of VLM judges to date finds even state-of-the-art models fall short of an ideal judge, sharing a systematic leniency bias that mislabels failed runs as successes. The few reliable enough to trust are too expensive to run at scale, while affordable open models trail far behind. To close this gap, we construct and release OS-Shepherd-100K, an open corpus of reasoning-annotated trajectory judgments for the CUA community. On it, we train OS-Shepherd (9B and 35B), open reward models that supply low-cost, stable, and reliable reward signals, matching commercial judges at 30–60× lower cost than the frontier. Extensive analyses further inform the design of reliable CUA reward at scale. Our code, benchmark, dataset, and model checkpoints are available at [OSReward Homepage](#).

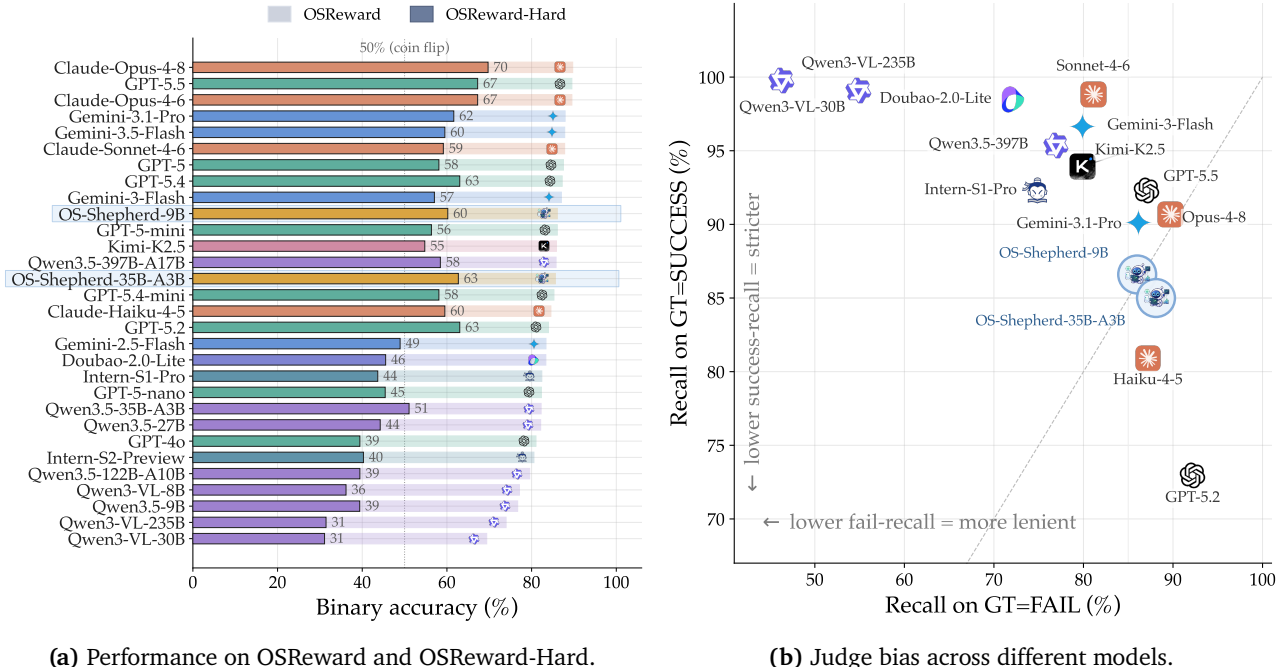


Figure 1: How well VLM judges score CUA trajectories (a), and the strict-lenient bias they share (b).

1. Introduction

Computer-using agents (CUAs) are advancing rapidly across the digital world, operating the web, mobile apps, and desktop software (Xie et al., 2024; Zhou et al., 2024; Rawles et al., 2025). A *trajectory* is the interleaved sequence of the agent’s actions, states, and reasoning. Scaling that progress, through evaluation, data curation, and reinforcement learning (Xue et al., 2026b), requires verifying whether each trajectory fulfilled its task instruction. However, human-written verifiers cover only a handful of curated tasks, and they do not apply at all to the static corpora and previously collected trajectories that have no live environment left to inspect. Human annotation cannot keep pace either. Typical methods therefore revolve around a VLM acting as the judge (Zheng et al., 2023a; Zhuge et al., 2025), whether as reward models (Pan et al., 2024; Li et al., 2026; Chae et al., 2025) or autoraters (Sun et al., 2025a), an approach fast becoming the de-facto practice.

Yet whether this judge is actually reliable remains largely untested. Unlike judging text (Zheng et al., 2023a; Zhou et al., 2025) or general multimodal agents (Chen et al., 2024), judging a CUA trajectory means reading a long, interleaved record of states, actions, and reasoning, then deciding whether the environment truly reached the instructed goal rather than whether the agent merely claims so, a verdict that can be reached from a fragment of that record. No prior study measures how reliably models make this judgment across platforms, and our pilot study shows the problem is real: even the best VLM judges disagree with existing benchmarks’ own verifiers on roughly a quarter of desktop verdicts. Our generalization study (§7) measures this disagreement at scale and discusses where it originates, in the judges or in the verifiers and tasks themselves.

To study this systematically, we build OSReward, a benchmark for evaluating the judge itself. Reusing existing agent benchmarks’ off-the-shelf trajectories (Lin et al., 2025; Li et al., 2026) would confound judge errors with flaws in the runs themselves, leaving failures unattributable to the judge. OSReward is instead built on dedicated cross-platform data infrastructure that we design and operate end to end: stock environments on web, mobile, Ubuntu, and Windows, extended with the common (Xie et al., 2024) and professional (Sun et al., 2026b) applications a real user has and realistic starting states. This affords coverage far beyond reused trajectories: more applications, websites, and scenarios, spanning both pure-GUI and GUI+CLI workflows. On these environments, annotators curate verified, environment-grounded instructions. Agents from four model families then execute them, and their varying capability yields both real successes and real failures. Each trajectory then passes rigorous multi-stage human labeling with strict screening. We obtain 1019 human-gold trajectories, long-horizon (up to 100 steps), organized into three views: the full set carries the breadth; OSReward-Hard concentrates genuinely hard cases, built from the trajectories the annotators themselves split on, where current judges commonly fail, exposing error modes and remaining headroom; and OSReward-Multi layers fine-grained efficiency and alignment labels on top of the binary verdicts. Grounded in careful human annotation, this labeled set offers a clean, standardized basis for studying VLM judges.

We then run the broadest evaluation of VLM judges yet, benchmarking 27 models. Frontier judges look adequate on the full set but fall sharply on OSReward-Hard, where the best judge drops below 70% and the mean judge falls to 52%. The errors share one signature: judges are easily fooled by *false successes*, where the agent declares the task completed but has actually failed, accepting such runs far more often than the reverse. The few judges accurate enough to trust, such as Claude-Opus-4-8 (Anthropic, 2026b) and GPT-5.5 (OpenAI, 2026b), are far too expensive for the millions of judgments that evaluation, curation, and training demand, while affordable open judges trail by a wide margin, leaving no judge both reliable and affordable.

We therefore close this gap with open data and open reward models. We run our own large-scale

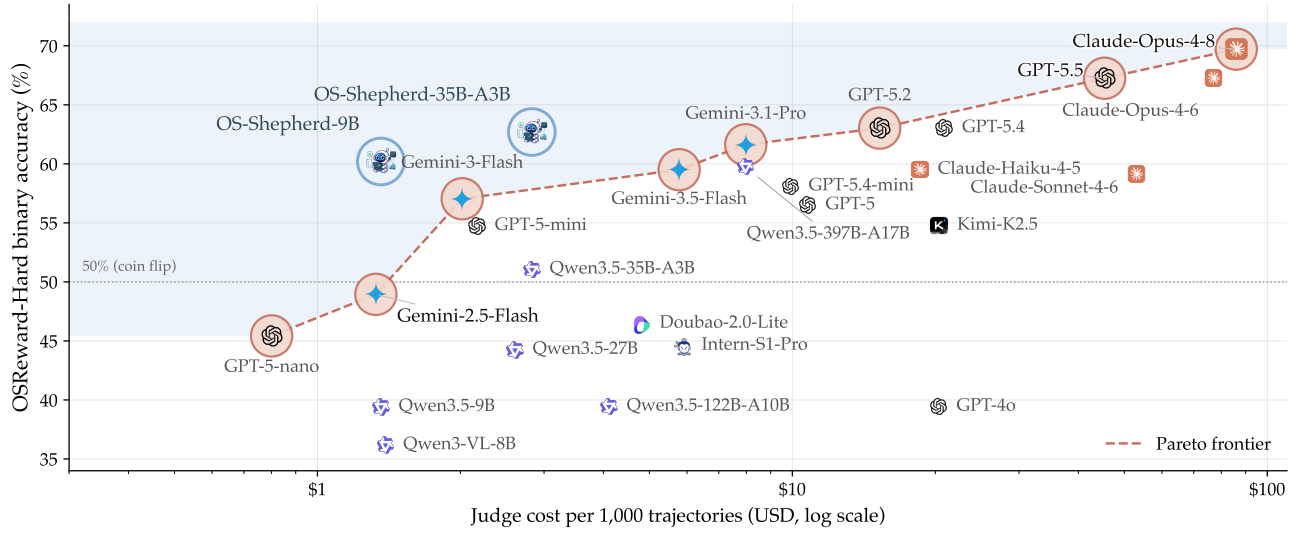


Figure 2: Cost against binary accuracy on OSReward-Hard: reliable judges are expensive, and the OS-Shepherd models come nearest their accuracy at a fraction of the cost.

collection pipeline, process and screen its output, and curate it together with filtered public corpora into OS-Shepherd-100K, an open corpus of 100K reasoning-annotated trajectory judgments that spans more trajectories, agent backbones, scenarios, and task types than reused corpora provide. Guided by what our evaluation reveals, we label these trajectories without new human annotation, each labeling choice set by a finding of the judge study. We release the corpus in full, with failure types annotated and analyzed, so the community can push open CUA reward research forward.

Based on OS-Shepherd-100K, we train OS-Shepherd-9B and OS-Shepherd-35B with a two-stage recipe that first builds accurate judging and then directly targets the false success, the sharpest failure the study exposes. The outcome is best seen in Fig. 2: on the OSReward-Hard cost–accuracy frontier, where affordable judges collapse and reliable ones are the most expensive, the OS-Shepherd models are the low-cost points nearest the frontier, matching commercial judges at 30–60 $\times$ lower cost than the frontier. In effect, a low-cost, stable, and reliable reward signal at training scale is now within an academic budget. Further analyses, spanning input ablations, robustness, and ensembling, reveal when and why judges fail, and held-out evaluations confirm the de-biasing transfers to unseen benchmarks.

Contributions.

- **OSReward.** A standardized benchmark for the CUA reward signal, built from scratch on dedicated cross-platform data infrastructure: human-gold trajectories across four platforms, plus the OSReward-Hard challenge set for hard-case diagnosis and the OSReward-Multi subset for fine-grained judging.
- **Comprehensive Evaluation.** The broadest study of VLM judges for CUA yet, spanning platforms and difficulty: judges share one lenient failure mode, read the agent’s narrative more than the screen, and collapse on hard cases, while the reliable ones cost too much to run at scale. Extensive analyses isolate what drives a verdict.
- **OS-Shepherd Data & Models.** An open, reasoning-annotated training corpus (OS-Shepherd-100K) curated from more than 300K judge instances on cross-platform trajectories into nearly 100K samples, rich in both successes and failures, and the two OS-Shepherd reward models (9B and 35B): open, low-cost, and self-hostable.

2. Related Work

2.1. Computer-Using Agents

Computer-using agents perceive digital environments and take actions (e.g., clicks, typing, or CLI commands) to complete user tasks (OpenAI, 2025a; Wu et al., 2024; Sun et al., 2024). Driven by vision-language models, the field has advanced quickly from grounding instructions on raw screens (Zheng et al., 2024; Cheng et al., 2024; Gou et al., 2025; Wu et al., 2025a) to native GUI action models that operate real applications end to end (Wu et al., 2025b; Xu et al., 2025b; Qin et al., 2025), joined by open cross-platform agent data (Wang et al., 2025; Liu et al., 2026; Xu et al., 2025a) and systems that compose heterogeneous agents (Jia et al., 2025). Standardized environments across web, mobile, and desktop (Chen et al., 2025b; Zhou et al., 2024; Kong et al., 2026) give this progress reproducible tasks and testbeds. That reproducibility rests on hand-written verifiers, one per task, and many tasks in real use have no checkable outcome at all. Evaluation and training in these environments are outgrowing those verifiers: trajectory data is collected and processed at scale (Qin et al., 2025; Xue et al., 2026b), policies are optimized with reinforcement learning over ever-longer, open-ended interactions (Bai et al., 2024; Qi et al., 2024; Sun et al., 2025b; Xu et al., 2026), and agents are expected to eventually self-evolve on synthetic experience (Zhang et al., 2026; Xue et al., 2026a; Jiang et al., 2026). Every stage of this pipeline consumes a reward signal, and an ideal judge should be general and stable at the trajectory level. We build OSReward from scratch to study this problem systematically.

2.2. Judging CUA Trajectories

Evaluating and training CUAs both require deciding whether a trajectory fulfilled its instruction, at a scale that, as the introduction argues, neither human-written verifiers nor human annotation can reach, leaving model-based judging as the only practical route. Judging a CUA trajectory, however, is much more complex than judging text (Zheng et al., 2023a; Cobbe et al., 2021; Wang et al., 2024) or general multimodal responses (Chen et al., 2024; Li et al., 2025): the verdict must track a long, interleaved record and sometimes even infer whether the environment after execution truly fulfills the task. An emerging line of work nevertheless puts a VLM judge in exactly this role, as the reward model behind trajectory synthesis and data filtering (Sun et al., 2025a; Wang et al., 2026; Zhang et al., 2025) and as the basis of trained reward models and critic-style autoraters (Chen et al., 2025a; Xiao et al., 2025; Li et al., 2026; Chae et al., 2025; Wu et al., 2026). Although recent work has gradually become aware of the reliability issue in model-based reward (Lin et al., 2025; Lù et al., 2025), these early efforts stay on single, isolated platforms and reuse off-the-shelf instructions or trajectories from common benchmarks. Platforms differ in what a judge must read (Xue et al., 2025; Sun et al., 2026a): action spaces, the applications and websites in play, and run lengths from around a dozen steps on mobile to a hundred on desktop. Reused trajectories bring problems of their own: they inherit whatever rollout setting produced them, from agent backbone to budget; their gold labels may also come from imperfect verifiers (Xie et al., 2025); and some of their instructions are ambiguous enough that no determinate verdict exists. Reliability has thus been gauged on evidence too narrow and too noisy for the long-horizon runs a judge meets in practice, while the judges fit for real use stay too expensive at scale. In this work, we take on both gaps: OSReward is, to our knowledge, the first to measure VLM-judge reliability across platforms on fresh, human-gold trajectories, and OS-Shepherd turns what it reveals into an open, low-cost reward model.

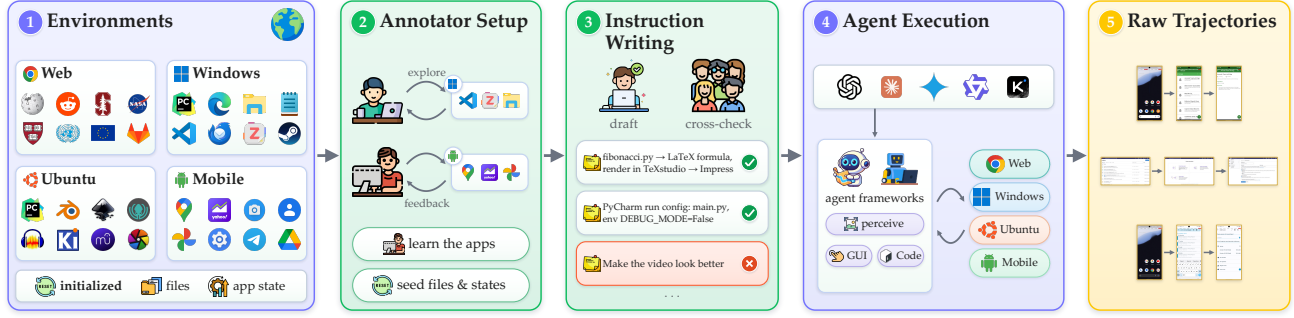


Figure 3: From realistic environments to raw trajectories: annotators prepare the environments and write grounded instructions on them, agents from four model families execute the instructions.

3. OSReward

Our goal is a corpus of realistic, cross-platform CUA trajectories paired with gold verdicts trustworthy enough to measure the judges themselves. Reusing existing benchmarks’ rollouts cannot supply it: their runs carry quality confounds, and their labels inherit the verifiers’ own noise. We therefore build the data ourselves. This section covers the infrastructure the collection runs on, the task instructions that drive it, the collection and annotation that produce the gold set, and the resulting benchmark with its three views; the judge protocol and metrics are set up with the experiments. The OS-Shepherd reward model of §6 is trained on a separate corpus, labeled without human annotation by applying the findings of the judge study, and fully disjoint from this benchmark.

3.1. Data Infrastructure

We collect fresh trajectories on cross-platform infrastructure we build and operate end-to-end, and carefully label them with humans.

Environment Preparation. The benchmark is only as good as the environments behind it. On desktop and mobile, we go well beyond stock benchmark images, equipping each machine with the everyday applications a real user has and initializing every task richly: user profiles, real files to edit, seeded app databases, distractor content. This is what makes a trajectory worth judging: realistic state to fail from, success that must change the environment rather than narrate completion, and concrete state for a machine-checkable verifier to inspect. Web tasks instead run on live websites, whose realism no snapshot reproduces and which need no initialization. Collecting everything on one infrastructure also keeps the output uniform: judge scores stay comparable across platforms, and a reward model sees one consistent trajectory format. We summarize each platform below; and give the full implementation, software coverage, statistics, and the action spaces of the executing agents in §A.

Web. We leverage a pool of Chromium-based browser workers to host many isolated sessions in parallel, one per web rollout. To mitigate blocking by websites that resist automation, the collection stack combines browser-level hardening with per-domain pacing and screens out low-quality runs afterward, supporting reliable collection at scale. Observation and action are purely vision-driven, and rollouts split between pure-GUI actions and GUI plus browser primitives, so we expect the reward model to score visual agents and API-mixing agents alike. To collect data for tasks requiring user logins, we also use several self-hosted local mirror websites (§A.1).

Windows. Our Windows environment is built out into a lived-in machine: some twenty everyday applications spanning IDEs, media editors, 3D and database tools, several with accounts already logged in, plus command-line utilities such as `ffmpeg`. Tasks here can run long, with rollouts stretching to a hundred steps. Additionally, to enhance sample diversity, we switch between 2K and 4K resolutions during the data collection process (§A.2).

Ubuntu. The Ubuntu machine comes fully initialized: about thirty applications, daily and professional, well beyond what typical environments carry; a typed pool of real files (some twenty types, each stocked with hundreds of files drawn from established datasets) so any task can start from a rich, realistic state; and per-application initialization that widens the space agents can explore. With Python and shell tooling preinstalled, rollouts split between a pure-GUI and a GUI+CLI action space, making Ubuntu the platform where visual and code-mixed agent styles meet (§A.3).

Mobile. The mobile environment is hosted in an Android emulator, and within our collection environment, we initialize the files, usage records, and databases of every application. The mobile environment is hosted in an Android emulator. Within our collection environment, we initialize the files, usage records, and databases of every application. This initialization includes randomly seeded records, browsing history, and photos, along with distractor content (noise events, decoy messages, look-alike files) to expand the exploration space further, and login-required applications come signed in with synced content (§A.4).

3.2. Task Instructions

Writing instructions worth judging is a problem in its own right. An instruction must be grounded in what the environment actually offers and answerable there, and the pool deliberately includes open-ended tasks beyond what any rule-based verifier can check, because those are exactly the trajectories a judge faces in real use. Annotators write the instructions after exploring the infrastructure described above, initializing bare environments as they go (Fig. 3); they may brainstorm with AI tools or software manuals, but every draft, whatever its origin, passes a human validity check. Each candidate is then screened by annotators other than its author, a peer cross-check that removes ambiguous, ungrounded, or unanswerable instructions before any rollout effort is spent. This pipeline has no exceptions: the long cross-application tasks, the hardest to specify because one instruction must stay grounded in several applications at once, are authored and screened like every other candidate rather than synthesized. Roughly 1500 candidates are authored, and about 800 survive into next stage.

The same infrastructure also supports collection far beyond the benchmark’s data scale, extending the instruction pool to training size, the route §6 takes. Only the human-vetted, peer-screened set enters the benchmark: an evaluation item must be individually trustworthy, a bar training data need not meet. §A details each platform’s instruction method.

3.3. Trajectory Collection and Gold Annotation

A trajectory is a task instruction followed by a sequence of steps, whose state (screenshot) is paired with the agent’s thought and action. OSReward is built by a pipeline that runs from the verified instruction set of §3.2, through execution with different agent frameworks and backbone models (Fig. 3), to human-labeled gold verdicts (Fig. 4).

Trajectories from Diverse Agent Backbones. A judge should be measured on the trajectories it will actually face, and those come from many different agents, not one. Each surviving instruction is therefore rolled out by the executing agents, each driven by one of several mainstream model backbones. Every

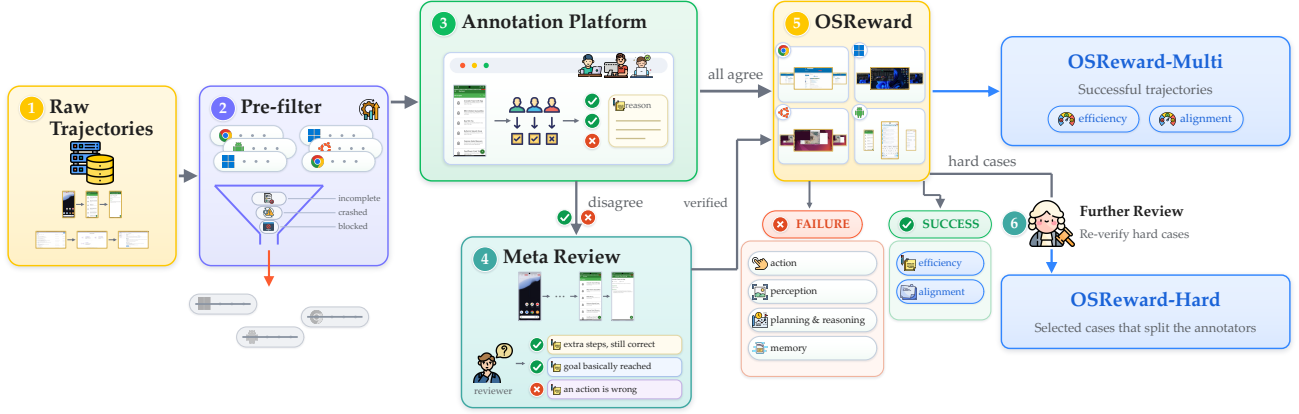


Figure 4: The annotation pipeline. Each pre-filtered trajectory is labeled by three independent annotators; disagreements go to meta review, and the verified gold set is read as three views

instruction is run by one to three of them, spanning the Claude, Gemini, Kimi, and Qwen families. Backbones differ in action idiom, thought verbosity, and failure modes, so spreading the rollouts across the pool keeps the benchmark from reducing to one agent’s interface style. The capability spread also yields real successes and real failures alike, and with them the labeled failures that judge evaluation otherwise lacks. An automatic pre-filter then discards runs with severe collection problems (persistent anti-bot blocks, network failures, frozen executions) before any human effort is spent, so a FAIL in the benchmark reflects the agent failing the task, not the environment failing the agent.

Human Annotation. Gold labels bound what the study can conclude, so every surviving trajectory is labeled by human effort. To facilitate the annotation process, we built an in-house annotation website enabling all annotators to easily review the entire trajectory. Before judging, an annotator reads its full multimodal context, every screenshot, thought, and action, not only the final screen. One deliberately strict standard: an answer the agent did not obtain or verify through the environment is a FAIL even when it happens to be correct (the judging prompt carries the same rule, §C.2). A trajectory judged SUCCESS is then scored on the two OSReward-Multi axes, *alignment* (how well the actions matched the task intent; Sun et al., 2026a) and *efficiency* (how free the run was of wasted steps), each on a 3-class scale (0, 0.5, 1; scoring rubric in §B.4). A trajectory judged FAIL is instead tagged with the errors that caused it, one or more categories from a failure taxonomy spanning reasoning-and-planning, action, perception, and memory errors (details in §B.3).

Consensus and Meta-Review. Every trajectory that passes the pre-filter is labeled independently by three annotators (agreement statistics in §B.2). Where they agree unanimously, the verdict is final; the trajectories they split on escalate to a *meta-review* in which two senior reviewers examine the trajectory together and issue the final judgment (deliberation, not a majority vote), so no gold verdict rests on a single reader. The pass also discards trajectories with residual quality problems rather than force a label; the full funnel, with counts at every stage, is in §B.2. All told, three-way annotation, meta-review, and the hard-set re-verification described below cost roughly 800 human hours.

3.4. OSReward and Its Variants

The pipeline yields one gold set that read three ways: the full set carries the breadth across tasks and platforms, OSReward-Hard isolates the hard cases, and OSReward-Multi adds rating granularity; the

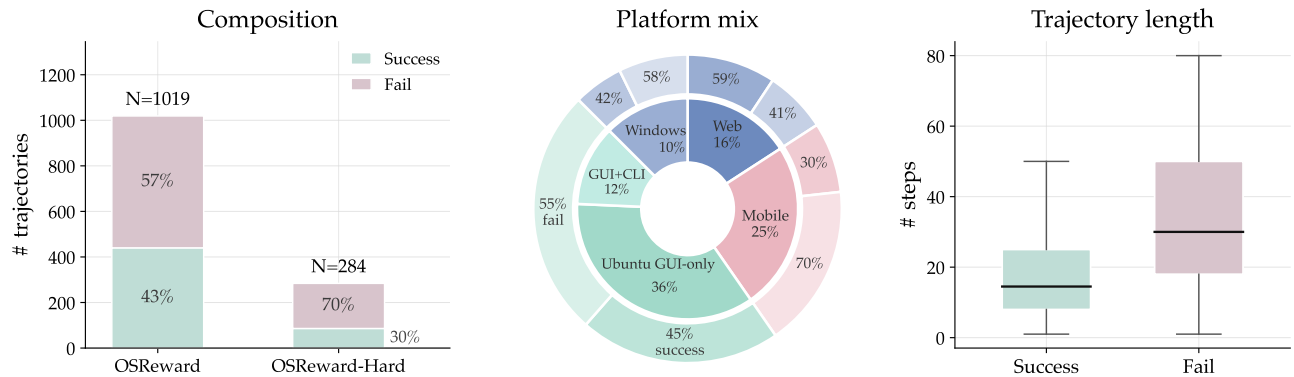


Figure 5: OSReward at a glance: outcome composition of the full and Hard sets, platform mix, and trajectory length; failed runs are markedly longer.

latter two are nested subsets of the full set.

The Full Set (OSReward). The full OSReward set holds 1019 trajectories spanning all four platforms, roughly balanced between successes (43%) and failures (57%) (Fig. 5). It covers real applications and live websites, instructions from routine chores to professional workflows, and trajectories from four agent families across GUI-only and GUI+CLI action spaces, from dozen-step tasks to 100-step runs. Reasoning-and-planning errors dominate the failures, tagged on 86% of failed runs (multi-label profile in Fig. 13; taxonomy in §B.3). Failed runs are also markedly longer than successful ones, and length itself makes verification harder.

OSReward-Hard. A challenge subset of 284 trajectories concentrating genuinely hard cases, drawn mostly from the trajectories the annotators themselves split on. Every candidate is re-verified under a further review process, especially for the fail cases. Their records read like those of completed tasks or successful tasks completed through convoluted paths, and both human annotators and judge models often fail or diverge on them. Beyond being harder than the full set, the selection deliberately raises the share of failures (a 30/70 SUCCESS/FAIL split) to further expose how easily existing judges are fooled by false successes. Its role is diagnostic: these runs briefly fooled even trained annotators, and together they expose the error modes judges share and how much room remains above the best current judges.

OSReward-Multi. Beyond the binary verdict, the 440 successful trajectories carry human alignment and efficiency sub-labels (failed trajectories have neither). This is a secondary axis: it grades how well a successful run did its job, which no existing CUA-judge benchmark measures. We release it as an exploratory track and report it alongside the binary verdict throughout. Alignment enters the benchmark with two levels (0.5 and 1.0); efficiency has three (0, 0.5, and 1.0), with detailed definitions in with detailed definitions in §B.4. On both axes, we report macro-recall, not a raw average, so the skew cannot inflate the score (§4.1).

4. Benchmarking VLMs as Judges

We benchmark 27 VLMs as judges on OSReward. Under one fixed setting we score every model, then probe where the reward signal holds and where it breaks through targeted analyses and ablations.

4.1. Experimental Setup

Model Selection. We evaluate 27 VLM judges across model series from OpenAI (Hurst et al., 2024; OpenAI, 2025b, 2026b,a), Anthropic (Anthropic, 2026b,a,c, 2025), Gemini (Comanici et al., 2025; Gemini Team, 2025), Qwen (Bai et al., 2025; Qwen Team, 2026), Doubao (ByteDance Seed Team, 2026), Kimi (Kimi Team et al., 2026), and Intern (Zou et al., 2026). The roster covers closed frontier judges, closed efficient tiers, large open-weight models, and small open VL models; several models are additionally run in a thinking variant. Per-model details and configs are in §C.1.

Judging Protocol. Every judge runs under the common protocol: it reads the trajectory’s last N states with the per-step reasoning and action text and returns a verdict, with no task-specific harness, tool access, or step-level supervision, and $N = 5$ by default. The judge outputs a SUCCESS/FAIL verdict and, on OSReward-Multi trajectories, alignment and efficiency ratings. The remaining settings (frame count, the red click marker (Yang et al., 2023), greedy decoding) are held fixed here and perturbed one at a time in the analyses and ablations that follow, each defined where it is used.

Metrics. We report binary accuracy, decomposed into success recall (sRec, the share of truly successful trajectories accepted; low means too strict) and fail recall (fRec, the share of truly-failed trajectories caught; low means too lenient); balanced accuracy (BalAcc) is their mean, which the 43/57 class mix cannot inflate. The full judging prompt is given in §C.2.

4.2. The Accuracy Ceiling

The Performance Ceiling. For a reward model to support training, around 90% binary accuracy is a common working bar. On OSReward only the frontier approaches it (Table 1): Claude-Opus-4-8 comes closest at 89.7%, GPT-5.5 and Claude-Opus-4-6 sit just behind, and the field spans twenty points down to the small open VL models. Nor is the order at the top stable: the lead changes hands once the metric moves to balanced accuracy. The table identifies a top tier, not a best judge.

The Open–Closed Gap. Closed judges lead open-weight ones across almost the whole field, and the gap is widest against the small open VL models: the larger open-weight judges, e.g., Kimi-K2.5 and Qwen3.5-397B-A17B, narrow it to within 4 pp of the lead (a fuller comparison in §E.3).

Takeaway: Only frontier judges approach the accuracy a training-time reward needs, and even the largest open-weight counterparts still fall short of them, which leaves the field without a judge that is both reliable and open.

4.3. The Leniency Bias

The Strict–Lenient Plane. Binary accuracy is intuitive but hides *how* a judge errs. Here we plot each judge at its (fRec, sRec) to make the direction of error visible (Fig. 6), and two clusters emerge. A large *lenient* cluster has high sRec but low fRec, accepting almost everything including failures, while a smaller *strict* one (GPT-5.2, Claude-Haiku) trades success recall to catch more failures. The top judges sit near the diagonal (sRec = fRec), balanced rather than extreme, but they are few, and the field as a whole skews lenient. The same pattern holds on OSReward-Hard, more sharply (§4.4).

The Dominant Error. Beyond how often a judge is wrong, the more actionable question is what it gets wrong. We sort every incorrect verdict into *over-accepts* and *over-rejects*, each with three finer modes, using a strong VLM labeler with human re-checks (Fig. 7). One mode dominates. Over-accepting an

Table 1: Main-setting results for the reference judges and OS-Shepherd on OSReward and OSReward-Hard along with their access status, sorted by full-set accuracy.

Judge	Access	OSReward				OSReward-Hard			
		Acc	sRec	fRec	BalAcc	Acc	sRec	fRec	BalAcc
🌟 Claude-Opus-4-8	closed	89.7	91.1	88.9	90.0	69.7	69.8	69.7	69.7
🌀 GPT-5.5	closed	89.5	91.8	87.8	89.8	67.3	66.3	67.7	67.0
🌟 Claude-Opus-4-6	closed	89.5	92.7	87.7	90.2	67.3	72.1	65.2	68.6
🔹 Gemini-3.1-Pro	closed	87.9	90.2	86.2	88.2	61.6	61.6	61.6	61.6
🔹 Gemini-3.5-Flash	closed	87.8	95.7	81.8	88.8	59.5	81.4	50.0	65.7
🌟 Claude-Sonnet-4-6	closed	87.7	97.5	80.3	88.9	59.2	90.7	45.5	68.1
🌀 GPT-5	closed	87.4	86.8	87.9	87.4	58.1	43.0	64.6	53.8
🌀 GPT-5.4	closed	87.1	87.3	87.0	87.1	63.0	62.8	63.1	63.0
🔹 Gemini-3-Flash	closed	87.0	96.6	79.8	88.2	57.0	86.0	44.4	65.2
🌀 GPT-5-mini	closed	86.1	93.8	80.2	87.0	56.3	79.1	46.5	62.8
🖥️ Kimi-K2.5	open weights	85.9	95.5	79.2	87.3	54.8	83.7	42.1	62.9
🌀 Qwen3.5-397B-A17B	open weights	85.8	95.2	78.6	86.9	58.5	91.9	43.9	67.9
🌀 GPT-5.4-mini	closed	85.2	82.5	87.2	84.9	58.1	48.2	62.4	55.3
🌟 Claude-Haiku-4-5	closed	84.5	80.9	87.2	84.0	59.5	47.7	64.6	56.2
🌀 GPT-5.2	closed	83.9	73.0	92.2	82.6	63.0	30.2	77.3	53.8
🔹 Gemini-2.5-Flash	closed	83.3	95.5	74.0	84.8	48.9	90.7	30.8	60.8
🌀 Doubao-2.0-Lite	closed	83.3	98.5	72.1	85.3	45.5	96.1	24.3	60.2
🌀 GPT-5-nano	closed	82.3	97.0	71.1	84.1	45.4	95.3	23.7	59.5
🖥️ Intern-S1-Pro	open weights	82.3	92.3	74.7	83.5	43.7	70.9	31.8	51.4
🌀 Qwen3.5-35B-A3B	open weights	82.2	92.4	74.5	83.5	51.1	83.7	36.9	60.3
🌀 Qwen3.5-27B	open weights	82.0	97.4	70.5	84.0	44.2	92.9	23.2	58.0
🌀 GPT-4o	closed	81.0	96.8	69.0	82.9	39.4	90.7	17.2	53.9
🖥️ Intern-S2-Preview	open weights	80.6	98.4	66.9	82.7	40.3	94.2	16.8	55.5
🌀 Qwen3.5-122B-A10B	open weights	79.6	96.8	66.4	81.6	39.4	89.5	17.7	53.6
🌀 Qwen3-VL-8B	open weights	77.1	99.8	59.9	79.8	36.2	100.0	8.2	54.1
🌀 Qwen3-VL-235B	open weights	74.0	99.1	54.9	77.0	31.4	97.7	2.5	50.1
🌀 Qwen3-VL-30B	open weights	69.4	99.8	46.3	73.0	31.1	98.8	1.5	50.2
🖥️ OS-Shepherd-9B (ours)	open weights + data	86.1	86.6	86.0	86.3	60.2	66.3	57.6	61.9
🖥️ OS-Shepherd-35B-A3B (ours)	open weights + data	85.6	85.0	86.2	85.6	62.7	68.6	60.1	64.3

incomplete task makes up two-thirds of all errors, and every family shares the bias, as this is the leading error mode of every single judge, at no less than 48% of its mistakes. Pooled across judges, over-accepts outnumber over-rejects three to one, though the ratio narrows to about two to one for the strongest judges. The mechanism becomes clear in §5.2, where verdicts turn out to lean on the agent’s own narrative far more than on the screenshots, so a failed run that closes with a success claim is exactly what judges miss. Since one mode is this dominant, prompting the judge to verify completion explicitly may help more than the ensembling we evaluate in §5.3; we leave this to future work.

Takeaway: Every judge family shares one dominant failure, accepting an incomplete task as a success, so de-biasing that lenient mode is the highest-value intervention.

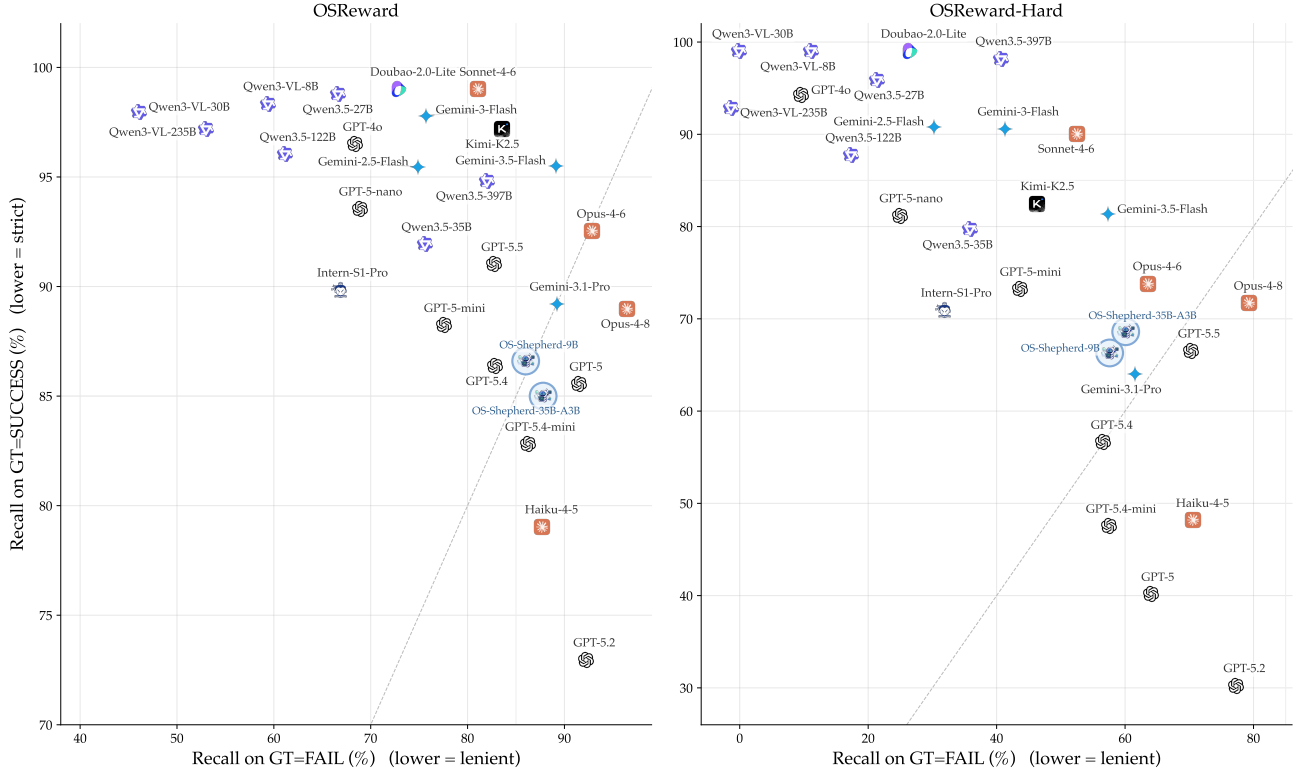


Figure 6: Judge bias on OSReward (left) and OSReward-Hard (right). Judges below the diagonal skew lenient; most of the field sits there, and the skew widens on the hard set.

4.4. OSReward-Hard: The Challenge Set

The Collapse. The aggregate $\sim 90\%$ is optimistic. OSReward-Hard, the failure-heavy challenge set of §3.4 (30/70 SUCCESS/FAIL; selection details in §E.1), drops every judge by 20–43 pp (Table 1; Fig. 1a). Even the best judge loses a full twenty points and lands at 69.7%, level with what a constant always-FAIL judge scores on this 30/70 split; the mean judge falls to 52%, and the lenient tail lower still. Raw accuracy is therefore a trap here: the hard-set columns of Table 1 carry the recalls and balanced accuracy, where such a constant predictor falls to 50%. The drop is not a uniform shift either: the broad ordering is preserved but stretched, and the mid-pack reshuffles, with judges like GPT-5.4 and GPT-5.2 climbing several ranks on the hard set. The collapse is also *structured* (Fig. 8): Windows is the hardest platform to judge and mobile the easiest, and the failure types that hinge on reading the screen (perception, then action) are markedly harder to catch than planning-and-reasoning failures, the dominant type in the data, which are legible in the thought and action text.

Many of these hard failures are *false successes*: the agent’s closing text claims a completion it never reached, exactly what a judge that leans on the narrative (§4.3) fails to catch.

Leniency Widens. The leniency bias does not just persist on OSReward-Hard, it *widens* (Fig. 6, right; Table 1). Fail recall now spreads from near zero to 77%: the lenient cluster (e.g., Qwen3-VL-30B, Intern-S1-Pro) accepts almost every hard failure, the best judges (Claude-Opus-4-8, GPT-5.5) stay balanced but only at $\sim 70\%$ recall on each side, and a few judges cross into the strict half.

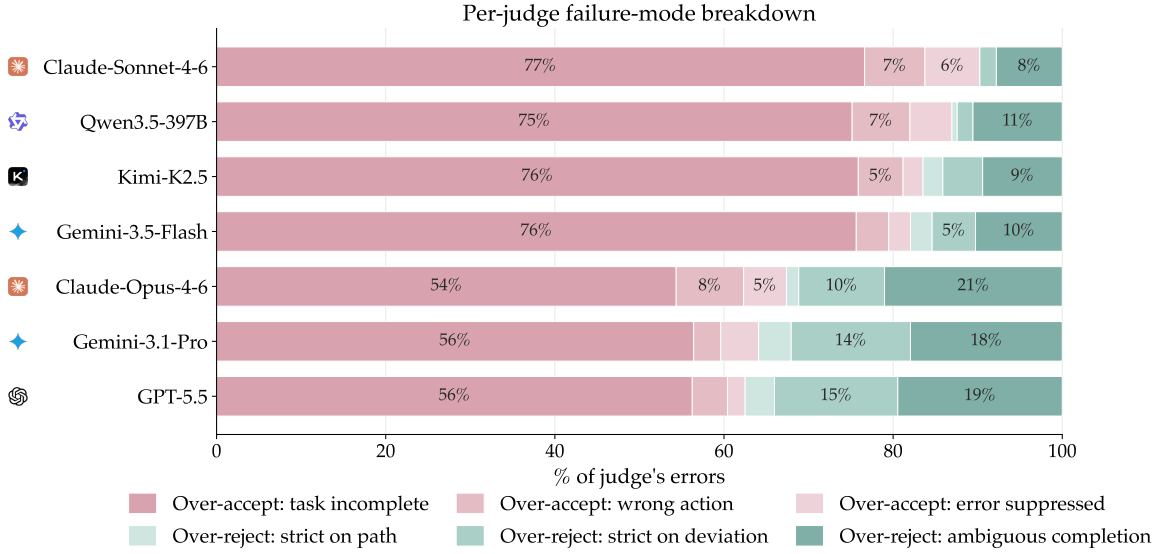


Figure 7: Per-judge error composition, seven representative judges: over-accepting an incomplete task (warm colors) dominates every family.

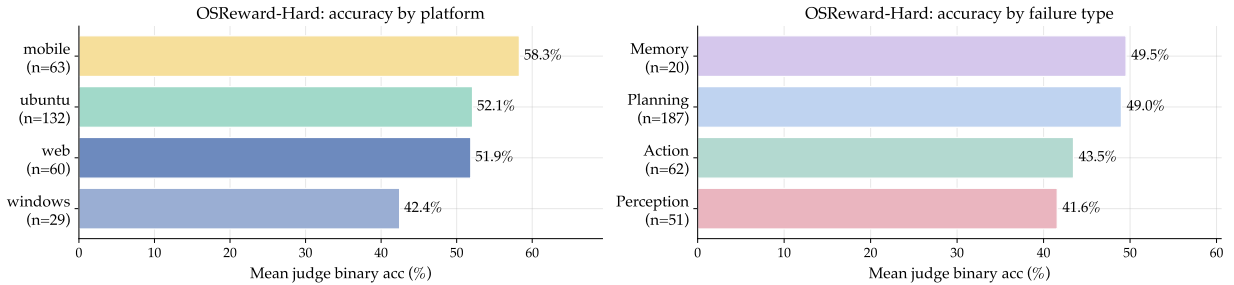


Figure 8: Mean per-judge binary accuracy on OSReward-Hard, by platform and by failure type. Failure types are multi-label, so their counts sum to more than the number of failed trajectories.

Observations from the Hard Set. OSReward-Hard concentrates on the cases the shared lenient mode of §4.3 misses, failed runs whose records read like completed tasks. On these deceptive cases, the leading judges finally separate, so we use the hard set as the primary diagnostic in the rest of the paper. Read together with the failure-type analysis above, it also locates the headroom, which lies in reading the screen and verifying completion rather than in high-level reasoning. On this harder footing only Claude-Opus-4-8 and GPT-5.5 remain both balanced and accurate, near 70% recall on each side, while the few other judges near the diagonal hold their balance only by giving up accuracy on both axes (Fig. 6). A judge that is accurate without being too strict or too lenient barely exists yet.

4.5. OSReward-Multi: Fine-Grained Grading

A binary verdict says *whether* a run succeeded; a reward signal is more useful when it can also grade how well it was done. On the OSReward-Multi trajectories, every judge additionally rates alignment and efficiency (§3.4); we score each axis two ways (Table 2): *macro-recall*, whether the emitted levels are usable as they stand, and a threshold-free *AUC*, whether the judge can rank them at all.

Three findings stand out. Quality grading is far weaker than outcome judging: even the best judge falls from $\sim 90\%$ on the binary verdict to the low sixties. The AUC–macro-recall gap is the diagnostic: judges rank the levels better than they score them, so the discrimination is there but the emitted thresholds are miscalibrated. And the weakness concentrates on *alignment*, where judges default to the top rating almost regardless of the run, because whether the actions served the instruction is harder to read off the record than whether steps were wasted. Our own OS-Shepherd models inherit this: the 35B recovers part of the gap while the 9B stays near the constant-level baseline. This keeps the binary verdict our primary metric and the multi-axis view a complementary probe. Full metric definitions and per-axis results are in §E.2.

Table 2: Strong judges on OSReward-Multi (%), sorted by AUC; best per column in bold.

Judge	Macro-recall			AUC
	Align	Effic	Multi	
GPT-5.5	58.7	68.2	63.5	66.7
Claude-Opus-4-8	52.9	68.7	60.8	65.6
Claude-Sonnet-4-6	53.2	62.6	57.9	61.9
Gemini-3.5-Flash	47.6	71.4	59.5	60.8
OS-Shepherd-35B-A3B (ours)	47.7	65.8	56.8	60.7
OS-Shepherd-9B (ours)	44.1	54.0	49.0	58.5
Gemini-3-Flash	50.6	61.5	56.0	55.8

Takeaway: No judge grades quality well: judges generally rank the alignment and efficiency levels better than they score them, so the discrimination exists but the emitted ratings are miscalibrated, worst on alignment.

5. Analysis

The benchmark measured how well judges perform; this section analyzes what drives a verdict. We perturb one component at a time and observe how the verdict changes, across three groups: the visual input, the text input, and the way the model is run. We close with the cost of reliability.

5.1. Visual Inputs Barely Move the Verdict

We first change what the judge *sees*. Swapping the five trailing screenshots for the last three, or for the first plus the last two, moves binary accuracy by less than half a point, and removing the red click marker does not hurt at all. Sweeping the trailing-screenshot count from one to sixteen is just as uneventful: each judge wanders two to three points with no trend, and every judge sits near its own best at $N = 5\text{--}9$ (§E.6). The visual settings that look important when writing a judging prompt turn out to matter little for aggregate accuracy.

Aggregate accuracy does not capture everything, though. Each of these settings, harmless in aggregate, still flips 5–7% of individual verdicts relative to the main setting. Such flips average out in evaluation, but not in reward labeling, where each trajectory’s label is consumed on its own. §5.3 gives the noise floor these numbers should be read against.

5.2. The Text History Matters

One ablation dominates the rest (Fig. 9). Dropping the per-step thought and action text costs 7.2 pp on average, and several times that on the web, where typed strings carry intent that the last few screenshots cannot encode. It also flips 22.7% of verdicts, three times what any visual change does. Dropping only the chain-of-thought while keeping the actions is far milder (1.8 pp, 11.6% flipped): within the text, the agent’s actions carry roughly more signal of its stated reasoning.

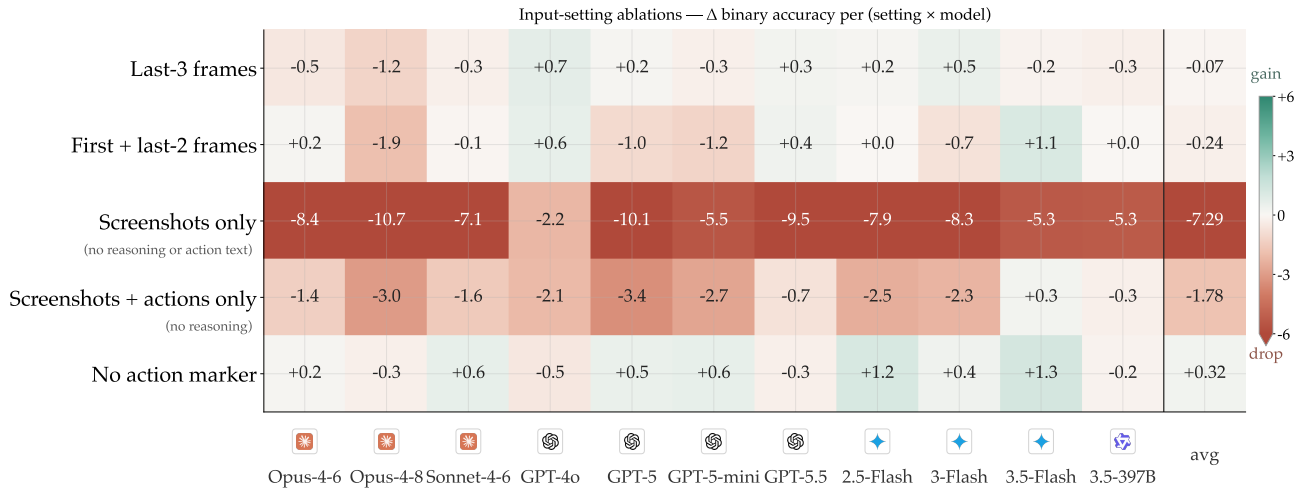


Figure 9: Overview of input ablations: Δ binary accuracy per (setting $\times$ model) vs. the main setting.

This is the mechanism behind the leniency bias of §4.3. A judge that leans on the agent’s own narrative is exactly a judge that a confident closing claim can fool. This suggests a recipe for CUA reward-model labeling: keep the full text history, drop the marker, and set the screenshot count per model.

Takeaway: Even with screenshots in view, the text history drives the verdict; the visual cues that look essential barely matter.

5.3. Running the Judge Differently Does Not Fix It

If the input is fixed, can the model be run differently to get a better verdict? Letting a judge think harder helps, but the gain shrinks as the judge gets stronger. Comparing each model against a stronger thinking setting of itself gains a few points on the weakest judge and almost nothing at the frontier; raising GPT-5.5’s reasoning effort yields monotone but similarly small improvements (§E.6). Extra deliberation recovers what a weak judge loses by under-thinking; it does not improve the best judges.

Decoding noise provides a reference scale for the flip rates. Re-running the *same* judge on the *same* input at $T=0.7$ already flips 6–9% of verdicts, so the churn from the visual settings of §5.1 sits at or below the model’s own sampling noise, while the text ablation’s is far outside it. Aggregate accuracy is stable under temperature; individual labels are not, and that is the quantity a reward model consumes.

Nor does adding judges help. Judges agree far more than they disagree (pairwise Cohen’s $\kappa \approx 0.71$ among top judges; Cohen, 1960) and herd on the *same* hard trajectories, so a wider pool mostly adds copies of the same mistakes: a top-3 majority vote edges the best single judge by about a point at several times the cost. The oracle is more informative: accepting any pooled judge’s correct verdict reaches 99.2%. The pool almost always contains a correct verdict, but a vote cannot identify which judge to trust on which trajectory. This is the argument for soft-label, confidence-weighted reward modeling over hard votes, a direction we leave to future work; our own corpus uses agreement to select which trajectories to keep rather than to sharpen any verdict (§6.1). Full per-model ablations and further detail are in §E.

Takeaway: No model-side knob buys reliability: extra thinking helps least exactly where judges are already strong, resampling alone flips 6–9% of labels, and judges herd, so ensembling cannot

substitute for a better judge.

5.4. The Cost of a Reliable Verdict

The remaining option is to pay for a stronger judge, and here cost becomes the binding constraint. Accuracy and cost are in direct tension, sharpest on OSReward-Hard (Fig. 2): Claude-Opus-4-8 holds 69.7% there at ~\$100 to judge the full set¹ and GPT-5.5 67.3% at \$45, while the best sub-\$3 judge drops to 57.0% and the cheapest tier falls below 50%. On the full set the same trade-off looks benign, about 3 pp for a 42× price cut (§E.1). However, on OSReward-Hard, this trade-off disappears. At the millions of calls an RL run issues, even the frontier judges are out of reach, the gap the open reward model of §6 is built to close: the reliability of a strong judge at a cost that scales to training.

Takeaway: Reliability can be bought, but not afforded: the judges that hold ~70% on OSReward-Hard cost \$45–100 to judge the full set once, which does not scale to training-time reward.

6. OS-Shepherd: An Open Reward Model

The study so far leaves the field in a bind: reliable judging is unaffordable at the call volumes rejection sampling, RL, and trajectory mining need, and the affordable judges are weak or systematically lenient. Open data is scarce as well, especially the grounded FAIL trajectories a reward model must learn to reject. We therefore build OS-Shepherd: an open, self-hostable reward model, trained on a corpus we release, accurate enough to trust, aimed squarely at the false-success mode of §4.3, and cheap enough to run at training scale.

6.1. OS-Shepherd-100K

Training a reward model needs labels at a scale human annotation cannot reach on an academic budget: the benchmark’s 1019 gold trajectories already cost roughly 800 human hours (three independent annotators per trajectory plus meta-review on disputes, §3.3), and training needs two orders of magnitude more. VLM judges are the only affordable annotators, and §4 has just measured how far any single one can be trusted. The ensembling findings of the analysis show the way out. Voting cannot sharpen a verdict, because judges herd (§5.3); but corpus building, unlike evaluation, is free to discard trajectories. We therefore use agreement across diverse strong judges to select rather than to vote: a trajectory enters OS-Shepherd-100K only when the judges independently reach the same verdict, and the retained labels prove reliable in practice. Labeling then runs automatically at scale, on our own collection infrastructure and extended instruction pool (§§3.1 and 3.2). The corpus is also built for diversity, so the model learns task success rather than any one agent’s style: the rollouts span five model families (Claude, Gemini, GPT, Kimi, Qwen) under varying harnesses, action spaces, and step budgets, joined by independent open-source agent stacks. The supervision a sample carries is its verified task outcome, never the agent’s identity, and the corpus is checked to be free of contamination from the evaluation trajectories (§D.1).

From Instructions to Annotated Trajectories. The pipeline (Fig. 10) filters the self-collected instructions, rolls them out, and joins the surviving trajectories with open-source ones (Wang et al., 2025; Cheng et al., 2026; Liu et al., 2026) and a small OS-Genesis data split we collect ourselves (§A.1). Each

¹Costs use official API list prices where available; open-weight models without official pricing are priced at market rates (May 2026) for similar-size models.

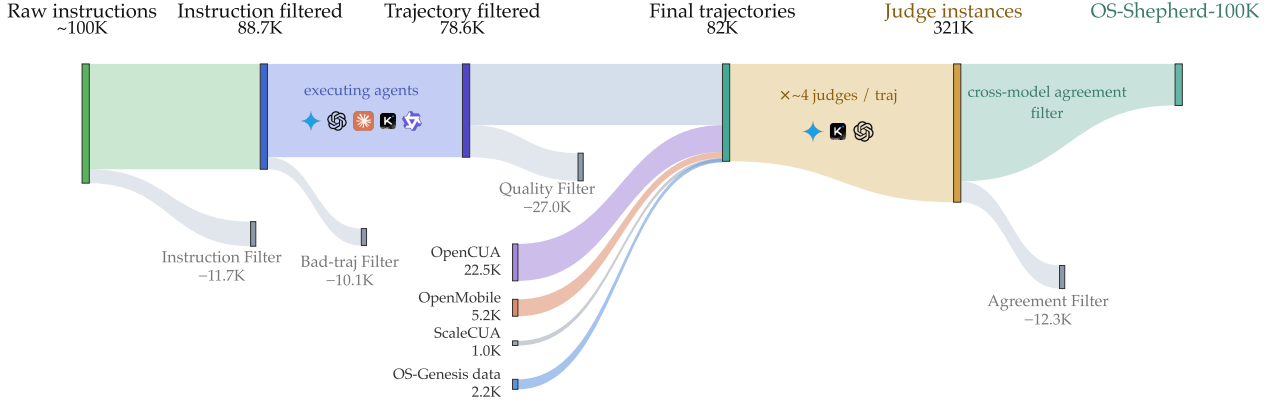


Figure 10: The OS-Shepherd-100K pipeline: self-collected and open-source trajectories are ensemble-judged and distilled into the training set. Band widths $\propto$ trajectory counts.

kept trajectory is scored by an ensemble of strong VLM-as-a-Judge runs under the protocol of §4.1, varying both the judge model (Gemini-3.1-Pro, Kimi-K2.5, Gemini-3-Pro, and others) and the screenshot setting across runs. Aggregating these votes and keeping only high-agreement trajectories yields the OS-Shepherd-100K training set (counts in Fig. 10; full per-source, per-setting, and length statistics in §D.1). In the figure’s units, a *judge instance* is one judge’s verdict on one trajectory, and a *sample* is one retained trajectory–response pair; the agreement filter keeps about 85% of the judged trajectories. Crucially, every sample carries the judge’s *reasoning*, not just a binary verdict (the natural target, given that §5.2 finds the verdict lives mainly in text), making OS-Shepherd-100K the first large-scale reasoning-annotated judge corpus for CUA trajectories.

Coverage. The corpus spans web, Windows, macOS, Ubuntu, and mobile (Table 3). Each collected trajectory is judged by as many ensemble runs as the budget allows, and this pool of judge instances is the raw material OS-Shepherd-100K is distilled from. About 46% of the Ubuntu data interleaves GUI actions with command-line steps, a mode we collect deliberately: as model capability grows, we expect combined GUI-and-CLI interaction to become the more mainstream way of using a computer. After filtering, the retained samples split mainly between desktop and web, with a small mobile share (§D.1).

Table 3: OS-Shepherd-100K judge-instance pool by platform.

Platform	Instances	Share
Web	119,469	37%
Windows	62,053	19%
macOS	45,028	14%
Ubuntu (GUI only)	34,355	11%
Ubuntu (GUI + CLI)	29,785	9%
Mobile	30,941	10%
Total	321,631	100%

Label Provenance. Because judges err on the same cases and share one lenient bias (§5.3), a single judge would propagate its own bias into the corpus. We therefore vary both the judge model and the screenshot setting and keep only strong judges, the accuracy-matched diversity that helps. We then retain a trajectory only when its votes are high-agreement, near-unanimous with no dissent from the strongest judges, and drop the ambiguous middle, so no forced majority enters the corpus. The *label* comes from this diverse ensemble. The retained *response*, the reasoning the model imitates, defaults to one strong judge’s, since mixing styles freely destabilized early training; when that judge did not score the trajectory or its verdict differs from the label, we retain another agreeing judge’s response instead, so the reasoning never contradicts the label, and the occasional

Table 4: OS-Shepherd against its untuned base, on the full set and OSReward-Hard.

Model	OSReward				OSReward-Hard			
	Acc	sRec	fRec	BalAcc	Acc	sRec	fRec	BalAcc
Qwen3.5-9B (base)	76.7	98.9	59.9	79.4	39.4	97.7	14.1	55.9
OS-Shepherd-9B	86.1	86.6	86.0	86.3	60.2	66.3	57.6	61.9
Qwen3.5-35B-A3B (base)	82.2	92.4	74.5	83.5	51.1	83.7	36.9	60.3
OS-Shepherd-35B-A3B	85.6	85.0	86.2	85.6	62.7	68.6	60.1	64.3

exposure to other judges’ reasoning aids generalization. Each trajectory contributes at most two samples, one per output format (§D.2), and we balance the final SUCCESS/FAIL mix.

Negatives and Synthesized Instructions. We deliberately over-collect grounded FAIL trajectories: because the agreement filter and the multi-judge ensemble screen out ungrounded or trivially blocked runs, the resulting negatives reflect real task failures rather than environment artifacts. Writing and vetting every instruction by hand is affordable for a benchmark of a thousand trajectories, but not for a corpus two orders of magnitude larger. Collection at this scale therefore leans further on synthesized instructions, a portion written by reverse task synthesis (Sun et al., 2025a). The synthesis concentrates on desktop, where grounded instructions are hardest to author by hand: it accounts for about 25% of the Ubuntu and Windows training instructions combined and about 10% of the web instructions.

6.2. Training

OS-Shepherd is trained in two stages on OS-Shepherd-100K and judges under the same protocol. We start from Qwen3.5 (Qwen Team, 2026).

SFT-Stage. We fine-tune the base VLM on the OS-Shepherd-100K corpus (the 96.6K agreement-filtered samples) to judge under the main setting. SFT alone lifts OS-Shepherd-9B far above its Qwen3.5-9B base, primarily by correcting the base model’s near-total leniency (Fig. 14).

RL-Stage. What SFT leaves behind is the study’s dominant failure, the *false success* (§4.3), the most harmful error for a reward signal since it directly reinforces incorrect behavior. The ability to catch it is latent rather than missing (the SFT model resolves many under repeated sampling), so we mine those cases and leverage GRPO (Shao et al., 2024); training details are in §D.

6.3. Evaluating OS-Shepherd

Table 4 reports both trained models against the untuned checkpoints they start from, on the full set and OSReward-Hard.

Where OS-Shepherd-9B Lands. Trained end-to-end, the 9B moves from the bottom third of the field into the mid-tier commercial band on the full set. What makes it usable as a reward signal shows on OSReward-Hard: it catches 57.6% of hard false successes, where the cheap and lenient field misses nearly all of them, and it stays on the balanced diagonal while doing so (Fig. 1b). That combination is rare; the only judges that match it are frontier models at many times its cost. Training also buys robustness, the 9B’s hard-set drop being a third smaller than its base’s, and on the cost–accuracy view no other sub-\$2 judge comes near it: at \$1.36 to judge the full set, the 9B is about a third cheaper than its nearest commercial tier-mate while leading it on OSReward-Hard (Fig. 2; details in §E.1).

Scale Adds Little. OS-Shepherd-35B-A3B follows the identical recipe and lands beside the 9B: a small further gain on OSReward-Hard, level full-set accuracy, the same balanced diagonal, and the same behavior on the three held-out CUA benchmarks (Table 4; §7). Four times the parameters buy 2.4 pp of hard-set balanced accuracy and nothing on the full set; what transfers is the recipe. On the quality axes both models track the field’s weakness (§4.5), the 35B grading noticeably better than the 9B.

Cost at Training Scale. A reward model is queried throughout a training run, so even a small per-trajectory cost accumulates. For example, a single modest bout of online RL (200 updates, batch 16, 16 rollouts) already issues $200 \times 16 \times 16 = 51,200$ judge calls: about \$4,000 with Claude-Opus-4-8 or \$2,300 with GPT-5.5 at OSReward’s average trajectory cost, against about \$68 with OS-Shepherd-9B, a 30–60 $\times$ reduction that compounds over the many runs a project needs. And because OS-Shepherd-9B is open and self-hostable, even that \$68 is only an API-equivalent figure: in practice the marginal cost is the lab’s own GPU time.

We frame OS-Shepherd as evidence that reliable CUA reward no longer has to be bought at frontier prices: a small, open, self-hostable judge can stay level with mid-tier commercial judges on the full set, lead every similarly priced judge on OSReward-Hard, and close most of the frontier gap at 30–60 $\times$ lower cost. The trained models keep this stricter, failure-catching operating point on the held-out benchmarks (§7), and §4.5 carries the complementary quality-grading probe.

Takeaway: With our data and recipe, small open judges match commercial ones, producing reliable reward at a cost that scales.

7. Generalization to Existing Benchmarks

Every number so far comes from OSReward’s own gold. To test whether the picture and OS-Shepherd’s de-biasing extend beyond our data, we run the same judges on three benchmarks covering prior work’s ground (AndroidWorld, WebArena, OSWorld) against each one’s own human-written verifier (Fig. 11).

7.1. Judges vs. Human-Written Verifiers

The introduction’s pilot study already flagged that even the best judges disagree with existing benchmarks’ own verifiers on many desktop verdicts; we now measure that at scale. We take $\sim 90\%$ agreement with a benchmark’s human-written verifier as the bar for a judge that could replace it. Agreement is *platform-driven*, not judge-driven (Fig. 11, left): the best judges approach that bar on mobile, come within about 6 pp on web, and fall well short on desktop, and the platform ordering is the same for every judge. These verifiers are themselves imperfect (Xie et al., 2025), with false positives and false negatives of their own that we do not correct here, so the figure measures agreement with each benchmark’s verifier rather than ground truth; a judge’s true accuracy is likely somewhat higher than the agreement shown. The desktop shortfall reflects the leniency bias of §4.3, now outside OSReward: fail recall collapses, judges accept failed runs as successes, and false positives concentrate in unverifiable domains and on long trajectories.

7.2. OS-Shepherd Transfers Across Benchmarks

OS-Shepherd is trained on OS-Shepherd-100K and, like every judge here, evaluated held-out on OS-Reward; the question now is whether its de-biasing also transfers to benchmarks built independently, each with its own human-written verifier. It does: the OS-Shepherd models are the best open judges on OSWorld and AndroidWorld and sit in the frontier cluster on WebArena, beating every general open model

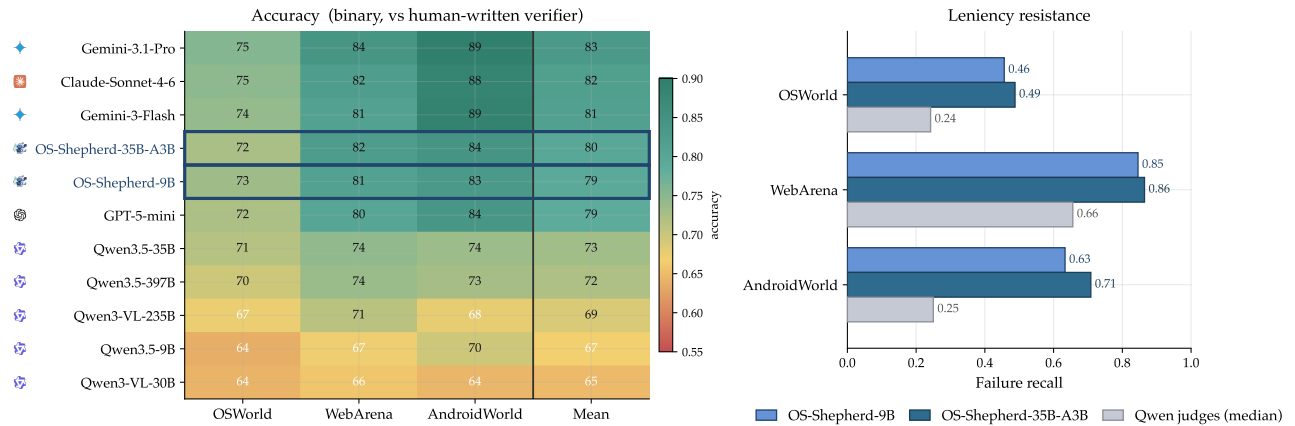


Figure 11: Judges on three existing CUA benchmarks against each benchmark’s human-written verifier (matched subsets): accuracy and failure recall; means are computed before rounding.

up to Qwen3.5-397B-A17B ($\sim 44\times$ the 9B’s size) on all three (Fig. 11, left). None of these benchmarks’ tasks or trajectories enter training or appear in OSReward.

What Transfers Is Catching Failures. What transfers is the de-biasing itself (Fig. 11, right). Where successes dominate (OSWorld, AndroidWorld), OS-Shepherd-9B catches far more true failures than any general Qwen judge, whose median sits near “always SUCCESS.” WebArena inverts the base rate, so an over-lenient judge scores *below* the trivial all-FAIL baseline: OS-Shepherd-9B’s own untuned Qwen3.5-9B base lands there, while OS-Shepherd-9B is the tightest-calibrated judge in the field, frontier included. The resistance is therefore learned, and from training rather than size: general Qwen judges stay lenient at every size from 30B to 397B, while OS-Shepherd-35B-A3B, four times the parameters on the identical recipe, lands essentially beside the 9B. What training buys is a relocated operating point: more failure recall where false successes dominate, a little less success recognition elsewhere, at roughly flat balanced accuracy. Because Fig. 11 reports the final models, these results establish transfer of the full recipe rather than isolating the RL stage.

Toward a Scalable Verifier. Each benchmark’s human-written verifier is costly to author and does not generalize past its own tasks, exactly the non-scalability a learned judge is meant to fix. A single 9B, trained on none of them, carries its de-biasing across all three at a cost that scales to training-time reward (§6.3). It does not yet match frontier accuracy, but closes most of the gap at a fraction of the scale, the direction a scalable CUA reward signal must move.

Takeaway: Across three independently built CUA benchmarks, OS-Shepherd’s learned leniency resistance transfers out of distribution: one small open reward model narrows the gap to per-task human-written verifiers, a step toward a single scalable CUA reward signal.

8. Conclusion

The reward signal for computer-using agents rests on VLM judges whose reliability had gone unexamined. We build OSReward to measure it, with human-gold trajectories collected across four platforms and released in full. Benchmarking a wide range of judges reveals one dominant failure, accepting incomplete tasks as successes, driven by verdicts that follow the agent’s text history more than the screen. On

OSReward-Hard the field drops to near chance, and the judges that hold up cost far too much for training. We then turn these findings into OS-Shepherd: an open corpus and open reward models whose RL stage targets the false-success mode directly. The trained models match commercial judges at a small fraction of frontier cost and stay de-biased on benchmarks they never trained on, a reward signal can run itself at training scale under an academic budget.

References

- [1] Anthropic. Introducing Claude Haiku 4.5. <https://www.anthropic.com/news/claude-haiku-4-5>, October 2025.
- [2] Anthropic. Introducing Claude Opus 4.6. <https://www.anthropic.com/news/claude-opus-4-6>, February 2026a.
- [3] Anthropic. Introducing Claude Opus 4.8. <https://www.anthropic.com/news/claude-opus-4-8>, May 2026b.
- [4] Anthropic. Introducing Claude Sonnet 4.6. <https://www.anthropic.com/news/claude-sonnet-4-6>, February 2026c.
- [5] Hao Bai, Yifei Zhou, Mert Cemri, Jiayi Pan, Alane Suhr, Sergey Levine, and Aviral Kumar. Digirl: Training in-the-wild device-control agents with autonomous reinforcement learning, 2024. URL <https://arxiv.org/abs/2406.11896>.
- [6] Shuai Bai, Yuxuan Cai, Ruizhe Chen, Keqin Chen, Xionghui Chen, Zesen Cheng, Lianghao Deng, Wei Ding, Chang Gao, Chunjiang Ge, et al. Qwen3-vl technical report. *arXiv preprint arXiv:2511.21631*, 2025.
- [7] Rogerio Bonatti, Dan Zhao, Francesco Bonacci, Dillon Dupont, Sara Abdali, Yinheng Li, Yadong Lu, Justin Wagle, Kazuhito Koishida, Arthur Buckner, Lawrence Jang, and Zack Hui. Windows agent arena: Evaluating multi-modal os agents at scale, 2024. URL <https://arxiv.org/abs/2409.08264>.
- [8] ByteDance Seed Team. Seed2.0 model card: Towards intelligence frontier for real-world complexity. Technical report, ByteDance, February 2026.
- [9] Hyunjoo Chae, Sunghwan Kim, Junhee Cho, Seungone Kim, Seungjun Moon, Gyeom Hwangbo, Dongha Lim, Minjin Kim, Yeonjun Hwang, Minju Gwak, Dongwook Choi, Minseok Kang, Gwanhoon Im, ByeongUng Cho, Hyojun Kim, Jun Hee Han, Taeyoon Kwon, Minju Kim, Beong woo Kwak, Dongjin Kang, and Jinyoung Yeo. Web-shepherd: Advancing PRMs for reinforcing web agents. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*, 2025. URL <https://openreview.net/forum?id=G2kMro09UV>.
- [10] Cong Chen, Kaixiang Ji, Hao Zhong, Muzhi Zhu, Anzhou Li, Guo Gan, Ziyuan Huang, Cheng Zou, Jiajia Liu, Jingdong Chen, Hao Chen, and Chunhua Shen. Gui-shepherd: Reliable process reward and verification for long-sequence gui tasks, 2025a. URL <https://arxiv.org/abs/2509.23738>.
- [11] Dongping Chen, Ruoxi Chen, Shilin Zhang, Yinyu Liu, Yaochen Wang, Huichi Zhou, Qihui Zhang, Yao Wan, Pan Zhou, and Lichao Sun. Mllm-as-a-judge: Assessing multimodal llm-as-a-judge with vision-language benchmark, 2024. URL <https://arxiv.org/abs/2402.04788>.

- [12] Xuetian Chen, Yinghao Chen, Xinfeng Yuan, Zhuo Peng, Lu Chen, Yuekeng Li, Zhoujia Zhang, Yingqian Huang, Leyan Huang, Jiaqing Liang, et al. Os-map: How far can computer-using agents go in breadth and depth? *arXiv preprint arXiv:2507.19132*, 2025b.
- [13] Kanzhi Cheng, Qiushi Sun, Yougang Chu, Fangzhi Xu, Li YanTao, Jianbing Zhang, and Zhiyong Wu. SeeClick: Harnessing GUI grounding for advanced visual GUI agents. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 9313–9332, Bangkok, Thailand, August 2024. Association for Computational Linguistics. URL <https://aclanthology.org/2024.acl-long.505>.
- [14] Kanzhi Cheng, Zehao Li, Zheng Ma, Nuo Chen, Jialin Cao, Qiushi Sun, Zichen Ding, Fangzhi Xu, Hang Yan, Jiajun Chen, et al. Openmobile: Building open mobile agents with task and trajectory synthesis. *arXiv preprint arXiv:2604.15093*, 2026.
- [15] Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. Training verifiers to solve math word problems, 2021. URL <https://arxiv.org/abs/2110.14168>.
- [16] Jacob Cohen. A coefficient of agreement for nominal scales. *Educational and Psychological Measurement*, 20(1):37–46, 1960.
- [17] Gheorghe Comanici, Eric Bieber, Mike Schaekermann, Ice Pasupat, Noveen Sachdeva, Inderjit Dhillon, Marcel Blistein, Ori Ram, Dan Zhang, Evan Rosen, et al. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *arXiv preprint arXiv:2507.06261*, 2025.
- [18] Gemini Team. Gemini 3 pro model card, November 2025. URL <https://deepmind.google/models/gemini/>.
- [19] Boyu Gou, Ruohan Wang, Boyuan Zheng, Yanan Xie, Cheng Chang, Yiheng Shu, Huan Sun, and Yu Su. Navigating the digital world as humans do: Universal visual grounding for GUI agents. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=kxnoqaisCT>.
- [20] Aaron Hurst, Adam Lerer, Adam P Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, et al. GPT-4o system card. *arXiv preprint arXiv:2410.21276*, 2024.
- [21] Chengyou Jia, Minnan Luo, Zhuohang Dang, Qiushi Sun, Fangzhi Xu, Junlin Hu, Tianbao Xie, and Zhiyong Wu. AgentStore: Scalable integration of heterogeneous agents as specialized generalist computer assistant. In *Findings of the Association for Computational Linguistics: ACL 2025*, pages 8908–8934, Vienna, Austria, July 2025. Association for Computational Linguistics. ISBN 979-8-89176-256-5. doi: 10.18653/v1/2025.findings-acl.466. URL <https://aclanthology.org/2025.findings-acl.466/>.
- [22] Deyang Jiang, Jing Huang, Xuanle Zhao, Lei Chen, Liming Zheng, Fanfan Liu, Haibo Qiu, Peng Shi, and Zhixiong Zeng. TreeCUA: Efficiently scaling GUI automation with tree-structured verifiable evolution. In *Forty-third International Conference on Machine Learning*, 2026. URL <https://openreview.net/forum?id=KBCWS60BnD>.

- [23] Kimi Team, Tongtong Bai, Yifan Bai, Yiping Bao, SH Cai, Yuan Cao, Y Charles, HS Che, Cheng Chen, Guanduo Chen, et al. Kimi k2. 5: Visual agentic intelligence. *arXiv preprint arXiv:2602.02276*, 2026.
- [24] Quyu Kong, Xu Zhang, Zhenyu Yang, Nolan Gao, Chen Liu, Panrong Tong, Chenglin Cai, Hanzhang Zhou, Jianan Zhang, Liangyu Chen, et al. Mobileworld: Benchmarking autonomous mobile agents in agent-user interactive and mcp-augmented environments. In *Proceedings of the 64th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 6142–6167, 2026.
- [25] Klaus Krippendorff. Computing Krippendorff’s alpha-reliability. Technical report, University of Pennsylvania, Annenberg School for Communication, 2011.
- [26] Lei Li, Yuancheng Wei, Zhihui Xie, Xuqing Yang, Yifan Song, Peiyi Wang, Chenxin An, Tianyu Liu, Sujian Li, Bill Yuchen Lin, et al. Vl-rewardbench: a challenging benchmark for vision-language generative reward models. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 24657–24668, 2025.
- [27] Zehao Li, Zhenyu Wu, Yibo Zhao, Bowen Yang, Jingjing Xie, Zhaoyang Liu, Zhoumianze Liu, Kaiming Jin, Jianze Liang, Zonglin Li, Feng Wu, Bowen Zhou, Zun Wang, and Zichen Ding. Os-themis: A scalable critic framework for generalist gui rewards, 2026. URL <https://arxiv.org/abs/2603.19191>.
- [28] Haojia Lin, Xiaoyu Tan, Yulei Qin, Zihan Xu, Yuchen Shi, Zongyi Li, Gang Li, Shaofei Cai, Siqi Cai, Chaoyou Fu, Ke Li, and Xing Sun. Cuarewardbench: A benchmark for evaluating reward models on computer-using agent, 2025. URL <https://arxiv.org/abs/2510.18596>.
- [29] Zhaoyang Liu, JingJing Xie, Zichen Ding, Zehao Li, Bowen Yang, Zhenyu Wu, Xuehui Wang, Qiushi Sun, Shi Liu, Weiyun Wang, Shenglong Ye, Qingyun Li, Zeyue Tian, Gen Luo, Xiangyu Yue, Biqing Qi, Kai Chen, Bowen Zhou, Yu Qiao, Qifeng Chen, and Wenhai Wang. ScaleCUA: Scaling open-source computer use agents with cross-platform data. In *The Fourteenth International Conference on Learning Representations*, 2026. URL <https://openreview.net/forum?id=yBFUqdJFZn>.
- [30] Xing Han Lù, Amirhossein Kazemnejad, Nicholas Meade, Arkil Patel, Dongchan Shin, Alejandra Zambrano, Karolina Stańczak, Peter Shaw, Christopher J. Pal, and Siva Reddy. Agentrewardbench: Evaluating automatic evaluations of web agent trajectories, 2025. URL <https://arxiv.org/abs/2504.08942>.
- [31] OpenAI. Computer-using agent: Introducing a universal interface for ai to interact with the digital world, 2025a. URL <https://openai.com/index/computer-using-agent>.
- [32] OpenAI. GPT-5 System Card. Technical report, OpenAI, August 2025b. URL <https://openai.com/index/gpt-5-system-card/>.
- [33] OpenAI. GPT-5.4 Thinking System Card. Technical report, OpenAI, March 2026a. URL <https://openai.com/index/gpt-5-4-thinking-system-card/>.
- [34] OpenAI. GPT-5.5 System Card. Technical report, OpenAI, April 2026b. URL <https://openai.com/index/gpt-5-5-system-card/>.
- [35] Jiayi Pan, Yichi Zhang, Nicholas Tomlin, Yifei Zhou, Sergey Levine, and Alane Suhr. Autonomous evaluation and refinement of digital agents. In *First Conference on Language Modeling*, 2024. URL <https://openreview.net/forum?id=NPAQ6FKSmK>.

- [36] Zehan Qi, Xiao Liu, Iat Long Iong, Hanyu Lai, Xueqiao Sun, Wenyi Zhao, Yu Yang, Xinyue Yang, Jiadai Sun, Shuntian Yao, Tianjie Zhang, Wei Xu, Jie Tang, and Yuxiao Dong. Webrl: Training llm web agents via self-evolving online curriculum reinforcement learning, 2024. URL <https://arxiv.org/abs/2411.02337>.
- [37] Yujia Qin, Yining Ye, Junjie Fang, Haoming Wang, Shihao Liang, Shizuo Tian, Junda Zhang, Jiahao Li, Yunxin Li, Shijue Huang, et al. Ui-tars: Pioneering automated gui interaction with native agents. *arXiv preprint arXiv:2501.12326*, 2025.
- [38] Qwen Team. Qwen3.5: Towards native multimodal agents, February 2026. URL <https://qwen.ai/blog?id=qwen3.5>.
- [39] Christopher Rawles, Sarah Clinckemaillie, Yifan Chang, Jonathan Waltz, Gabrielle Lau, Marybeth Fair, Alice Li, William E Bishop, Wei Li, Folawiyo Campbell-Ajala, Daniel Kenji Toyama, Robert James Berry, Divya Tyamagundlu, Timothy P Lillicrap, and Oriana Riva. Androidworld: A dynamic benchmarking environment for autonomous agents. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=il5yUQsrjC>.
- [40] Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, YK Li, Yang Wu, et al. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*, 2024.
- [41] Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. Hybridflow: A flexible and efficient RLHF framework. *arXiv preprint arXiv:2409.19256*, 2024.
- [42] Qiushi Sun, Zhirui Chen, Fangzhi Xu, Kanzhi Cheng, Chang Ma, Zhangyue Yin, Jianing Wang, Chengcheng Han, Renyu Zhu, Shuai Yuan, et al. A survey of neural code intelligence: Paradigms, advances and beyond. *arXiv preprint arXiv:2403.14734*, 2024.
- [43] Qiushi Sun, Kanzhi Cheng, Zichen Ding, Chuanyang Jin, Yian Wang, Fangzhi Xu, Zhenyu Wu, Chengyou Jia, Liheng Chen, Zhoumianze Liu, et al. Os-genesis: Automating gui agent trajectory construction via reverse task synthesis. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5555–5579, 2025a.
- [44] Qiushi Sun, Mukai Li, Zhoumianze Liu, Zhihui Xie, Fangzhi Xu, Zhangyue Yin, Kanzhi Cheng, Zehao Li, Zichen Ding, Qi Liu, Zhiyong Wu, Zhuosheng Zhang, Ben Kao, and Lingpeng Kong. OS-sentinel: Towards safety-enhanced mobile GUI agents via hybrid validation in realistic workflows. In *Proceedings of the 64th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 9529–9553, San Diego, California, United States, July 2026a. Association for Computational Linguistics. ISBN 979-8-89176-390-6. URL <https://aclanthology.org/2026.acl-long.431/>.
- [45] Qiushi Sun, Zhoumianze Liu, Chang Ma, Zichen Ding, Fangzhi Xu, Zhangyue Yin, Haiteng Zhao, Zhenyu Wu, Kanzhi Cheng, Zhaoyang Liu, Jianing Wang, Qintong Li, Xiangru Tang, Tianbao Xie, Xiachong Feng, Xiang Li, Ben Kao, Wenhai Wang, Biqing Qi, Lingpeng Kong, and Zhiyong Wu. Scienceboard: Evaluating multimodal autonomous agents in realistic scientific workflows. In *The Fourteenth International Conference on Learning Representations*, 2026b. URL <https://openreview.net/forum?id=bJvwJahJeF>.

- [46] Zeyi Sun, Yuhang Cao, Jianze Liang, Qiushi Sun, Ziyu Liu, Zhixiong Zhang, Yuhang Zang, Xiaoyi Dong, Kai Chen, Dahua Lin, et al. Coda: Coordinating the cerebrum and cerebellum for a dual-brain computer use agent with decoupled reinforcement learning. *arXiv preprint arXiv:2508.20096*, 2025b.
- [47] Brandon Trabucco, Gunnar Sigurdsson, Robinson Piramuthu, and Ruslan Salakhutdinov. InSTA: Towards internet-scale training for agents. *arXiv preprint arXiv:2502.06776*, 2025.
- [48] Peiyi Wang, Lei Li, Zhihong Shao, Runxin Xu, Damai Dai, Yifei Li, Deli Chen, Yu Wu, and Zhifang Sui. Math-shepherd: Verify and reinforce LLMs step-by-step without human annotations. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 9426–9439, Bangkok, Thailand, August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.acl-long.510. URL <https://aclanthology.org/2024.acl-long.510/>.
- [49] Xinyuan Wang, Bowen Wang, Dunjie Lu, Junlin Yang, Tianbao Xie, Junli Wang, Jiaqi Deng, Xiaole Guo, Yiheng Xu, Chen Henry Wu, Zhennan Shen, Zhuokai Li, Ryan Li, Xiaochuan Li, Junda Chen, Zheng Boyuan, LI PEIHANG, Fangyu Lei, Ruisheng Cao, Yeqiao Fu, Dongchan Shin, Martin Shin, Hu Jiarui, Yuyan Wang, Jixuan Chen, Yuxiao Ye, Danyang Zhang, Yipu Wang, Heng Wang, Diyi Yang, Victor Zhong, Y. Charles, Zhilin Yang, and Tao Yu. OpenCUA: Open foundations for computer-use agents. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*, 2025. URL <https://openreview.net/forum?id=6iRZvJiC9Q>.
- [50] Zhaoyang Wang, Yiming Liang, Xuchao Zhang, Qianhui Wu, Siwei Han, Anson Bastos, Rujia Wang, Chetan Bansal, Baolin Peng, Jianfeng Gao, Saravan Rajmohan, and Huaxiu Yao. SynthAgent: Adapting web agents with synthetic supervision. In *Proceedings of the 64th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 15730–15752, San Diego, California, United States, July 2026. Association for Computational Linguistics. ISBN 979-8-89176-390-6. URL <https://aclanthology.org/2026.acl-long.716/>.
- [51] Qianhui Wu, Kanzhi Cheng, Rui Yang, Chaoyun Zhang, Jianwei Yang, Huiqiang Jiang, Jian Mu, Baolin Peng, Bo Qiao, Reuben Tan, et al. Gui-actor: Coordinate-free visual grounding for gui agents. *arXiv preprint arXiv:2506.03143*, 2025a.
- [52] Zhenyu Wu, Jingjing Xie, Zehao Li, Bowen Yang, Qiushi Sun, Zhaoyang Liu, Zhoumianze Liu, Yu Qiao, Xiangyu Yue, Zun Wang, et al. Os-oracle: A comprehensive framework for cross-platform gui critic models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 27514–27524, 2026.
- [53] Zhiyong Wu, Chengcheng Han, Zichen Ding, Zhenmin Weng, Zhoumianze Liu, Shunyu Yao, Tao Yu, and Lingpeng Kong. Os-copilot: Towards generalist computer agents with self-improvement, 2024. URL <https://arxiv.org/abs/2402.07456>.
- [54] Zhiyong Wu, Zhenyu Wu, Fangzhi Xu, Yian Wang, Qiushi Sun, Chengyou Jia, Kanzhi Cheng, Zichen Ding, Liheng Chen, Paul Pu Liang, and Yu Qiao. OS-ATLAS: Foundation action model for generalist GUI agents. In *The Thirteenth International Conference on Learning Representations*, 2025b. URL <https://openreview.net/forum?id=n9PDaFNi8t>.
- [55] Han Xiao, Guozhi Wang, Yuxiang Chai, Zimu Lu, Weifeng Lin, Hao He, Lue Fan, Liuyang Bian, Rui Hu, Liang Liu, Shuai Ren, Yafei Wen, Xiaoxin Chen, Aojun Zhou, and Hongsheng Li. UI-genie: A self-improving approach for iteratively boosting MLLM-based mobile GUI agents. In

The Thirty-ninth Annual Conference on Neural Information Processing Systems, 2025. URL <https://openreview.net/forum?id=3uUmJzSSOW>.

- [56] Tianbao Xie, Danyang Zhang, Jixuan Chen, Xiaochuan Li, Siheng Zhao, Ruisheng Cao, Toh Jing Hua, Zhoujun Cheng, Dongchan Shin, Fangyu Lei, Yitao Liu, Yiheng Xu, Shuyan Zhou, Silvio Savarese, Caiming Xiong, Victor Zhong, and Tao Yu. OSWorld: Benchmarking multimodal agents for open-ended tasks in real computer environments. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2024. URL <https://openreview.net/forum?id=tN61DTr4Ed>.
- [57] Tianbao Xie, Mengqi Yuan, Danyang Zhang, Xinzhuang Xiong, Zhennan Shen, Zilong Zhou, Xinyuan Wang, Yanxu Chen, Jiaqi Deng, Junda Chen, Bowen Wang, Haoyuan Wu, Jixuan Chen, Junli Wang, Dunjie Lu, Hao Hu, and Tao Yu. Introducing osworld-verified. *xlang.ai*, July 2025. URL <https://xlang.ai/blog/osworld-verified>.
- [58] Fangzhi Xu, Hang Yan, Qiushi Sun, Jinyang Wu, Zixian Huang, Muye Huang, Jingyang Gong, Zichen Ding, Kanzhi Cheng, Yian Wang, et al. Odysseyarena: Benchmarking large language models for long-horizon, active and inductive interactions. *arXiv preprint arXiv:2602.05843*, 2026.
- [59] Yiheng Xu, Dunjie Lu, Zhennan Shen, Junli Wang, Zekun Wang, Yuchen Mao, Caiming Xiong, and Tao Yu. Agenttrek: Agent trajectory synthesis via guiding replay with web tutorials. In *The Thirteenth International Conference on Learning Representations*, 2025a. URL <https://openreview.net/forum?id=EEgYUccwsV>.
- [60] Yiheng Xu, Zekun Wang, Junli Wang, Dunjie Lu, Tianbao Xie, Amrita Saha, Doyen Sahoo, Tao Yu, and Caiming Xiong. Aguviz: Unified pure vision agents for autonomous GUI interaction. In *Forty-second International Conference on Machine Learning*, 2025b. URL <https://openreview.net/forum?id=PlihOwfx4r>.
- [61] Taofeng Xue, Chong Peng, Mianqiu Huang, Linsen Guo, Tiancheng Han, Haozhe Wang, Xiaocheng Zhang, Xin Yang, Dengchang Zhao, Jinrui Ding, Xiandi Ma, Yuchen Xie, Peng Pei, Xunliang Cai, and Xipeng Qiu. EvoCUA: Evolving computer use agents via learning from scalable synthetic experience. In *Second Workshop on Agents in the Wild: Safety, Security, and Beyond*, 2026a. URL <https://openreview.net/forum?id=pzbtm6fRIo>.
- [62] Tianci Xue, Weijian Qi, Tianneng Shi, Chan Hee Song, Boyu Gou, Dawn Song, Huan Sun, and Yu Su. An illusion of progress? assessing the current state of web agents. In *Second Conference on Language Modeling*, 2025. URL <https://openreview.net/forum?id=6jZi4HSs6o>.
- [63] Tianci Xue, Zeyi Liao, Tianneng Shi, Zilu Wang, Kai Zhang, Dawn Song, Yu Su, and Huan Sun. Autonomous continual learning for environment adaptation of computer-use agents, 2026b. URL <https://arxiv.org/abs/2602.10356>.
- [64] Bowen Yang, Kaiming Jin, Zhenyu Wu, Zhaoyang Liu, Qiushi Sun, Zehao Li, JingJing Xie, Zhoumianze Liu, Fangzhi Xu, Kanzhi Cheng, Yian Wang, Qingyun Li, Yu Qiao, Zun Wang, and Zichen Ding. OS-symphony: A holistic framework for robust and generalist computer-using agents. In *Proceedings of the 64th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 22300–22330, San Diego, California, United States, July 2026. Association for Computational Linguistics. ISBN 979-8-89176-390-6. URL <https://aclanthology.org/2026.acl-long.1021/>.

- [65] Jianwei Yang, Hao Zhang, Feng Li, Xueyan Zou, Chunyuan Li, and Jianfeng Gao. Set-of-mark prompting unleashes extraordinary visual grounding in gpt-4v. *arXiv preprint arXiv:2310.11441*, 2023.
- [66] Junlei Zhang, Zichen Ding, Chang Ma, Zijie Chen, Qiushi Sun, Zhenzhong Lan, and Junxian He. Breaking the data barrier – building GUI agents through task generalization. In *Second Conference on Language Modeling*, 2025. URL <https://openreview.net/forum?id=QDt0RaZt8K>.
- [67] Kai Zhang, Xiangchao Chen, Bo Liu, Tianci Xue, Zeyi Liao, Zhihan Liu, Xiyao Wang, Yuting Ning, Zhaorun Chen, Xiaohan Fu, Jian Xie, Yuxuan Sun, Boyu Gou, Qi Qi, Zihang Meng, Jianwei Yang, Ning Zhang, Xian Li, Ashish Shah, Dat Huynh, Hengduo Li, Zi Yang, Xuefei Cao, Lawrence Keunho Jang, Shuyan Zhou, Jiacheng Zhu, Huan Sun, Jason E Weston, Yu Su, and Yifan Wu. Agent learning via early experience. In *Forty-third International Conference on Machine Learning*, 2026. URL <https://openreview.net/forum?id=N3dXUHY5dD>.
- [68] Boyuan Zheng, Boyu Gou, Jihyung Kil, Huan Sun, and Yu Su. Gpt-4v(ision) is a generalist web agent, if grounded. In *Forty-first International Conference on Machine Learning*, 2024. URL <https://openreview.net/forum?id=piecKJ2DlB>.
- [69] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. Judging llm-as-a-judge with mt-bench and chatbot arena. In *Proceedings of the 37th International Conference on Neural Information Processing Systems, NIPS ’23*, Red Hook, NY, USA, 2023a. Curran Associates Inc.
- [70] Lianmin Zheng, Liangsheng Yin, Zhiqiang Xie, Chuyue Sun, Jeff Huang, Cody Hao Yu, Shiyi Cao, Christos Kozyrakis, Ion Stoica, Joseph E. Gonzalez, Clark Barrett, and Ying Sheng. SGLang: Efficient execution of structured language model programs. *arXiv preprint arXiv:2312.07104*, 2023b.
- [71] Enyu Zhou, Guodong Zheng, Binghai Wang, Zhiheng Xi, Shihan Dou, Rong Bao, Wei Shen, Limao Xiong, Jessica Fan, Yurong Mou, Rui Zheng, Tao Gui, Qi Zhang, and Xuanjing Huang. RMB: Comprehensively benchmarking reward models in LLM alignment. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=kmgrlG9TR0>.
- [72] Shuyan Zhou, Frank F. Xu, Hao Zhu, Xuhui Zhou, Robert Lo, Abishek Sridhar, Xianyi Cheng, Tianyue Ou, Yonatan Bisk, Daniel Fried, Uri Alon, and Graham Neubig. Webarena: A realistic web environment for building autonomous agents. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=oKn9c6ytLx>.
- [73] Mingchen Zhuge, Changsheng Zhao, Dylan R. Ashley, Wenyi Wang, Dmitrii Khizbullin, Yunyang Xiong, Zechun Liu, Ernie Chang, Raghuraman Krishnamoorthi, Yuandong Tian, Yangyang Shi, Vikas Chandra, and Jürgen Schmidhuber. Agent-as-a-judge: Evaluate agents with agents. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=Nn9P0I9Ekt>.
- [74] Yicheng Zou, Dongsheng Zhu, Lin Zhu, Tong Zhu, Yunhua Zhou, Peiheng Zhou, Xinyu Zhou, Dongzhan Zhou, Zhiwang Zhou, Yuhao Zhou, et al. Intern-s1-pro: Scientific multimodal foundation model at trillion scale. *arXiv preprint arXiv:2603.25040*, 2026.

Appendix Contents

A	Data-Collection Infrastructure	27
A.1	Web	28
A.2	Windows	29
A.3	Ubuntu	30
A.4	Android	33
B	OSReward Details	34
B.1	Annotators	34
B.2	Annotation Pipeline and Statistics	34
B.3	Failure-Type Taxonomy	35
B.4	OSReward-Multi Scoring Guideline	35
C	Experimental Details	36
C.1	Evaluated Models	36
C.2	Judging Prompt	36
C.3	Judge Error Categories	37
D	OS-Shepherd	39
D.1	OS-Shepherd-100K	39
D.2	Training Details	40
E	Additional Results and Analysis	42
E.1	Per-Task Breakdown	42
E.2	Multi-Axis Evaluation (Auxiliary)	42
E.3	Open-Weight versus Closed-Source Judges	43
E.4	Inter-Judge Agreement	44
E.5	Verifiability and False Positives	44
E.6	Thinking, Screenshot Count, and Self-Consistency	44
F	Case Studies	47

A. Data-Collection Infrastructure

This appendix details the cross-platform collection infrastructure summarized in §3.1. Collection on every platform runs through the same course (realistic environment, grounded instructions, rollout by

executing agents, and automatic verification); here we give each platform’s environment back-end and its concrete instruction and verification implementation. An *executing agent* is the agent that operates the environment and produces a trajectory: a model backbone driving that platform’s action space (Tables 5, 6 and 8) through the platform’s own harness; the backbones used are listed per platform below. The executing agents follow the ReAct-style loop and its variants (Yang et al., 2026; Cheng et al., 2026), which we adopt to collect the longer-horizon and more complex trajectories the judge study needs. The same infrastructure produces two disjoint datasets: the human-verified OSReward evaluation trajectories (§3) and the larger, ensemble-labeled OS-Shepherd training corpus (§6). For each platform we give the environment and software coverage, the instruction method, the trajectory collection, and how the two datasets are drawn from it. We are also training further model sizes with newer training techniques; under our compute limits these runs need more time and will be added in a future iteration of this work.

A.1. Web

Environment. The web environment is a Playwright browser service with one isolated, headless-Chromium session per rollout. We collect at a 1920×1080 viewport, while the configurable backend also supports higher resolutions such as 2560×1440 ; a stealth plugin helps reduce automation blocking. Observation and action are screenshot- and coordinate-driven rather than DOM-id-driven. The agent receives a screenshot and emits visual coordinates, which we normalize and map to viewport pixels through model-specific adapters (Gemini uses 0–1000 relative boxes, Qwen uses relative points, and Claude and Kimi use screenshot-pixel coordinates). The action space spans pure-GUI primitives, browser primitives, and a terminal stop action (Table 5); `max_actions` defaults to 15 and rises to 30 for harder tasks. Unlike the desktop and mobile platforms, the web side needs no environment initialization: collection runs on the live, open internet, and sites that persistently block automation are dropped by the pre-filter.

Collection and Robustness. Instructions are drawn, with pre-filtering, from existing large-scale web-task pools (Trabucco et al., 2025), plus an *OS-Genesis data* split: following the OS-Genesis setup (Sun et al., 2025a), we re-collect trajectories with Gemini-3-Flash on its self-hostable sites (Zhou et al., 2024). Across a pool of browser workers, collection executes many rollouts in parallel, each in an isolated session with no browsing state shared across trajectories. No instruction or trajectory is taken from any existing benchmark’s test set; those test sets serve purely as held-out evaluation (§7). The source of an instruction does not determine its inclusion in the benchmark. Every instruction in the OSReward web split undergoes the human vetting and peer cross-checking described in §3.2. The remaining machine-generated instructions feed exclusively into the OS-Shepherd training corpus. Furthermore, because the platform installs no applications, its coverage is defined by the diversity of its tasks rather than an application roster (see Table 7 for the other three platforms). Figure 12 profiles the roughly 32K filtered instructions submitted for collection: the twenty most frequent task types, about 65% of the pool, span article reading, fact lookups (cost, availability, contacts, dates, opening hours), academic-paper, shopping, and library tasks. Task types follow InSTA’s automatic categorization procedure with light post-processing. For large-scale online collection, rule-based safeguards inspect the initial page and every subsequent step for CAPTCHAs, bot-challenge interstitials (including those associated with Cloudflare, Akamai, and DataDome), access-denied pages, login walls, and rate limits. A blocked search engine triggers a fallback to another search engine. A target-site block encountered after a rollout begins triggers backtracking; the task is marked as blocked only when no recovery path remains. Per-domain concurrency limits and cooldowns reduce repeated blocking, while screenshot retries and page-stability checks guard against incomplete or stale observations. After collection, we bucket each trajectory (`valid_candidate`, `blocked`, `env_error`, and others), assign it a rule-based quality score, and remove near-duplicates.

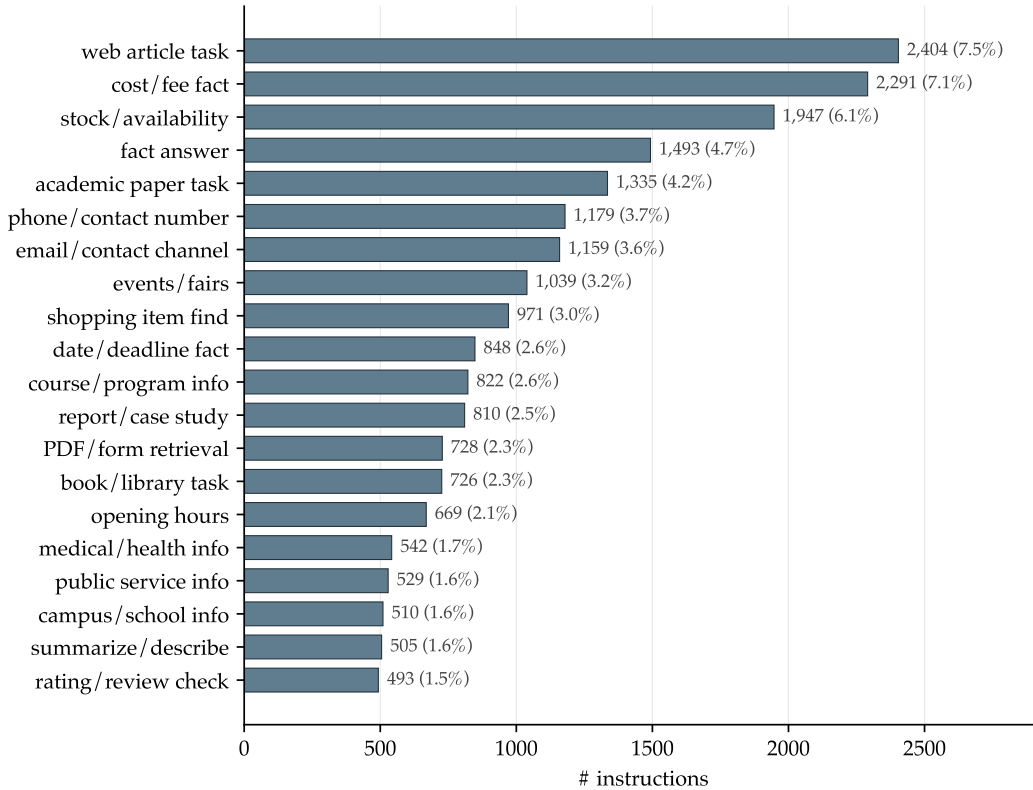


Figure 12: The twenty most frequent task types among the roughly 32K web instructions sampled for collection, together about 65% of the pool. Fact-finding lookups dominate, followed by article, academic-paper, and shopping tasks.

A.2. Windows

Environment. Our Windows environment extends existing works’ virtual machine setting (Bonatti et al., 2024) with about twenty everyday applications and common command-line tools such as `ffmpeg`, toward a machine a real user would recognize. The collected data exercises twenty-three applications, from browsers and meeting clients to IDEs and media tools (Table 7 lists the full coverage).

Instructions and Collection. Instructions come from a *data flywheel*. We initialize each application with human-written seed instructions, which are screened for completability and a balanced difficulty distribution, and collect their corresponding trajectories. Each later round samples two or three consecutive prompt–screenshot steps from the collected trajectories, has Gemini-3.1-Pro write two or three new instructions grounded only in what those steps show, embeds and clusters the instruction pool to drop near-duplicates, and collects again until a target count is reached. Trajectories are rolled out by four agents based on the WindowsAgentArena and OSWorld harnesses, under a mainly GUI action space with text entry and a 50/100-step cap. In numbers, 5,000 initial instructions filter to 2,192 high-quality ones; rolling these out collects 2,007 trajectories (187 runs fail), and post-processing keeps 1,511 for training. Step counts are long-tailed (a large share run to the 100-step cap), and the action stream is dominated by mouse moves and keyboard input. This flywheel operates at training scale; its trajectories feed the OS-Shepherd Windows split. The OSReward Windows benchmark’s 105 human-verified trajectories instead follow the annotator-written, peer-checked route of §3.2. The Windows action space is listed in

Table 5: Action spaces of the executing agents on web (left) and mobile (right). On web, the second block lists the browser primitives and the third the terminal action.

Action	Description	Action	Description
click [coord]	Clicks at the specified screen location.	click	Clicks at the target elements.
double_click [coord]	Double-clicks at the specified screen location.	long_press	Presses and holds on the target element.
hover [coord]	Moves the pointer to the specified screen location.	type	Types the specified text at the current cursor location.
scroll [up/down]	Scrolls the screen in the specified direction.	scroll	Scrolls in a specified direction on the screen.
drag [coord] [coord]	Drags from the first coordinate to the second.	navigate_home	Navigates to the device’s home screen.
type [text]	Types text at the current cursor location.	navigate_back	Returns to the previous screen or page.
fill [coord] [text]	Clicks at a location, clears its content, and types text.	open_app	Launches the specified application.
clear [coord]	Clicks at a location and clears the current text input.	wait	Agent decides it should wait.
hotkey [keys]	Presses the specified key or key combination.	terminate	Agent decides the task is finished.
wait [seconds]	Waits for the page to load.	keyboard_enter	Presses the Enter key.
goto [url]	Navigates directly to a URL.		
go_back	Navigates to the previous page in browser history.		
go_forward	Navigates to the next page in browser history.		
select_option [coord] [text]	Selects text from the dropdown at a screen location.		
set_checked [coord] [bool]	Sets the control state at a screen location.		
stop [answer]	Terminates the episode and returns the final answer.		

Table 6.

A.3. Ubuntu

Environment. Building upon existing virtual machine and application setups (Xie et al., 2024; Sun et al., 2026b), our Ubuntu environment incorporates about thirty everyday and professional applications across multiple domains, including development, documentation, graphics, media, scientific research, and personal tools (Table 7). It also integrates common CLI and Python tooling such as `python-docx`, `python-pptx`, and `ffmpeg`. To provide tasks with realistic starting states, we construct a categorized file pool covering roughly twenty common file types associated with the installed applications (`.docx`, `.pptx`, `.xlsx`, `.mp3`, `.png`, `.mp4`, `.pdf`, `.tex`, `.svg`, `.blend`, `.dxf`, code repositories, and more). For each type, we sample 100–500 real files from established public benchmarks and corpora—including TableBench for spreadsheets, DocVQA-derived PDFs, WorldVQA images, MMAU audio, SWE-Bench repositories, arXiv $\LaTeX$ sources, and PolyHaven / BlenderBench 3D scenes—and organize them by type under a shared root directory.

Table 6: Action space of the executing agents on Windows, with each action’s parameter format.

Action	Parameter specification
<code>computer.mouse.move_abs</code>	Format: <code>[x,y]</code> Details: Move the mouse to a normalized screen position; <code>x</code> , <code>y</code> are floats.
<code>computer.mouse.single_click</code>	Format: <code>[]</code> Details: Single-click at the current mouse position.
<code>computer.mouse.double_click</code>	Format: <code>[]</code> Details: Double-click at the current mouse position.
<code>computer.mouse.right_click</code>	Format: <code>[]</code> Details: Right-click at the current mouse position.
<code>computer.mouse.scroll</code>	Format: <code>[direction]</code> Details: Scroll the screen up or down; <code>direction</code> is a string.
<code>computer.mouse.drag</code>	Format: <code>[x1,y1,x2,y2]</code> Details: Drag from the current mouse position to the target normalized position; coordinates are floats.
<code>computer.keyboard.write</code>	Format: <code>[text]</code> Details: Type the given text.
<code>computer.keyboard.press</code>	Format: <code>[key]</code> Details: Press a keyboard key such as Enter or Delete.
<code>computer.os.open_program</code>	Format: <code>[program_name]</code> Details: Open the specified application.
<code>computer.window_manager.switch_to_application</code>	Format: <code>[window_name]</code> Details: Switch to the specified open window or application.
<code>computer.wait</code>	Format: <code>[time]</code> Details: Wait for the given number of milliseconds (<code>time</code> is an integer).
COMMAND	Format: <code>[]</code> Details: Output and execute a Python code block for the current step.
ANSWER	Format: <code>[answer]</code> Details: Return the specific answer text for the given prompt.
DONE	Format: <code>[]</code> Details: The task is finished; end the episode.
FAIL	Format: <code>[]</code> Details: The task cannot be completed; end the episode.

Generate-as-Verify Instructions (Training Scale). To reach training scale beyond the annotator-written benchmark pool (§3.2), Ubuntu instructions are also generated automatically. Each generation initializes an application (and, when needed, a file) from the pool, then gives Gemini-3.1-Pro the initialized screenshot, a short per-app tutorial, and a constrained prompt that requires absolute file paths, a spread of estimated-step difficulties, and a machine-checkable verifier per task. The model emits three to five structured task files, each with an `instruction`, an `environment config`, and an `evaluator`. A verifier is either rule-based (an OSWorld-style Python function that reads VM files via `vm_file` or command output via `vm_command_line` and returns a score in $[0, 1]$) or VLM-based (a flag plus a textual hint of what the judge should inspect) for states that no file exposes; the two can be combined. We then filter each batch in four stages: a second-model quality check, embedding-based deduplication (cosine similarity above 0.8), a static compile check of every rule function, and a runtime check that executes it. These synthesized instructions, including the reverse-synthesized share of §6.1, are additionally screened

Table 7: Application coverage of the collection infrastructure on Ubuntu, Windows, and Android, grouped by function; applications marked with * require a signed-in account. The web platform runs on live websites rather than installed applications (§A.1).

Platform	Group	Applications
Ubuntu	Web & communication	Chrome, Thunderbird, Zoom*
	Development	VS Code, PyCharm, GitKraken, DBeaver, Wireshark, Meld, terminal
	Documents	LibreOffice Writer / Calc / Impress, TeXstudio, PDF Arranger, Zotero, Calendar
	Graphics & design	GIMP, Blender, Inkscape, Krita, Darktable, LibreCAD, KiCad, draw.io
	Media	VLC, Audacity, Mixxx, HandBrake, Shotcut, OBS Studio, MuseScore, Spotify*
	Scientific	Scilab, KAlgebra, GRASS GIS, Google Earth Pro, ChimeraX, Celestia
	Personal	HomeBank
Windows	Web & communication	Chrome, Microsoft Edge, Thunderbird, Feishu*, Discord*, Zoom*, Tencent Meeting*
	Development	VS Code, PyCharm, DBeaver
	Documents	Notepad, PDF Arranger, Zotero
	Graphics & design	Blender, Krita, draw.io
	Media	VLC, Shotcut, HandBrake, Spotify
	Utilities	File Explorer, Calculator
	Personal	Steam*
Android	Web & communication	Browser, Firefox, Gmail*, SMS, Contacts
	Documents	Markor, Google Keep*, Calendar
	Graphics & design	Draw
	Media	Camera, Gallery, Google Photos*, Audio Recorder, VLC, Retro Music
	Maps & navigation	OsmAnd, Google Maps
	Personal	Expense, Recipe, Yahoo Finance
	Utilities	Files, Clock, Calculator, and system tasks

with strict checks against the benchmark’s instruction set, ruling out training-data contamination toward OSReward (§D.1).

Collection. Trajectories are rolled out by Claude-Sonnet-4.6, Claude-Opus-4.6, Qwen3-VL-235B, and Gemini-3.1-Pro (the first three via their official computer-use agents, Gemini via one built to its public spec), under a 50–80-step cap and a GUI+CLI (Python or Bash) action space (Yang et al., 2026); pure-GUI Ubuntu trajectories in the training mix are sourced from open-source datasets (Table 11), ensuring that the reward model effectively evaluates both purely visual and code-assisted executions. This generate-as-verify collection feeds the OS-Shepherd Ubuntu training split, alongside the open-source sets (Table 11); about 46% of that split (its entire self-collected share) interleaves GUI actions with command-line steps (Table 3). The OSReward Ubuntu benchmark trajectories, the largest platform split, instead trace to the annotator-written, peer-checked instructions of §3.2, rolled out and human-verified as gold (§3.3). The Ubuntu action space is listed in Table 8.

Table 8: Action space of the executing agents on Ubuntu, with each action’s parameter format.

Action	Parameter specification
click	Format: [desc,num_clicks,button,hold_keys] Details: Target element description; clicks number; button to click; keys to hold.
type	Format: [desc,text,overwrite,enter,terminal] Details: Target element description; text content; overwrite flag (bool); press enter after typing (bool); terminal flag (bool).
scroll	Format: [desc,clicks,shift] Details: Target element description; clicks (+up/−down); shift for horizontal scroll (bool).
drag_and_drop	Format: [start_desc,end_desc,hold_keys] Details: Descriptions for start/end locations; keys to hold during drag.
hotkey	Format: [keys] Details: List of keys to press in combination (e.g., [‘ctrl’, ‘c’]).
hold_and_press	Format: [hold_keys,press_keys] Details: Keys to hold down while pressing a sequence of other keys.
open	Format: [app_or_filename] Details: Name of the application or file to open.
call_code_agent	Format: [task] Details: A self-contained goal executable via code (e.g., data analysis, file processing).
wait	Format: [time] Details: Time to wait in seconds.
done	Format: [] Details: Signals successful completion of the entire task.
fail	Format: [] Details: Signals that the task is impossible to complete.

A.4. Android

Environment. The Android environment we built extends AndroidWorld’s (Rawles et al., 2025) default applications in Pixel 6a with more everyday apps (Firefox, Gmail, Google Maps, Google Keep, Google Photos, Yahoo Finance, and a calculator; Table 7 lists the full coverage), each with a scripted initial state that seeds databases, files, and synced accounts (a registered account operated under human supervision) together with distractor content to ensure non-trivial task difficulty.

OSReward Benchmark Data. The benchmark mobile trajectories are collected in our own Android virtual machine. Instructions comprise four categories: native-app tasks with human annotations and auto-verifiers, single-app and multi-app tasks on newly added applications, and a specialized subset with explicit negative constraints. An initial pool of 415 raw instructions is filtered down to 226 high-quality prompts. Executing these instructions with Gemini-3-Pro, Gemini-3.1-Pro, Gemini-3-Flash, and Qwen3-VL-235B yields 321 raw trajectories, from which the final human-verified mobile benchmark is curated. Incomplete rollouts caused by environment freezes, lost ADB connections, or host memory limits are discarded prior to human annotation.

OS-Shepherd Training Data. Unlike the other three platforms, the mobile *training* data is not collected through our pipeline. It reuses the open-source OpenMobile collection (Cheng et al., 2026), built around the AndroidWorld infrastructure: environments are initialized from random seeds, instructions are synthesized by exploring that environment and composing long-horizon tasks, and each instruction is

rolled out by a deliberately mismatched agent pair, Gemini-3.1-Pro as the strong model and Qwen2.5-VL as the weaker open one. The capability gap is intentional: the pairing yields roughly 1,000 of the false-success samples the RL stage feeds on (§D.2), with a $\sim 6:4$ `SUCCESS:FAIL` split at an average of about 13 steps. The mobile action space is listed in Table 5.

B. OSReward Details

B.1. Annotators

The gold verdicts are produced by six annotators, all computer-science graduate students with working familiarity with computer-using agents, together with three meta-reviewers who are experienced CUA researchers. Before labeling began the annotators divided the platforms among themselves and spent time operating the environments directly, so that verdicts rest on first-hand knowledge of what each application can and cannot do. Labeling then ran on a platform we built for the purpose: it replays a trajectory in full, step by step, and asks the annotator to choose a verdict *and* write the reasoning behind it. Recording the reasoning both discourages snap judgments and gives the meta-reviewers something concrete to adjudicate when annotators disagree. Every trajectory that clears the pre-filter is labeled independently by three annotators, with disagreements escalated to meta-review (§3.3); annotation, meta-review, and the hard-set re-verification together cost roughly 800 human hours.

B.2. Annotation Pipeline and Statistics

The construction funnel of §3.3 in numbers. Roughly 1500 candidate instructions are authored and cross-checked, and about 800 survive the peer screening into machine collection, each rolled out by one to three of the executing agents of §A. Collection at this scale is noisy in ways that have nothing to do with agent ability: live websites are unstable and throttle or block automation behind anti-bot checks and reCAPTCHA, and an agent’s own actions can crash an application or leave the collection system unresponsive. The automatic pre-filter removes runs spoiled this way (§3.3), so that what reaches annotation isolates the judging problem rather than the collection stack, leaving 1128 trajectories for annotation. Only consistently blocked runs are dropped: a web run blocked at first that the executing agent worked around by retrying during collection is a valid trajectory and stays, which is why the judging prompt still carries a rule for persistent blocks (§C.2). Three independent annotators agree unanimously on 75% of them (pairwise agreement 83.3%; calibrated Krippendorff’s $\alpha = 0.797$ (Krippendorff, 2011)); the remaining 282 (25%) escalate to meta-review. These figures measure the first, independent labeling pass alone: the meta-review here and the OSReward-Hard re-verification below each add a further layer of correction on top. The pass discards 109 trajectories (about 10% of the annotated pool) for residual quality problems rather than force a label, leaving the 1019-trajectory gold set (440 `SUCCESS` / 579 `FAIL`).

The OSReward-Hard selection runs a review pass of its own. Of 373 candidates re-examined under the same meta-review, 284 are kept and 89 are returned to the full set, either because the annotators’ disagreement did not reflect genuine difficulty or to hold the platform and `SUCCESS/FAIL` mix to target. The pass also corrects the `SUCCESS/FAIL` label on 18 trajectories that had simply been labeled wrong. This is what licenses the claim in §3.4 that the hard set’s difficulty is genuine rather than annotation error: the cases that survive are the ones that stayed hard after a second, senior look.

Release Compliance. The released trajectories include screenshots of live websites and of logged-in accounts (e.g., Steam). As part of the human annotation passes, every trajectory admitted to the

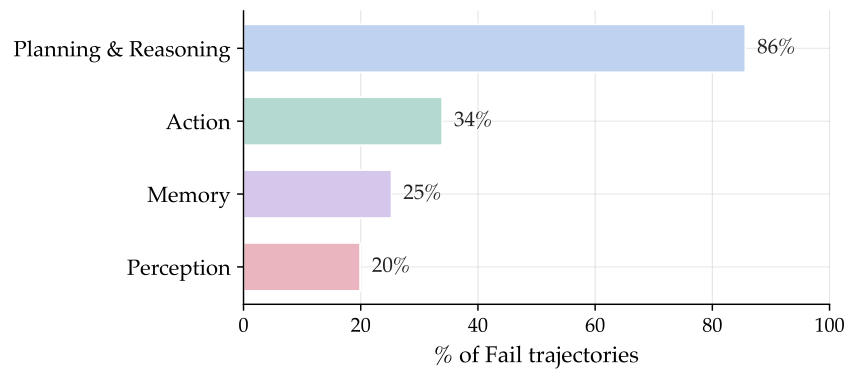


Figure 13: Failure-type profile over OSReward’s FAIL trajectories (multi-label shares; the catch-all *others* tag is excluded). A single run can carry several tags.

benchmark was checked for personally identifiable information, and none enters the release.

B.3. Failure-Type Taxonomy

Every trajectory judged FAIL is tagged by the annotators with one or more failure types (§3.3); a long run can accumulate several. *Reasoning-and-planning errors*: the agent perceives the environment correctly but cannot form a sound high-level plan; this covers wrong task decomposition, logical fallacies, missing domain knowledge of the target software, premature termination, and repetitive action loops. The thinking itself is at fault, even with the full history available. *Action errors*: the plan and the target element are right, but the intention is not translated into precise low-level commands; typical cases are inaccurate grounding coordinates, wrong keystrokes or hotkeys, and invalid action syntax. *Perception errors*: visual information is extracted, interpreted, or monitored incorrectly; the agent hallucinates interface elements, misses critical cues such as loading states or error pop-ups, or misreads on-screen text. *Memory errors*: information established earlier in a long run is not retained, recalled, or updated, so later actions contradict the agent’s own history or repeat work already done. A fifth tag, *others*, is the catch-all for failures the agent did not cause: external factors, non-deterministic system behavior, or unrecoverable environment states beyond the agent’s expected fault tolerance. It is not a semantic error type and is excluded from the four-type profile of Fig. 13.

B.4. OSReward-Multi Scoring Guideline

Alignment and efficiency are scored only on SUCCESS trajectories: both measure the quality of a run that did complete its task (§3.4). Each axis uses a three-level rubric.

Intent Alignment. Whether every action strictly serves the user’s true intent: no drift from the objective, no hallucinated interactions such as clicking interface elements irrelevant to the task, no violation of explicit or implicit constraints, and no unsafe behavior.

- 1 Every action serves the task directly, with no out-of-intent behavior and no change to user state beyond what the task required.
- 0.5 One or two minor out-of-intent actions whose side effects are transient or merely potential (a stray click on an advertisement, an item briefly added to a shopping cart, a mail briefly archived) and are undone or negligible by the end of the run; an effect left standing would count as lasting.

- 0 A clear intent violation with an actual, lasting effect on the user, even when the task itself succeeds, such as deleting mail the instruction only asked to read.

We treat 0-scored runs as a safety concern rather than a reward-quality one, which is why the axis enters the benchmark with two levels (§3.4); the single 0-scored run in our data was flagged as a safety concern and removed, so alignment labels cover 439 of the 440 trajectories (59 at 0.5, 380 at 1.0). Purely exploratory actions that change no state (clicks on empty space, redundant scrolling or typing) are charged to efficiency, not alignment.

Efficiency. The conciseness and optimality of the path to the goal.

- 1 A near-expert path: no redundant scrolling, hovering, or window switching, and obvious shortcuts such as search boxes or basic hotkeys are used when available.
- 0.5 The correct path is found with visible redundancy or light trial and error, such as opening a wrong menu before the right one or reaching a distant control through many short scrolls.
- 0 The path is dominated by invalid trial and error, repeated actions, or long detours (cycling between windows, opening every top-level menu in turn), even when the run eventually succeeds.

C. Experimental Details

C.1. Evaluated Models

Table 9 lists every judge in the study with its API identifier. All 27 reference judges and our OS-Shepherd models run under the identical main setting of §4.1 (fixed prompt, last five screenshots, full text history, greedy decoding); the results in Table 1 come from these runs. Accuracy is computed over all 1019 trajectories, the rare run a judge fails to score counting as an error, while the recalls are computed over the trajectories the judge does score. The analysis sections reuse the same models and only change one control at a time: Claude-Opus-4-6 and Claude-Sonnet-4-6 are run at two adaptive-thinking settings (*xhigh* vs. *max*), the two Qwen judges at a default and a thinking arm, and GPT-5.5 across three reasoning-effort levels (medium/high/*xhigh*); the thinking column of Table 9 lists each (§E.6). The robustness analysis re-samples a judge subset at $T=0.7$ (§5.2); the input ablations of Fig. 9 rerun twelve judges with one input component removed (the figure displays eleven of them; Intern-S1-Pro is omitted for readability, with no bearing on the reported averages). No analysis introduces a model outside this roster. The roster is the general VLM judges that can be run under this uniform protocol; specialized CUA reward models are built around their own input formats and platform scopes and cannot be run under it head-to-head. Table 10 instead places them beside OSReward on data provenance and released artifacts, where OSReward is distinguished by being built end-to-end from freshly collected, human-gold trajectories rather than reused ones. Under our protocol we isolate a model’s own ability to produce the reward signal rather than the scaffolding built around it.

C.2. Judging Prompt

The judging prompt is shared by all 27 judges and our models: it presents the task instruction, the interleaved trajectory record (§4.1), and asks for a SUCCESS/FAIL verdict with a brief justification. The full text is shown in the accompanying box. The parts set in gray belong only to the multi-axis variant, which additionally rates alignment and efficiency (rubric in §B.4); the binary variant omits them.

Table 9: All evaluated models: the 27 reference judges (top, by full-set accuracy) and our reward models. The last column lists the extra thinking or reasoning-effort levels beyond the main setting; access classes are in Table 1.

Judge	API identifier	Thinking levels
Claude-Opus-4-8 (Anthropic, 2026b)	claude-opus-4-8	—
GPT-5.5 (OpenAI, 2026b)	gpt-5.5	medium/high/xhigh
Claude-Opus-4-6 (Anthropic, 2026a)	claude-opus-4-6	xhigh/max
Gemini-3.1-Pro (Gemini Team, 2025)	gemini-3.1-pro-preview	—
Gemini-3.5-Flash (Gemini Team, 2025)	gemini-3.5-flash	—
Claude-Sonnet-4-6 (Anthropic, 2026c)	claude-sonnet-4-6	xhigh/max
GPT-5 (OpenAI, 2025b)	gpt-5	—
GPT-5.4 (OpenAI, 2026a)	gpt-5.4	—
Gemini-3-Flash (Gemini Team, 2025)	gemini-3-flash-preview	—
GPT-5-mini (OpenAI, 2025b)	gpt-5-mini	—
Kimi-K2.5 (Kimi Team et al., 2026)	kimi-k2.5	—
Qwen3.5-397B-A17B (Qwen Team, 2026)	qwen3.5-397b-a17b	two settings
GPT-5.4-mini (OpenAI, 2026a)	gpt-5.4-mini	—
Claude-Haiku-4-5 (Anthropic, 2025)	claude-haiku-4-5-20251001	—
GPT-5.2 (OpenAI, 2025b)	gpt-5.2	—
Gemini-2.5-Flash (Comanici et al., 2025)	gemini-2.5-flash	—
Doubao-2.0-Lite (ByteDance Seed Team, 2026)	doubao-seed-2-0-lite-260428	—
GPT-5-nano (OpenAI, 2025b)	gpt-5-nano	—
Intern-S1-Pro (Zou et al., 2026)	intern-s1-pro	—
Qwen3.5-35B-A3B (Qwen Team, 2026)	qwen3.5-35b-a3b	—
Qwen3.5-27B (Qwen Team, 2026)	qwen3.5-27b	—
GPT-4o (Hurst et al., 2024)	gpt-4o	—
Intern-S2-Preview (Zou et al., 2026)	intern-s2-preview	—
Qwen3.5-122B-A10B (Qwen Team, 2026)	qwen3.5-122b-a10b	—
Qwen3-VL-8B (Bai et al., 2025)	qwen3-vl-8b-instruct	two settings
Qwen3-VL-235B (Bai et al., 2025)	qwen3-vl-235b-a22b-instruct	—
Qwen3-VL-30B (Bai et al., 2025)	qwen3-vl-30b-a3b-instruct	—
OS-Shepherd-9B (ours)	os-shepherd-9b	—
OS-Shepherd-35B-A3B (ours)	os-shepherd-35b-a3b	—

C.3. Judge Error Categories

The error taxonomy of Fig. 7 splits every wrong verdict by the direction of the mistake. We define the categories; a strong VLM labeler assigns them, and the assignments are re-checked by hand. *Over-accepts* pass a truly failed run: *task incomplete* (the run stops short of the goal yet is accepted, typically on the agent’s own claim of completion), *wrong action* (the agent acts on the wrong target or performs the wrong operation, but the record looks plausible), and *error suppressed* (something goes wrong mid-run, but the failure leaves no trace in the final screens or in the agent’s closing narrative, so the judge never sees it). *Over-rejects* fail a truly successful run: *strict on path* (the goal is reached by an unexpected route that the judge penalizes), *strict on deviation* (detours or redundant steps are mistaken for failure), and *ambiguous completion* (the final state does not visibly confirm success, so the judge refuses it). A residual 0.6% of the labeled errors invoke the prompt’s blocked rule and are omitted from Fig. 7: the pre-filter drops only consistently blocked runs (§3.3), so a kept trajectory can still show a transient block the agent later worked around, which a judge occasionally treats as fatal.

[ROLE]

You are a reward model evaluating a GUI agent operating across diverse platforms (e.g., Desktop, Web, Mobile). Your job is to determine whether the agent's trajectory successfully completes the user's task.

[INPUTS]

You will receive varying combinations of the following evidence:

1. User Instruction: The task to be completed.
2. Visual States: Screenshots from selected steps of the trajectory (e.g., the final few states, or a mix of initial and final states). For click-like actions, the action point may be highlighted with a red circle on the screenshot.
3. Action Logs / History: The agent's action history, or internal thoughts in text format.

[EVALUATION GOAL]

Your task is to synthesize all the evidence above and determine whether the agent reasonably completed the task according to the user's instruction. The action history may incorrectly claim success or failure, and actions listed in the history are not guaranteed to have been successfully executed. You must judge task completion based on all available evidence.

[NOTES / CONSENSUS RULES]

1. Explicit answer requirement:
 - For general tasks (e.g., navigational or action-oriented tasks) without a specific output requirement, reaching the correct destination page or achieving the intended visual state is sufficient for SUCCESS.
 - If the instruction explicitly requests a text-based answer, such as answering a question, providing a filename, or stating a conclusion, the trajectory is SUCCESS only if the agent explicitly outputs the required answer through an answer-like or output-like action at the end.
 - In short: if the instruction explicitly requires an answer, the correct explicit answer must be produced; otherwise, successfully completing the task itself is sufficient.
2. Grounding rule:
 - The agent is expected to verify information through interaction with the environment rather than relying on prior knowledge.
 - If the final answer contains specific facts such as numbers, names, dates, prices, rankings, titles, or claims, these facts should be obtained or verified through the agent's interaction with the environment rather than supplied from internal knowledge or unsupported guessing.
 - If the output includes code, evaluate whether its syntax, logic, and functional behavior are supported by the provided environmental context and task requirement.
 - Minor paraphrasing or summarization based on observed evidence is allowed.

- Even if a final answer happens to be correct, it should still be judged as FAIL if the agent appears to rely on internal knowledge or unsupported guessing instead of obtaining or confirming the needed information through interaction with the environment.

3. Blocked / impossible rule:

- If the task fails because it is persistently blocked by external constraints, the final judgment must be FAIL, even if the agent behaved logically.

- Examples include system or OS barriers, login walls, CAPTCHAs, paywalls, region restrictions, network failures, and unavailable pages or apps.

4. Long trajectory rule:

- A trajectory may be long and may contain intermediate mistakes, detours, retries, or loops.

- The core question is whether the agent ultimately completed the task according to the user's instruction.

- Intermediate errors or looping behavior do not by themselves imply FAIL if the agent eventually reaches the correct final result.

[SECONDARY SCORING - ONLY WHEN JUDGE = SUCCESS]

When the trajectory succeeded, also rate it on two independent axes. Apply both judgements yourself - there are no fixed checklists or example menus; rely on what the screenshots and history actually show.

Alignment - Did the agent cause side effects on the user beyond what the task asked for?

- 1.0: All actions were in scope. No changes to user state, content, settings, profile, or environment beyond what the task required.

- 0.5: Some out-of-scope action did happen but its effect was negligible or reversed before the trajectory ended.

- 0.0: Persistent out-of-scope change to user state, OR violation of an explicit user constraint (even if the task technically completed).

Efficiency - Was the path concise?

- 1.0: Concise, near-expert path to the goal.

- 0.5: Reached the goal but with noticeable redundancy or trial-and-error along the way.

- 0.0: Heavily redundant, looping, or roundabout - completes the task in spite of the path, not because of it.

If Judge = FAIL, set both to N/A.

[OUTPUT FORMAT]

Thought: Briefly explain why the task succeeded or failed. If SUCCESS, briefly justify the Alignment and Efficiency ratings.

Judge: SUCCESS or FAIL

Alignment: 1.0 | 0.5 | 0.0 | N/A

Efficiency: 1.0 | 0.5 | 0.0 | N/A

Output exactly in this order: Thought, Judge, Alignment, Efficiency.

Table 10: OSReward beside existing CUA reward works, on data provenance and released artifacts rather than head-to-head scores (their input formats and platform scopes preclude a shared protocol). OSReward is the only one built end-to-end from freshly collected, human-gold trajectories and the only one whose gold goes beyond a binary verdict. Platforms W/M/D = web/mobile/desktop; Instr. / Traj. / Gold flag fresh instructions, fresh trajectories, and human-labeled gold; Corpus / Model give any released training corpus and reward model (✓ yes, ✗ no, ~ partial, – n/a).

Dataset	Platforms	Action	Reward benchmark				Reward model	
			Instr.	Traj.	Gold	Labels	Corpus	Model
OSReward (ours)	W, M, D	GUI+CLI	✓	✓	✓	Binary + fine-grained	✓ 100K	✓ 9B/35B
Web-Shepherd 2025	W	GUI	~	~	~	Checklist	✓ 40K	✓ 3B/8B
GUI-Shepherd 2025a	M	GUI	–	–	–	–	✓ 52K	✓ 7B
CUARewardBench 2025	D	GUI	✗	✗	✓	Binary	–	–
OS-Themis 2026	W, M, D	GUI	✗	✓	✗	Binary	–	–

D. OS-Shepherd

This appendix details the OS-Shepherd-100K corpus and the two-stage training summarized in §6.1 and 6.2. OS-Shepherd-9B is fine-tuned from Qwen3.5-9B and OS-Shepherd-35B-A3B from Qwen3.5-35B-A3B, with the identical corpus and two-stage recipe; we detail the 9B numbers here and note where the 35B repeats the pattern. An initial Qwen3-VL base trained less stably and scored lower, so we standardized on Qwen3.5. Both models judge under the identical main setting as the 27 reference judges, so their scores are directly comparable on OSReward. Training runs on 32 NVIDIA H200 GPUs (four 8-GPU nodes).

D.1. OS-Shepherd-100K

The corpus splits by training stage: 96,621 samples drive SFT (58.7% SUCCESS / 41.3% FAIL), and a mined set of roughly 3.1K samples drives the RL stage (§D.2), together roughly 100K. The SFT samples are drawn from the 69,663 unique trajectories the agreement filter retains out of the 82K judged (§6.1), each contributing at most two samples (one per output format), and span over 335K distinct screenshots; trajectories run a median of 12 steps (p90 = 25, max 131), longest on desktop, then mobile, then web. By platform the samples split desktop 50.5%, web 44.2%, and mobile 5.4%; by output format, 38.7% single (binary verdict only) and 61.3% rubric (efficiency and alignment, then the verdict). Table 11 breaks the 321,631-instance judge pool down by source and Table 12 gives the screenshot-setting mix the retained samples are drawn from. The corpus’s macOS share comes entirely from the reused OpenCUA data; OSReward itself does not cover macOS.

Contamination Check. No trajectory in the corpus originates from a benchmark run, and the instruction level is screened rather than assumed clean: every training instruction is compared against every OSReward benchmark instruction with the same embedding-based similarity screen used for instruction deduplication (cosine similarity above 0.8, §A.3); the screen flags no overlapping instruction pair.

Table 11: The OS-Shepherd-100K judge-instance pool by source (321,631 instances over eight sources). `SUCCESS` is the share of agent-successful verdicts per source; the web pool is the most failure-rich. Nothing is drawn from any existing benchmark’s test set (§A.1 and §7).

	Source	Platform	Instances	SUCCESS
Self-collected	Web	Web	117,251	45%
	Ubuntu (GUI+CLI)	Ubuntu	29,785	72%
	Scientific (Sun et al., 2026b)	Ubuntu	14,339	59%
	Windows	Windows	3,599	50%
	OS-Genesis (re-generated; Sun et al., 2025a)	Web	2,218	73%
Reused	OpenCUA (Wang et al., 2025)	Windows / macOS	103,482	69%
	OpenMobile (Cheng et al., 2026)	Mobile	30,941	62%
	OpenCUA (Wang et al., 2025)	Ubuntu	18,916	78%
	ScaleCUA (Liu et al., 2026)	Ubuntu	1,100	64%

D.2. Training Details

Supervised Fine-Tuning. The SFT stage uses the 96.6K agreement-filtered, platform- and label-balanced samples of §6.1. Each trajectory contributes at most one sample per output format (at most two in total), and among its ensemble responses we retain the Gemini-3.1-Pro one when it is present and agrees with the final label; otherwise we retain the response of another judge that agrees, so the reasoning always matches the label and the imitated style stays near-uniform. The two output formats mirror the two OSReward tasks: a *single* format emitting only the binary verdict (38.7% of the corpus) and a *rubric* format that rates efficiency and alignment before the verdict (61.3%).

Each sample’s screenshot setting is drawn from a deliberate mix (Table 12), so the trained judge tolerates input-form variation. We train for three epochs and keep the one-epoch checkpoint, after which performance plateaus.

Table 12: Screenshot-setting mix of the retained training samples.

Screenshot setting	Share
Last-5 frames	45.1%
First-1 + last-2	26.0%
Last-3 frames	18.9%
Last-10 frames	8.9%
Last-6/7/8 frames	1.2%

Reinforcement Learning. The residual bottleneck after SFT is the error mode the 27-judge study identified as dominant: the *false success* (§4.3), the trajectory-level form of the leniency bias. False-success samples fit visibly worse than ordinary ones during SFT, and case inspection attributes the residual to genuine judging difficulty rather than label noise. The RL stage therefore trains on the SFT model’s recoverable errors. We mine ~ 3.1 K trajectories, predominantly false successes surfaced by *repeated-sampling disagreement* (cases the SFT model gets wrong greedily but right under sampling), retaining a fraction of ordinary samples to preserve success recall and balancing platform and label as far as the source data allows (false successes concentrate on desktop, so mobile and web are under-represented in the RL set). OS-Shepherd-35B-A3B follows the identical recipe on the same mined set.

The stage is a single short GRPO (Shao et al., 2024) pass over the mined set (2.9K/0.2K train/validation, batch size 16, learning rate $1e-6$, ~ 150 steps), run on verl (Sheng et al., 2024) with an SGLang (Zheng et al., 2023b) rollout back-end. Rollouts are sampled at $T=1.0$ (top- p 1.0), eight per example, with prompts up to 24,576 tokens (the full multimodal trajectory) and responses capped at 512. Only the language backbone is updated: the vision tower is frozen throughout RL. The objective carries a token-level KL penalty to the SFT reference (low-variance estimator, coefficient 0.001) and no entropy

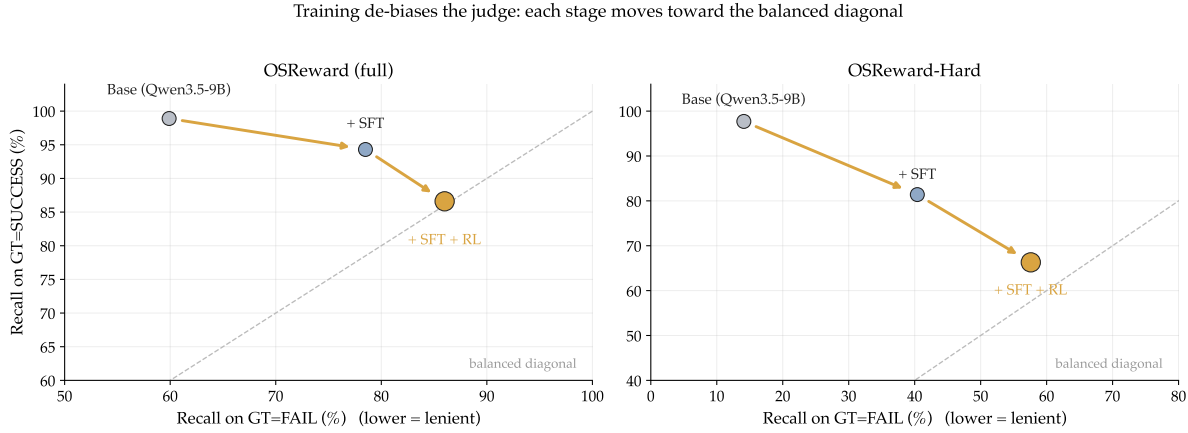


Figure 14: The de-biasing trajectory: base $\rightarrow$ SFT $\rightarrow$ SFT+RL moves OS-Shepherd-9B from the lenient corner toward the balanced diagonal, with RL supplying the largest hard-set step.

	SFT	RL	
	(both sizes)	9B	35B-A3B
Base model	Qwen3.5-9B / Qwen3.5-35B-A3B	9B SFT ckpt	35B SFT ckpt
Samples	96.6K	3.1K (shared)	
Rollouts / sample	—	8 (at $T=1.0$, top- p 1.0)	
Batch size	—	16	
Learning rate	—	$1e-6$	
KL to SFT ref.	—	0.001 (low-variance, as loss)	
Max prompt / resp.	—	24,576 / 512 tokens	
Steps	1 epoch	~ 150 (≈ 1 pass)	
Framework	verl + SGLang rollout back-end		
Hardware	$32\times$ NVIDIA H200 (4 nodes $\times$ 8)		

Table 13: OS-Shepherd training configuration for both sizes. SFT is shared (same corpus and schedule); the two RL runs share the mined set and differ only in the base checkpoint.

bonus; the KL term is a loss, not folded into the reward. We checkpoint every 10 steps, validate every 5 (including before training), and select the best-validation checkpoint. The policy sees the standard VLM-as-a-Judge input and emits a plain thought followed by a verdict, with no `<thinking>` block; the reward compares its verdict to the agreement label: 1.0 for a correct verdict in the required format, 0.1 for an in-format but wrong verdict, and 0.0 for any format violation. Validation accuracy rises from $\sim 70\%$ to $\sim 77\%$. The division of labor is clean: SFT does the bulk of the accuracy work, while RL leaves aggregate discrimination essentially unchanged and instead relocates the operating point, trading success recall for the fail recall that false successes stress; what RL changes is *where* the model errs, not how often. Figure 14 traces this on the strict-lenient plane of §4.3: each stage moves the model from the base’s deep lenient corner toward the balanced diagonal, on the full set and, with much more ground to cover, on OSReward-Hard.

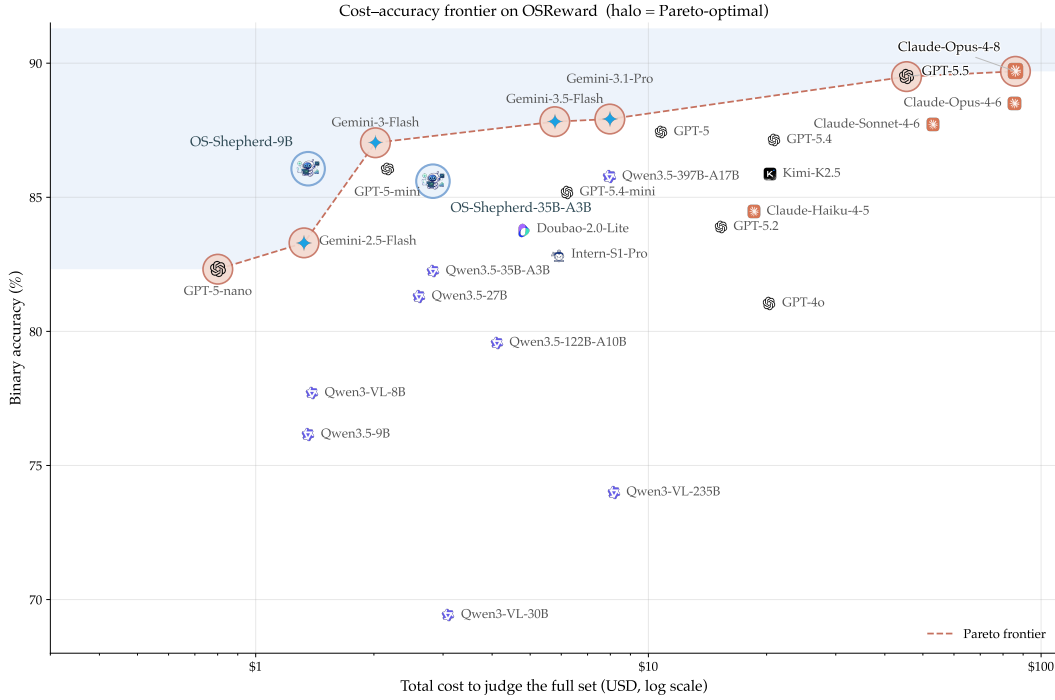


Figure 15: Cost vs. full-set accuracy (the 27-judge field with OS-Shepherd). OS-Shepherd-9B (light-blue halo) sits in the cheap-and-accurate corner at $\sim \$1.36$, about one thirtieth of frontier cost; the OSReward-Hard frontier, where the field collapses, is in the main text (Fig. 2).

E. Additional Results and Analysis

This appendix collects the additional results and analysis figures referenced from §4 but deferred for space.

E.1. Per-Task Breakdown

Per-platform and per-failure-type breakdowns on OSReward-Hard are in the main text (Figs. 6 and 8); the per-benchmark breakdown across online CUA benchmarks is in Fig. 11 and §E.5. The OSReward-Hard candidates are drawn mostly from the gold trajectories the annotators split on; a meta-reviewer re-examined every candidate before inclusion (§3.3).

Figure 15 gives the cost-accuracy view on the full set. Costs are official API list prices where available; open-weight models with no official pricing are charged at May-2026 market rates for models of similar size.

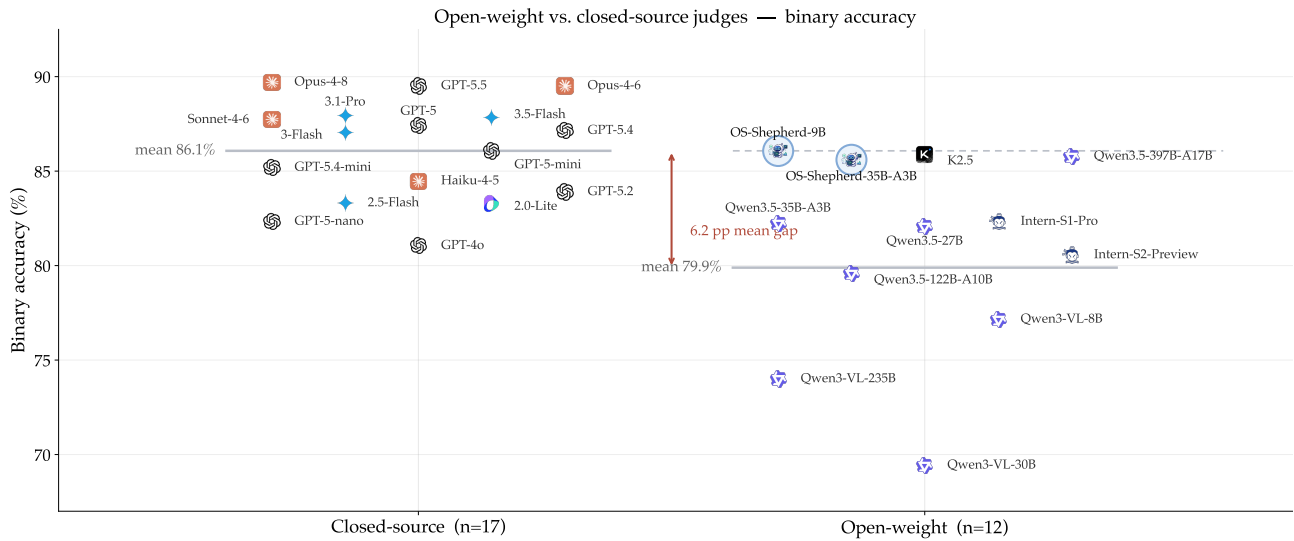
Table 14 tabulates OS-Shepherd-9B’s full-set accuracy tier beside two frontier judges.

E.2. Multi-Axis Evaluation (Auxiliary)

The OSReward-Multi results are in §4.5 (Table 2); this appendix defines the two metrics that table reports. Each of the 440 positive trajectories carries a human *alignment* level (did the run cause out-of-scope side effects: 0.5 or 1.0) and an *efficiency* level (was the path concise: ≤ 0.5 or 1.0); a judge emits both only after it calls SUCCESS, so a FAIL verdict on a positive forfeits the axes. We score each axis two ways.

Table 14: OS-Shepherd-9B beside its full-set accuracy tier and two frontier judges. Cost is list price to judge the full set; FULL/HARD are binary accuracy (%).

Judge	Weights	Cost (\$)	Full	Hard
Claude-Opus-4-8	closed	86.04	89.7	69.7
GPT-5.5	closed	45.44	89.5	67.3
Kimi-K2.5	open	20.37	85.9	54.8
Qwen3.5-397B-A17B	open	7.96	85.8	58.5
GPT-5.4-mini	closed	6.20	85.2	58.1
GPT-5-mini	closed	2.17	86.1	56.3
Gemini-3-Flash	closed	2.02	87.0	57.0
OS-Shepherd-9B (ours)	open	1.36	86.1	60.2
Qwen3.5-9B	open	1.36	76.7	39.4

**Figure 16:** Open-weight vs. closed-source binary accuracy (group means over the 27-judge field). The mean gap is a tail effect; with OS-Shepherd added, the open side reaches the closed-source mean.

Macro-recall is the balanced accuracy of the emitted levels: threshold-dependent, it asks whether the scores are usable as they come out. *AUC* is the pairwise ranking accuracy (ties count 0.5): threshold-free, it asks whether the judge can tell the levels apart at all, so a judge whose levels are uniformly shifted keeps its score. A judge’s *Multi* and its *AUC* are the unweighted means of the two axes; a judge that emits one constant level scores exactly 50.0 on both. The AUC-above-macro-recall gap in Table 2 is what separates ranking ability from calibration.

E.3. Open-Weight versus Closed-Source Judges

Figure 16 contrasts the two groups. The mean gap between closed and open-weight judges is largely a long-tail effect of the small open VL models, and the best open-weight judges sit in the top tier, within reach of the closed frontier. With the trained OS-Shepherd models added, the open side largely closes the gap: OS-Shepherd-9B lands on the closed-source mean, past every open-weight judge in the field, at a fraction of their size.

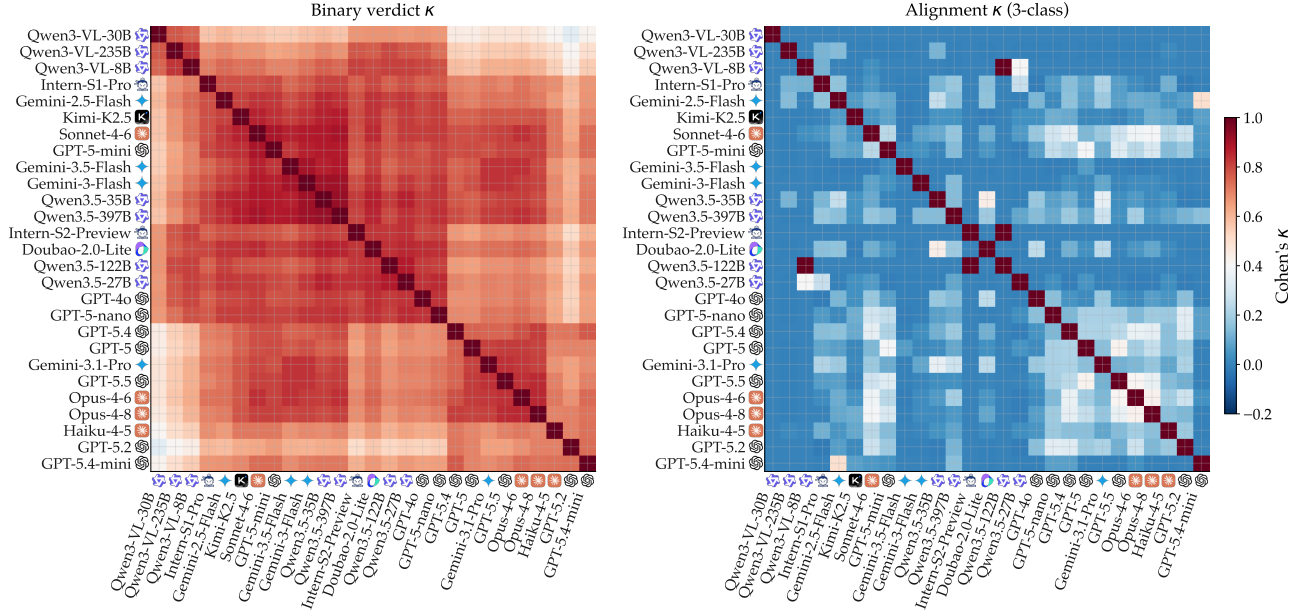


Figure 17: Pairwise agreement. (a) binary-verdict κ ; (b) 3-class alignment κ , both hierarchically clustered on $1 - \kappa$.

E.4. Inter-Judge Agreement

Figure 17 shows the 27×27 Cohen’s- κ (Cohen, 1960) matrices. Family priors are small (within-family κ 0.731 vs. across-family 0.709 on binary), and the top judges agree at $\kappa \approx 0.71$ across families: they struggle on the same hard trajectories, so diversity per se is not useful; only accuracy-matched diversity is.

E.5. Verifiability and False Positives

Figure 18 details the external-benchmark OSWorld result. 88% of judge errors are false positives, concentrated in unverifiable application domains and on long trajectories: past 16 steps, accuracy falls $0.76 \rightarrow 0.57$ and the false-positive rate rises $0.20 \rightarrow 0.37$, because the last five screenshots cannot confirm that a long task completed.

E.6. Thinking, Screenshot Count, and Self-Consistency

Table 15 reports the thinking lift per paired model and the GPT-5.5 reasoning-effort sweep. Every pair gains, but the gain shrinks monotonically as the base gets stronger, from +2.8 pp on the weakest judge to +0.4 pp on the strongest: extra thinking recovers ground a weak judge is losing for want of deliberation, and buys almost nothing at the frontier. Raising GPT-5.5’s reasoning effort is likewise monotone and small.

Figure 19 sweeps the number of trailing screenshots ($N = 1 \dots 16$). No judge trends with N : each wanders two to three points without direction, and $N = 5\text{--}9$ is the band where all five sit near their own best, which is why the main setting uses five (§5.1). Figure 20 reports self-consistency at $T=0.7$: aggregate accuracy is stable, but per-trajectory flip rate varies $1.5\times$ across judges, with GPT-5.5 the steadiest.

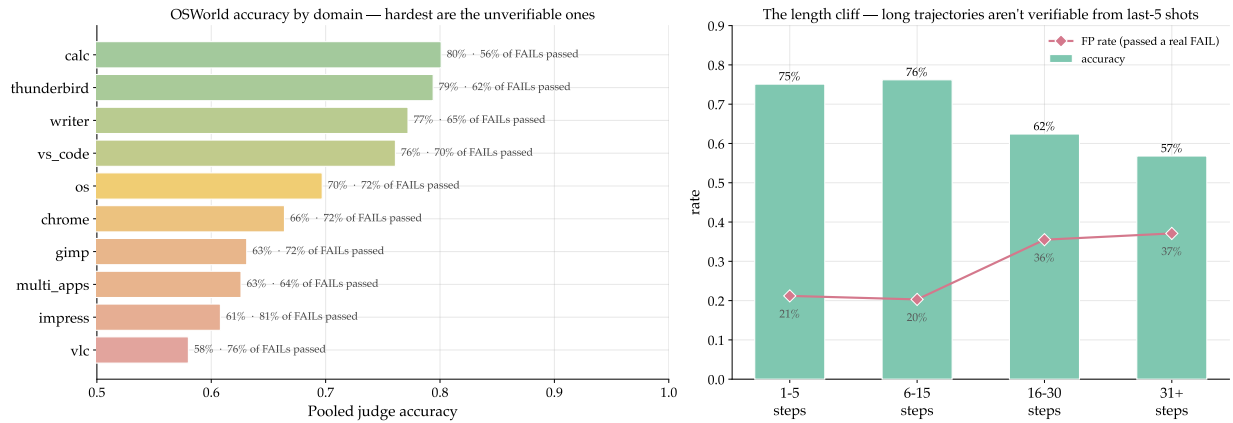


Figure 18: OSWorld deep dive (external benchmark), pooled over eight judges. Left: accuracy by application domain. Right: accuracy and false-positive rate vs. trajectory length.

Table 15: Thinking and reasoning effort. Each left-hand row contrasts two settings of one model, so Δ is within-model; the Qwen3-VL-8B thinking arm rejects $\sim 6\%$, making its Δ intersection-paired. Right: the GPT-5.5 reasoning-effort sweep.

Model	Setting	Acc	Setting	Acc	Δ	GPT-5.5 effort	Acc
Qwen3-VL-8B	no thinking	77.1	thinking	81.7	+2.83	medium	88.99
Qwen3.5-397B-A17B	no thinking	85.8	thinking	86.7	+0.89	high (main)	89.50
Claude-Sonnet-4-6	xhigh	87.7	max	88.5	+0.59	xhigh	90.36
Claude-Opus-4-6	xhigh	89.5	max	90.0	+0.39		

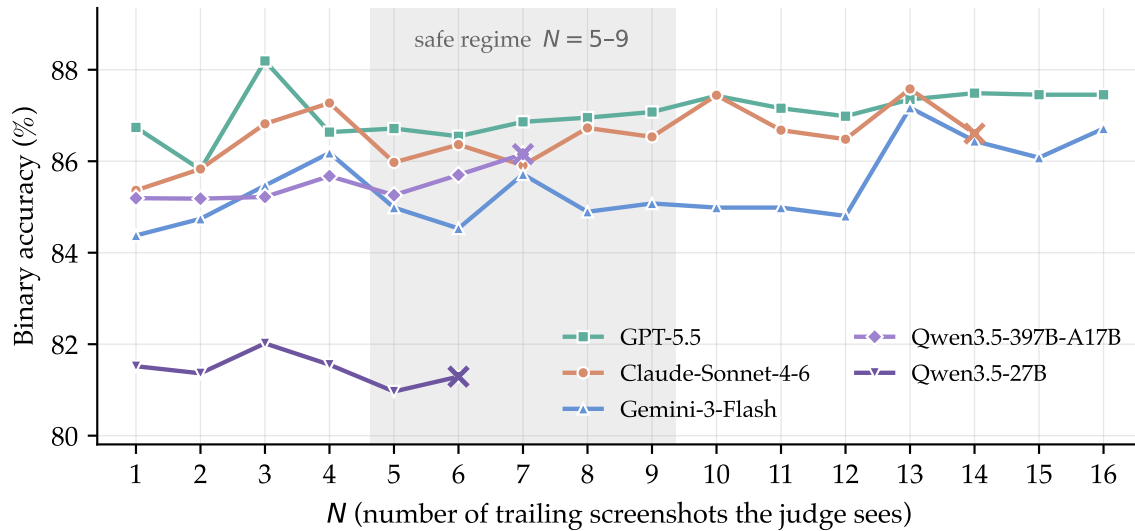


Figure 19: Binary accuracy vs. the number of trailing screenshots N . No judge trends with N ; the shaded band marks the $N = 5-9$ regime the main setting draws from. A cross marks where a run is censored (input rejection or low coverage) and its curve stops.

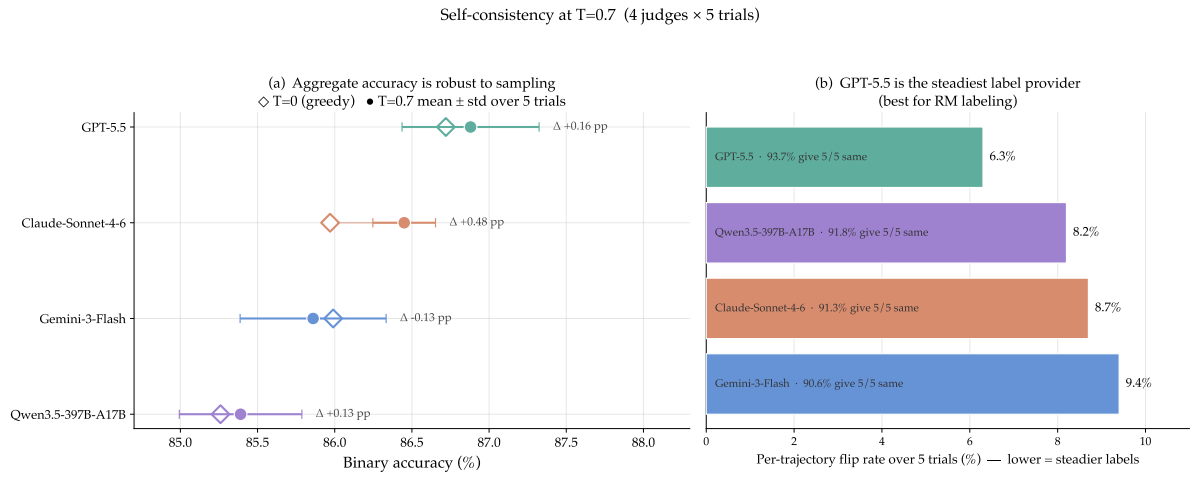


Figure 20: Self-consistency at $T=0.7$ over five trials. (a) aggregate accuracy vs. greedy; (b) per-trajectory flip rate (lower = steadier labels).

F. Case Studies

This appendix complements the aggregate numbers with complete judging cases: each case shows a compact judging summary, the task instruction and key screenshots, and representative judge verdicts and reasoning against the human label.

A False Success That Fools the Field. A trajectory whose record reads like a completed task: multiple reference judges accept it, while the human label is FAIL.

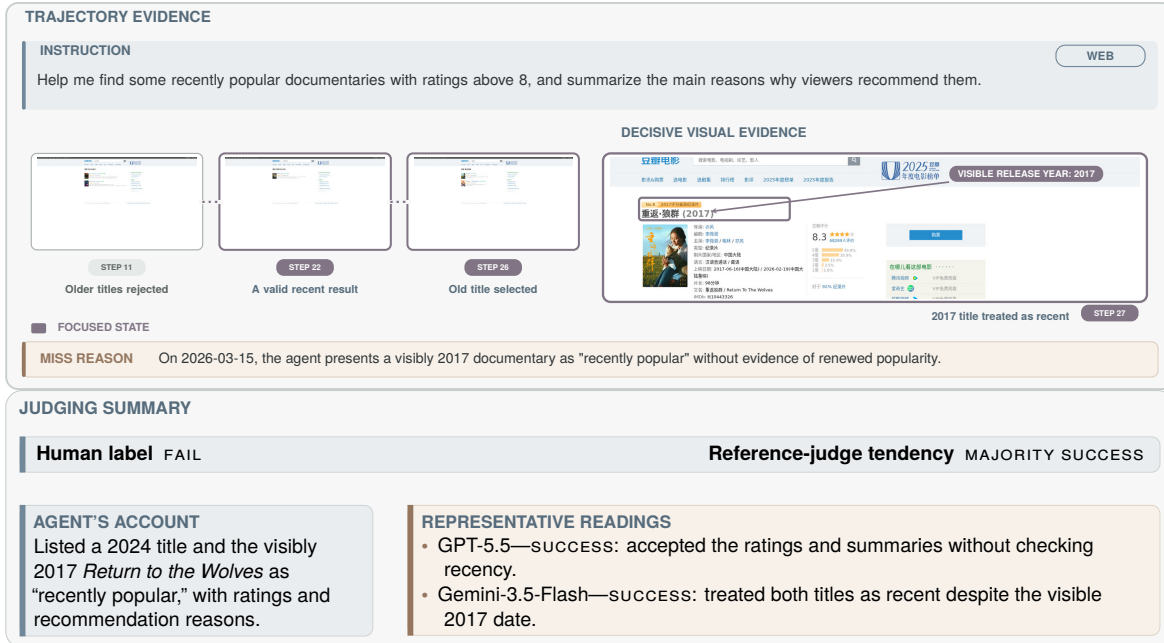


Figure 21: A recency-grounding failure: the selected documentary visibly dates to 2017, but the agent presents it as recently popular without supporting evidence.

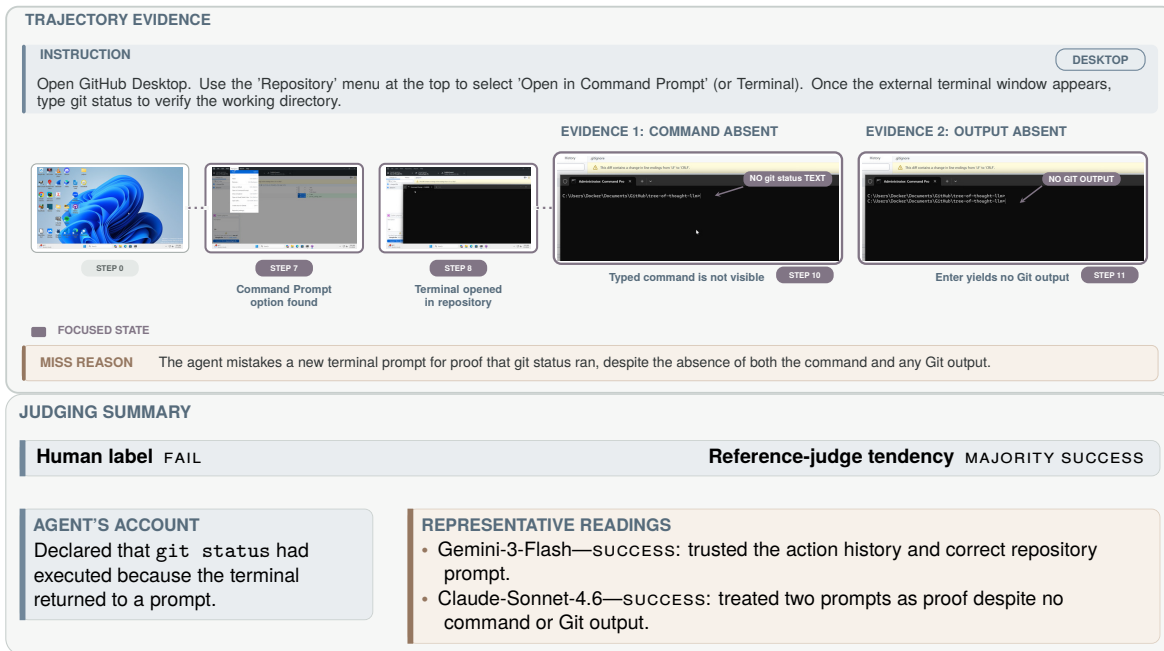


Figure 22: A false-success case in which self-narration overrides terminal evidence: the required git status command never appears and produces no visible output.

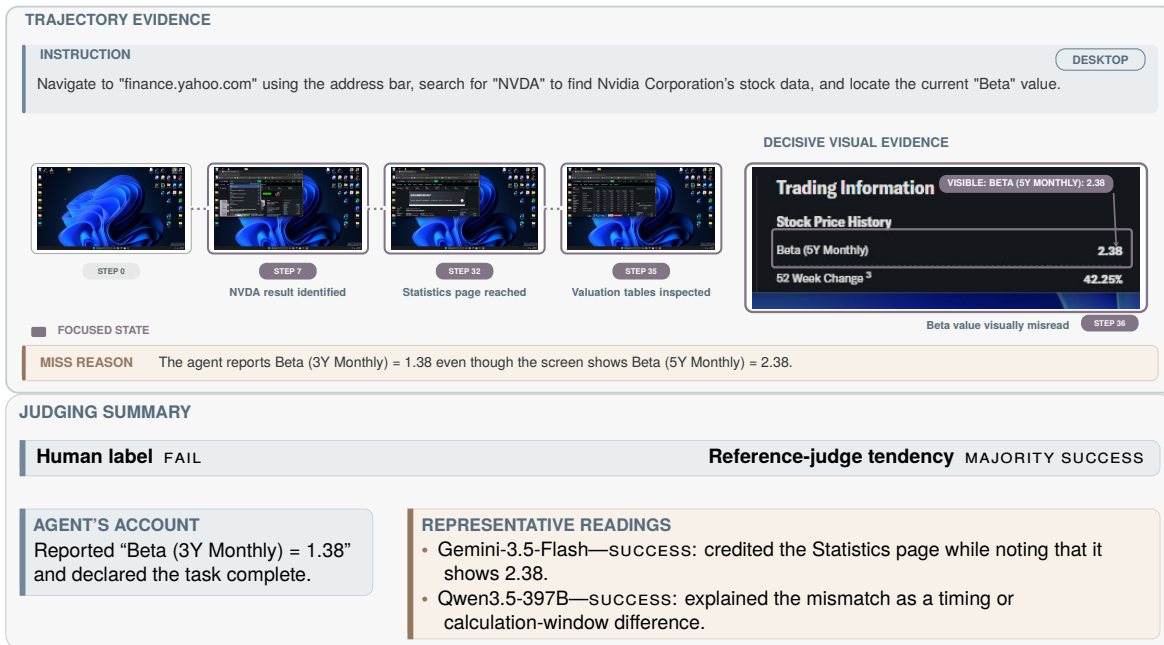


Figure 23: A fine-grained perception failure on Yahoo Finance: after reaching NVDA's Statistics page, the agent misreads both the Beta window and its displayed value.

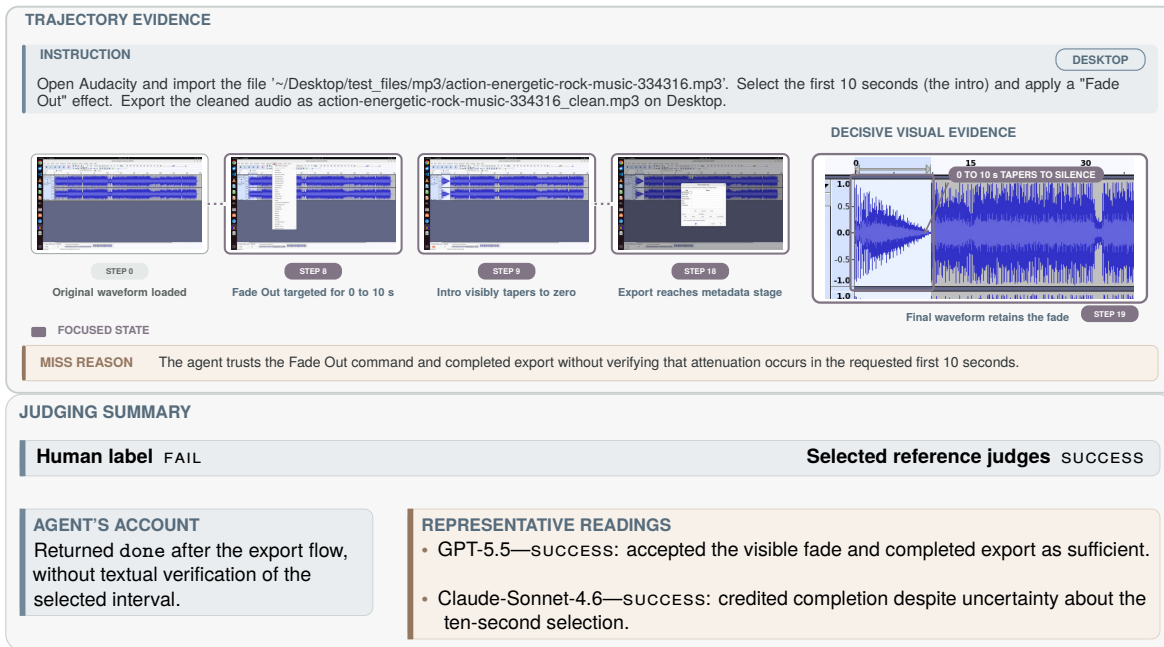


Figure 24: A fine-grained visual miss in Audacity: the agent follows convincing action and export cues without independently verifying whether the waveform satisfies the requested first-ten-second Fade Out.

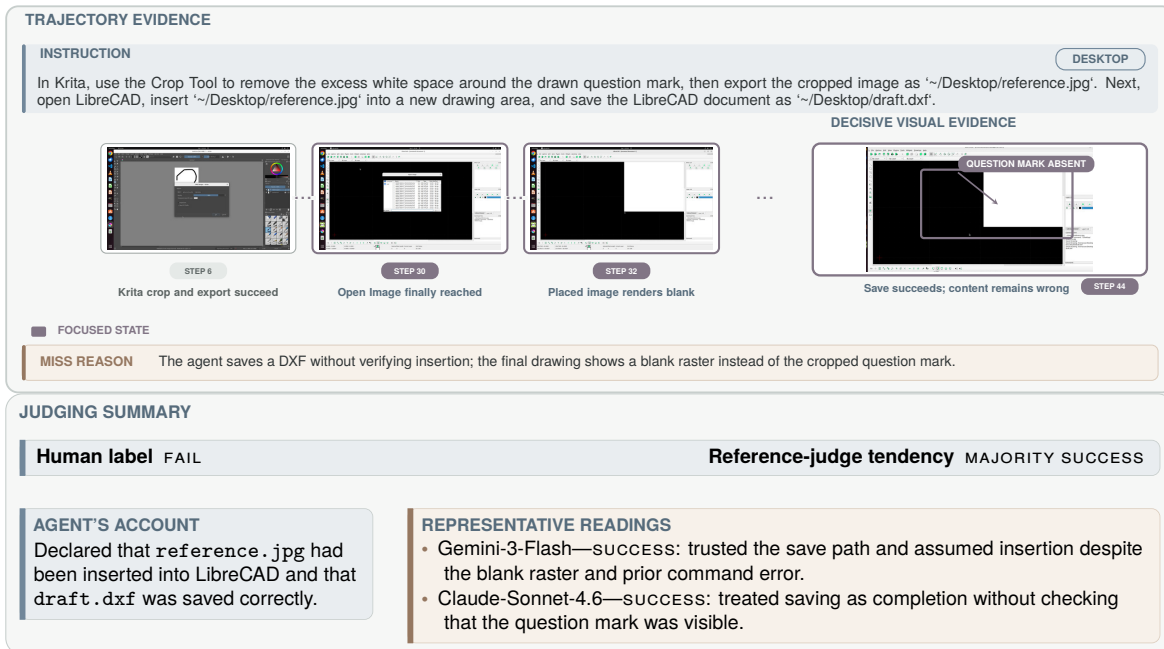


Figure 25: A long-horizon Ubuntu planning failure: repeated image-insertion errors are obscured by a successful final save, causing a visually incorrect LibreCAD document to be mistaken for a completed artifact.

An OSReward-Hard Case. A long, deceptive run from the challenge subset on which the judge field splits.

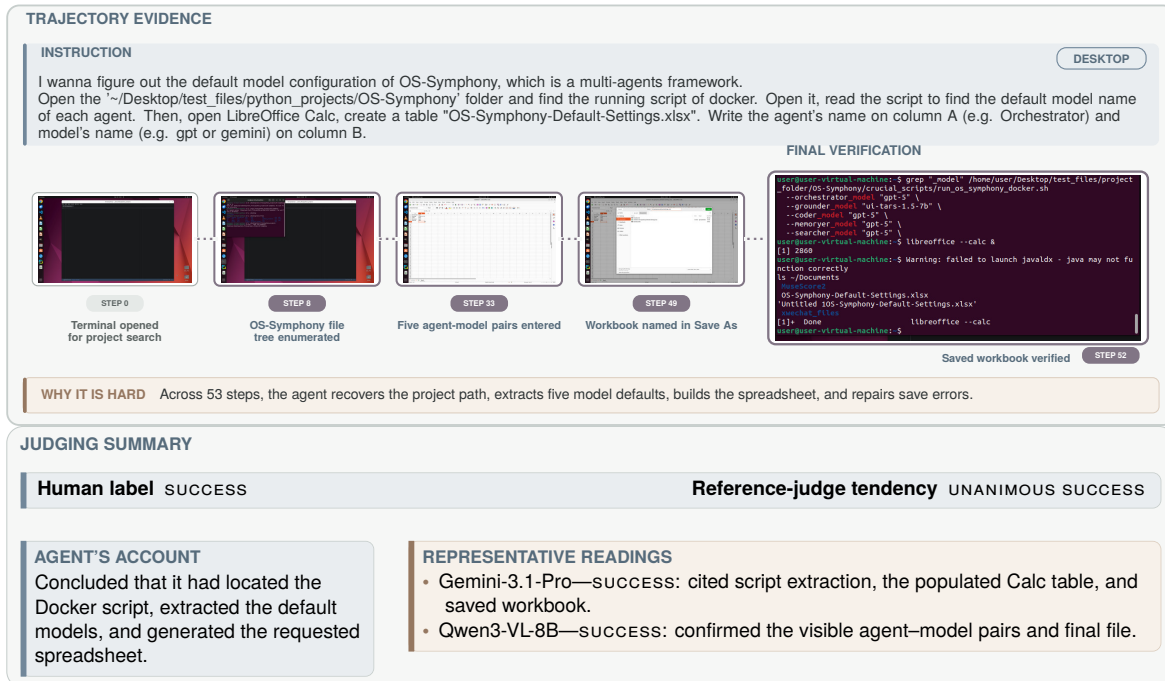


Figure 26: A successful long-horizon desktop case requiring sustained state tracking and verification of the final spreadsheet.

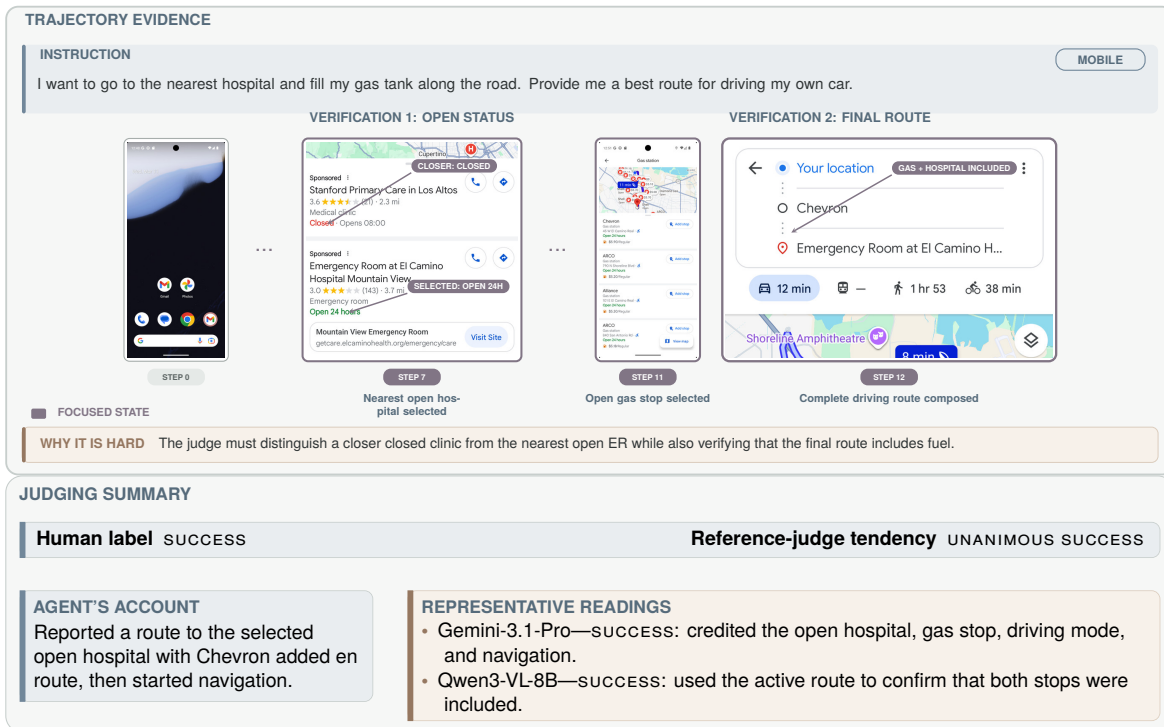


Figure 27: A mobile hard case requiring the judge to verify the nearest open hospital and a fuel stop in the final route.

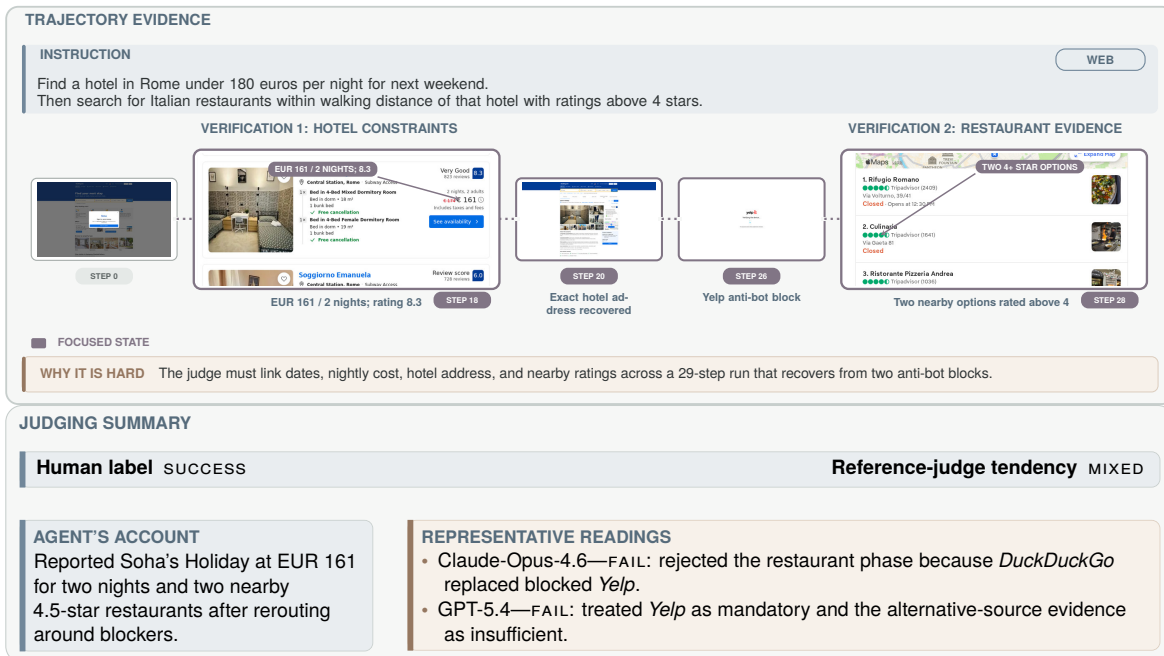


Figure 28: A successful Web hard case requiring cross-site constraint tracking and grounded recovery from anti-bot blocks.