

Article

Not peer-reviewed version

AI Agents for Web Testing: A Case Study in the Wild

[Naimeng Ye](#)^{*,†}, Xiao Yu[†], Ruize Xu[†], [Tianyi Peng](#), Zhou Yu

Posted Date: 26 August 2025

doi: 10.20944/preprints202508.1890.v1

Keywords: AI agent; automated web testing; vLM



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a Creative Commons CC BY 4.0 license, which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Article

AI Agents for Web Testing: A Case Study in the Wild

Naimeng Ye ^{†,*}, Xiao Yu [†], Ruize Xu [†], Tianyi Peng and Zhou Yu

Columbia University, USA

* Correspondence: ny2336@columbia.edu

[†] These authors contributed equally to this work.

Abstract

Automated web testing plays a critical role in ensuring high-quality user experiences and delivering business value. Traditional approaches primarily focus on code coverage and load testing, but often fall short of capturing complex user behaviors, leaving many usability issues undetected. The emergence of large language models (LLM) and AI agents opens new possibilities for web testing by enabling human-like interaction with websites and a general awareness of common usability problems. In this work, we present WebProber, a prototype AI agent-based web testing framework. Given a URL, WebProber autonomously explores the website, simulating real user interactions, identifying bugs and usability issues, and producing a human-readable report. We evaluate WebProber through a case study of 120 academic personal websites, where it uncovered 29 usability issues—many of which were missed by traditional tools. Our findings highlight agent-based testing as a promising direction while outlining directions for developing next-generation, user-centered testing frameworks.

Keywords: AI agent; automated web testing; vLM

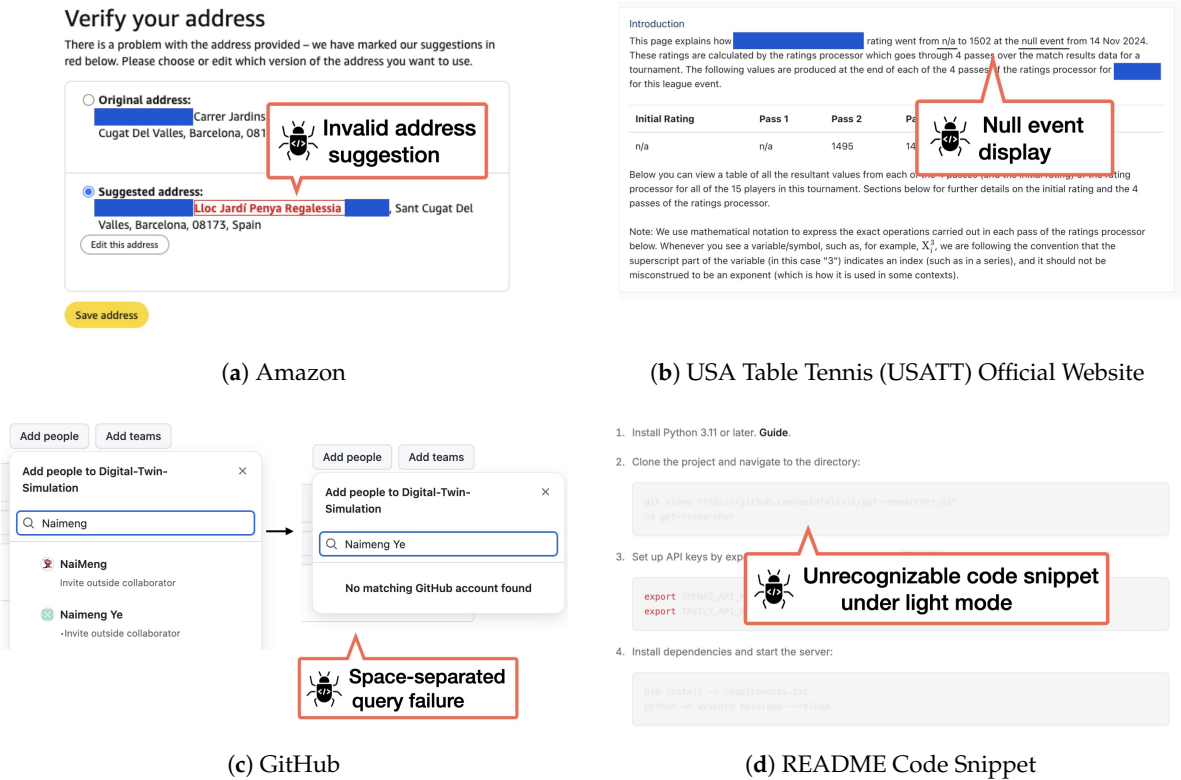
1. Introduction

The modern web hosts billions of websites [1], offering rich services and content that span nearly every aspect of daily life. Common web applications include e-commerce websites such as Amazon, social media platforms like Facebook, information portals like Wikipedia along with a vast number of personal websites.

To ensure the quality and reliability of these web applications, automated web testing has become a critical component of modern web development cycles. Traditional web testing approaches, such as static and dynamic analysis [2], have been crucial in mitigating common issues and vulnerabilities such as layout and functional bugs. These traditional approaches mainly rely on verifying code paths, automating scripted UI interactions, and measuring load performance using established tools like Cypress, Puppeteer, and JMeter [3–5]. Despite these efforts, such approaches face significant challenges in detecting real-world, user-facing issues. Since real users’ actions are highly diverse and context-dependent, software-based methods often fail to cover test-cases that capture the full spectrum of user behavior. This results in many undetected bugs and missing features that degrade user experience (see Figure 1 for real-world examples).

We introduce **WebProber**, a highly extensible web testing framework that leverages AI agents to simulate complex human behaviors on the web. Unlike existing approaches [6–8] that use large language models to generate test cases or interact with post-processed HTML files, WebProber employs powerful visual language models (VLMs) [9] to interact directly with visual webpages like human testers. Given a URL, WebProber explores the webpage for common user-side bugs by performing actions such as clicking, typing, and scrolling. It generates a comprehensive report of unexpected website behaviors based on its interaction history. We illustrate this workflow in Figure 2, which consists of three stages.(1) a proposal module that suggests error-prone features to investigate, guided by a bug database, (2) an interaction module that simulates user experience guided by VLMs, and (3)

a report generation module that examines the full interaction history to identify user-side bugs and suggest potential UI/UX improvements.



2. Related Work

Browser-Use Agents

With the advent of visual language models (VLMs), many works have explored the use of powerful models like GPT-4o and Claude-3.7 for web navigation tasks [10–12]. Earlier efforts used the accessibility tree or a screenshot of the webpage as input to the VLM, and prompted it to generate actions such as clicking and typing [12,13]. Recent works have explored various strategies to further improve an agent’s decision-making process, such as iteratively prompting the model to improve its own output [14,15], or augmenting the agent’s decision process using search algorithms such as breadth- or depth-first search [16], best-first search [17], and Monte Carlo tree search [18,19]. However, these works typically focus on solving pre-defined tasks such as finding a specific item on a shopping website, or navigating to a specific webpage. Our work aims to use agents to *discover* bugs missed by existing automated testing tools on real-world websites.

Automated Web Testing

Automated web testing emphasizes systematically testing web applications with minimal human intervention. Traditional approaches aim to generate action trajectories and can be broadly categorized into three classes: (1) randomized testing, where action sequences are generated stochastically [20]; (2) model-based methods, which construct a state graph of the application and use graph traversal algorithms such as depth-first search to explore it [21–23]; and (3) techniques based on reinforcement learning, which generate action sequences while maximizing a reward signal [24,25].

More recently, the field has begun to incorporate LLMs in automated web testing. They are used to expand the test action space [26–28], and to guide navigation and interaction [29–31]. In parallel, similar trends have emerged in mobile app testing [32–36]. Our work differs from prior literature by emphasizing the simulation of realistic user behaviors powered by VLMs, and by targeting contextual bugs often overlooked by traditional techniques, such as critical typographical errors or misdirected links.

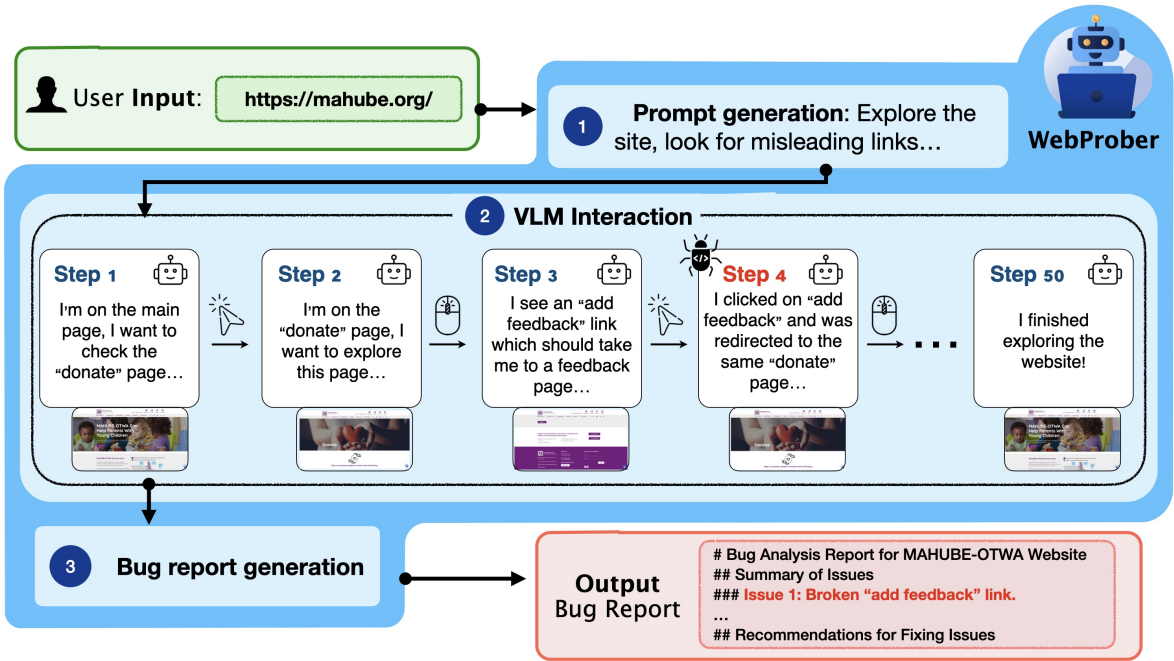


Figure 2. Workflow of WebProber. Given a user-provided URL, the agent generates a comprehensive bug report through three stages: (1) testing prompt generation, (2) VLM-guided interaction, and (3) bug report generation. For more details, refer to Section 3.

3. WebProber

We present WebProber, a web testing system based on AI agents. Given a website URL, WebProber returns a detailed report enumerating user-side bugs and UI/UX issues found during its interaction with the website. WebProber operates through a three-stage pipeline: (1) generating testing prompts that target common vulnerabilities for the particular class of website given, (2) simulating human-like web interactions, and (3) analyzing the interaction trajectory to generate comprehensive bug reports. We present an overview of this process in Figure 2, and describe this process in detail below.

Prompt Generation

Prompts guide an AI agent to look for common usability issues for different classes of web pages, enabling more targeted and efficient exploration by focusing on typical usability issue patterns. In this work, we created our testing prompts through an iterative refinement process and release the final high-quality prompt template in our repository. For each website type (e.g., personal websites), we begin with a preliminary prompt instructing the VLM on which features to test and what issues to detect. We then refine both the prompt and bug set through iterative cycles: applying WebProber to discover new bugs, manually verifying their reproducibility, and using a VLM to generate improved prompts based on the expanded bug set. This process continuously develops our prompt instructions while building a diverse collection of web usability bugs valuable for future evaluation. An example of prompt refinement is provided in Appendix B.1. While we demonstrate one effective approach to prompt generation, any method that produces high-quality, targeted prompts for usability testing would be compatible with our framework.

Interaction Simulation

Using the generated testing prompts from the previous stage, WebProber employs VLM-based agents to systematically interact with the website. Building on the Browser-Use Python package [37], our system iteratively (1) prompts a VLM for an action based on a website screenshot (e.g., clicking a button or entering text), (2) executes the action on the website, and (3) repeats until either the maximum step limit is reached or the target feature has been tested. Throughout this process, we preserve the complete interaction trajectory, including screenshots, reasoning traces, and actions.

Bug Report Generation

Finally, we generate detailed bug and usability reports by analyzing the full interaction trajectory with a VLM. Since usability issues typically emerge during interactive use, the complete interaction history is important for an accurate diagnosis. Detailed prompts for report generation are provided in Appendix B.2.

In our implementation, we used Claude-3.7 Sonnet ([38]) as the VLM for each stage of WebProber's pipeline. Each stage can be independently configured to use a different VLM, though exploring that is left as future work.

4. Experiments

To demonstrate the effectiveness of WebProber, we conducted a case study on real-world academic personal websites crawled from OpenReview author profiles. We collected 120 personal websites and applied WebProber on this dataset to detect usability issues. We then manually inspected the generated reports to analyze the detected bugs, specifically evaluating whether they represented genuine usability issues or false positives. The results are presented in Section 4.1.

In addition to measuring the capabilities of WebProber, we also investigated the coverage of bugs detectable by our framework. Since the total set of bugs on a website is unknown *a priori*, we manually inspected a representative subset of 80 websites to identify all potential bugs as a proxy for ground truth. We then ran WebProber on the same subset of websites to investigate both detected and undetected issues. The results and analysis are presented in the following sections.

4.1. Results

Our Approach Effectively Identifies Usability Issues That Impact User Experience

Across our dataset of 120 academic personal websites, WebProber successfully identified 29 usability issues (verified by the authors). In addition to bugs detectable by traditional techniques, e.g. rendering issues with images, our agent is also able to identify contextual bugs that are often overlooked by these methods. These issues span several categories, including link mistakes, rendering issues etc. We give a couple of representative examples of these usability issues as follows.

- **Broken or misdirected links** The most common class of bugs detected is broken or misdirected links. We present an example in Figure 3a: the agent identified that a project description was inconsistent with the paper linked through the "Read more here" button.
- **Logical inconsistencies** Finally, we find our WebProber is also able to detect logical inconsistencies in website contents, typically resulting from typographical errors. These errors sometimes lead to factual inaccuracies or user confusion. For example, in Figure 3b, the agent identified a spring course syllabus (determined by calendar dates) that incorrectly scheduled a "Fall break" week.

These results illustrate the capabilities of VLM-based web testing and provide insights into what types of issue can be automatically detected.

4.2. Discussion

While WebProber is able to identify real bugs and UI/UX issues in the wild, we also find numerous cases where human oversight is still needed for bug discovery. We present our findings below.

False Positives

While WebProber successfully discovered 29 usability issues, we found that 85% of all reported bugs across the 120 websites were false positives. The majority of these problems stemmed from technical limitations of the browser automation framework (the framework through which the agent applies actions on the webpage) rather than actual website issues. One common example is PDF access problems, which is often caused by the automation framework's security settings. However, the agent often incorrectly attributes these failures to website defects rather than automation constraints. Additionally, a small portion of false positives resulted from reasonable but incorrect assumptions of the agent, particularly when the agent lacked sufficient temporal or domain context to properly interpret website content. Figure 4 illustrates one such example.

Undetected Bugs

On our representative subset of 80 websites, we manually identified 32 bugs, of which WebProber successfully detected 19, achieving a coverage of 59.4%². The undetected bugs fell in two primary failure modes. First, and most frequently, bugs were often located deep within the website hierarchy, requiring navigation through multiple pages. Since we executed our pipeline only once per website, the agent often terminated exploration before encountering these deeply embedded issues. Second, certain pages containing bugs were inaccessible due to dynamic content rendering issues that our current implementation cannot handle effectively. These results suggest that effective bug detection requires improved exploration strategies capable of performing systematic, long-horizon traversals of website hierarchies and handling dynamic content. We defer these enhancements to future work.

² Since the authors may not have found all possible bugs, the actual coverage may be lower.

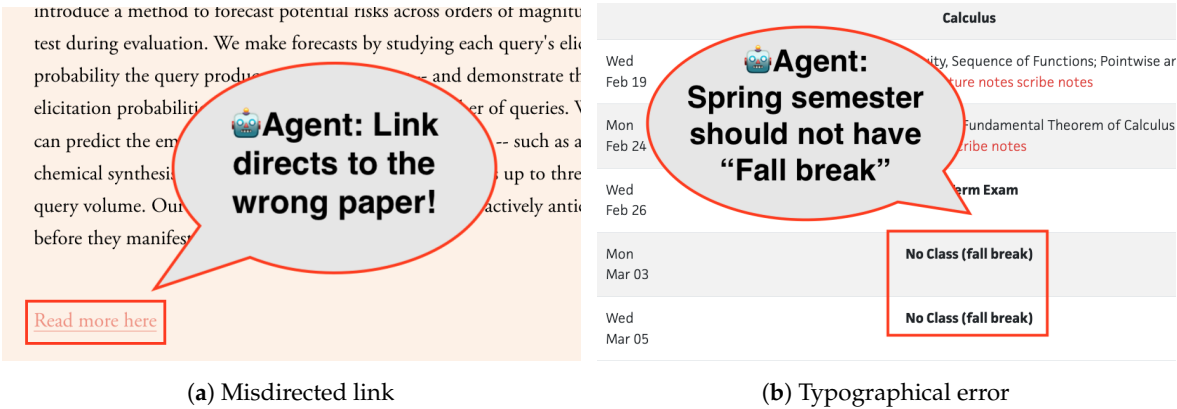


Figure 3. Examples of WebProber bug detection results. (a) A "Read more here" link for one research project incorrectly leads to a different paper. (b) A spring course syllabus mistakenly lists "fall break."

Recent News



Figure 4. Example of false positive: an announcement referencing a “2029” conference is flagged as a typo, but in reality, conferences such as ICCV do plan leadership roles and organizational details several years in advance.

5. Conclusion and Future Work

We introduced WebProber, an agent-based web testing framework. Applied to 120 academic personal websites, WebProber uncovered 29 usability issues—many missed by traditional tools—demonstrating the potential of agent-driven testing. This case study also revealed several challenges and future directions:

Agent-Browser Interaction. Agent interactions remain unreliable—misclicks, erratic navigation, and poor performance on complex sites contribute to false positives. Enhancing browser control fidelity is a key priority.

Bug Coverage and Training. Current agents are not optimized for bug discovery. Reinforcement learning and hybrid approaches incorporating traditional automated web testing tools may improve coverage and effectiveness.

Lack of Benchmarks. Progress is hindered by the absence of a standardized benchmark for web usability issues. Curating datasets like SWEBench would support training and evaluation.

Web Testing in Other Domains. Vibe-coded websites, startup landing pages, and non-profit websites often involve quick prototyping with limited budgets for thorough quality assurance. AI-generated sites, in particular, may contain bugs that their creators—often without professional development expertise—are unable to detect. While we believe AI agent-based web testing could significantly benefit these cases, we leave a rigorous field study for future work.

Appendix A. Example Bug Report Generated by WebProber

Figure A1 shows an example of bug report by our WebProber. Within 20 steps, the model correctly found the inconsistency between the dates of a course on the personal website and the course page. Via prompting, the report is well-formatted with bug descriptions and details such as common patterns or recurring issues, and recommendations for fixing the bugs.

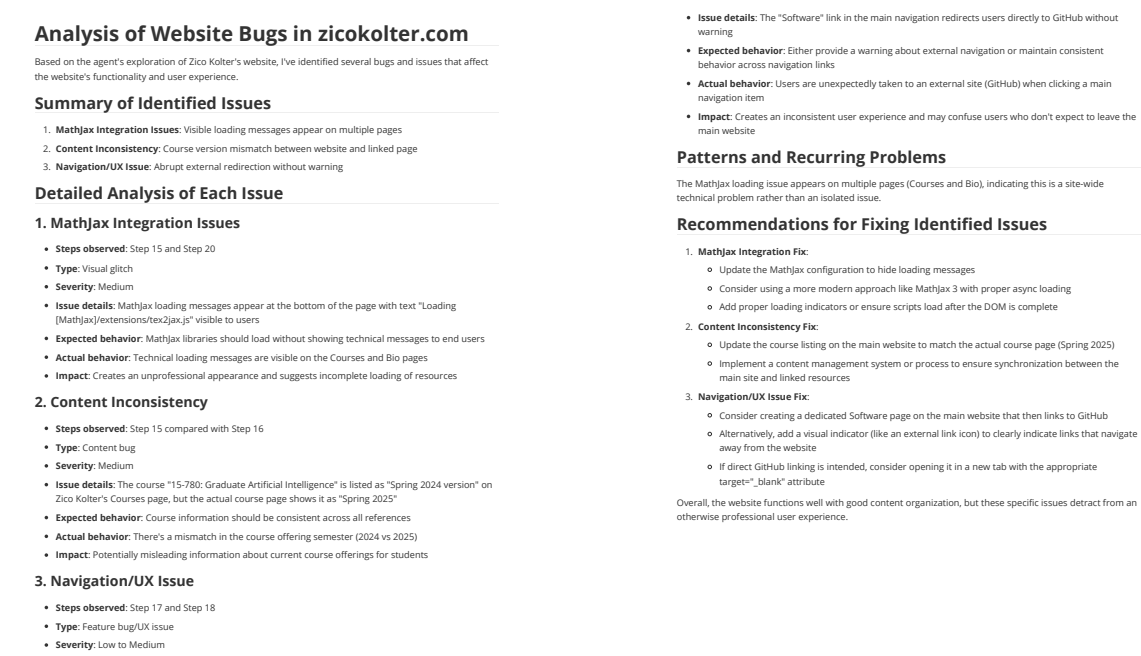


Figure A1. An Example Bug Report on a Personal Website.

Appendix B. Implementation Details

Appendix B.1. Proposing Potential Bugs

To expand the database of historical issues, we follow an iterative prompting pipeline. Specifically, given the bugs found with current prompts, we pick up highly representative and domain-specific ones. Then, these bug descriptions are fed into powerful models to iteratively improve the current prompts. The updated prompts contain potential bug proposals that encourage the model to explore areas that are often more likely to be buggy. Table A1 shows the initial prompt and the prompt iterated with ten historical bugs in the personal website domain.

Table A1. Initial and Iterated Prompts.

Iteration	Prompt
Initial	Go to the website [URL], a personal website. Explore the content, click on links, and occasionally pause to assess whether what is shown and linked on the website is coherent and appropriate. Unreasonable or problematic issues include, but are not limited to: broken or mismatched images/links, UI glitches/incapabilities, illogical or unfunctional web design, or textual errors, etc.
One Iteration	Go to the website [URL], a personal website. Systematically explore all accessible content, including menus, links, embedded media, interactive elements, and downloadable materials for WEBSITE BUGS. You should prioritize areas that are often more likely to contain issues. At each stage, critically evaluate whether the displayed information, layout, and behavior align with expectations for a functional and professional web experience. Carefully inspect for issues such as, but not limited to: (1) Broken elements: dead/missing links, 404 pages, failed image or video loads. (2) Interaction failures: non-responsive buttons, malfunctioning forms or filters, non-working download or redirect actions. (3) UI/UX flaws: lack of visual feedback, missing tooltips/ESC buttons, layout inconsistencies, uncustomized templates, poor mobile compatibility. (4) Content inconsistencies: outdated or contradictory data (e.g., dates or names), mismatched references or external links, typos or formatting errors. (5) Domain-specific bugs: for instance, broken external links to publications, projects, GitHub, Google Scholar, etc. Incorrect anchor links (e.g., internal navigation like #about or #projects not working). Outdated or dead email links (e.g., mailto: pointing to deprecated addresses). Missing or malformed citation info (e.g., BibTeX files, DOI links not rendering or downloading properly). Mismatched thumbnails or missing alt-text on research project previews. Videos or talks not embedded properly (e.g., iframe blocked by CORS). For each identified issue, consider its impact, repeatability, and specific trigger (e.g., "clicking X under condition Y leads to error Z")

Appendix B.2. Generating Bug Reports

Given the trajectories, we then prompt powerful models such as Claude-3.7 Sonnet to analyze the trajectories and generate a formatted bug report. We prompt the model to contain a summary of bugs, the steps where bugs occurred, short descriptions, and then more details such as common patterns or recurring issues, and recommendations for fixing the bugs. The full prompt is shown in Table A2.

Table A2. Prompt for Generating Bug Reports.

Prompt
Please analyze the following agent run trajectory and identify any potential bugs or glitches in the website being tested. Consider both feature bugs (missing or incorrect functionality) and glitch-like bugs (visual or behavioral anomalies). Note that the type of bug is not always obvious, so don't be afraid to make an assumption. For example, if the website does not support certain features that the agent is trying to use, that is a bug (e.g. the agent is trying to use the "add to cart" feature, but the website does not have a cart, or that the agent is searching in some language that the website does not support). For each step, I'll provide: 0. The screenshot of the current browser state 1. The agent's evaluation of the step 2. The next goal 3. The action taken Please analyze the entire sequence of steps and identify: 1. Any unexpected behaviors or errors of the website itself (*note: not the agent's actions*) 2. Missing or incorrect functionality 3. Visual glitches or UI inconsistencies 4. Any other anomalies that might indicate bugs Here's the step-by-step trajectory: [Trajectory] Based on the above trajectory, please provide: 1. A summary of any bugs or glitches identified 2. The specific steps where issues occurred 3. The nature of each issue (feature bug, visual glitch, etc.) 4. Any patterns or recurring problems 5. Recommendations for fixing the identified issues For each identified issue, please specify: - The step number where it occurred - Whether it's a feature bug or visual glitch - The severity of the issue - The expected behavior vs actual behavior

References

1. Chakarov, R. How Many Websites Are There? How many are active in 2023? <https://webtribunal.net/blog/how-many-websites>, 2023.
2. Ricca, F.; Tonella, P. Analysis and testing of Web applications. In Proceedings of the Proceedings of the 23rd International Conference on Software Engineering. ICSE 2001, 2001, pp. 25–34. <https://doi.org/10.1109/ICSE.2001.919078>.
3. Cypress.io. Cypress: Testing Frameworks for JavaScript. <https://www.cypress.io/>, 2025. Accessed: 2025-06-26.
4. Google Chrome Developers. Puppeteer: Headless Chrome Node.js API. <https://pptr.dev/>, 2025. Accessed: 2025-06-26.
5. Apache Software Foundation. Apache JMeter: Load Testing for Web Applications. <https://jmeter.apache.org/>, 2025. Accessed: 2025-06-26.
6. Le, N.K.; Bui, Q.M.; Nguyen, M.N.; Nguyen, H.; Vo, T.; Luu, S.T.; Nomura, S.; Nguyen, M.L. Automated Web Application Testing: End-to-End Test Case Generation with Large Language Models and Screen Transition Graphs, 2025, [arXiv:cs.SE/2506.02529].

7. Wang, D.; Hsu, T.Y.; Lu, Y.; Gu, H.; Cui, L.; Xie, Y.; Headean, W.; Yao, B.; Veeragouni, A.; Liu, J.; et al. AgentA/B: Automated and Scalable Web A/BTesting with Interactive LLM Agents, 2025, [[arXiv:cs.HC/2504.09723](https://arxiv.org/abs/cs/2504.09723)].
8. Lu, Y.; Yao, B.; Gu, H.; Huang, J.; Wang, Z.J.; Li, Y.; Gesi, J.; He, Q.; Li, T.J.J.; Wang, D. Uxagent: An llm agent-based usability testing framework for web design. In Proceedings of the Proceedings of the Extended Abstracts of the CHI Conference on Human Factors in Computing Systems, 2025, pp. 1–12.
9. Bordes, F.; Pang, R.Y.; Ajay, A.; Li, A.C.; Bardes, A.; Petryk, S.; Mañas, O.; Lin, Z.; Mahmoud, A.; Jayaraman, B.; et al. An introduction to vision-language modeling. *arXiv preprint arXiv:2405.17247* 2024.
10. Liu, X.; Yu, H.; Zhang, H.; Xu, Y.; Lei, X.; Lai, H.; Gu, Y.; Ding, H.; Men, K.; Yang, K.; et al. AgentBench: Evaluating LLMs as Agents, 2023, [[arXiv:cs.AI/2308.03688](https://arxiv.org/abs/cs/2308.03688)].
11. Zhou, S.; Xu, F.F.; Zhu, H.; Zhou, X.; Lo, R.; Sridhar, A.; Cheng, X.; Ou, T.; Bisk, Y.; Fried, D.; et al. WebArena: A Realistic Web Environment for Building Autonomous Agents, 2024, [[arXiv:cs.AI/2307.13854](https://arxiv.org/abs/cs/2307.13854)].
12. Koh, J.Y.; Lo, R.; Jang, L.; Duvvur, V.; Lim, M.C.; Huang, P.Y.; Neubig, G.; Zhou, S.; Salakhutdinov, R.; Fried, D. VisualWebArena: Evaluating Multimodal Agents on Realistic Visual Web Tasks, 2024, [[arXiv:cs.LG/2401.13649](https://arxiv.org/abs/cs/2401.13649)].
13. Yang, J.; Zhang, H.; Li, F.; Zou, X.; Li, C.; Gao, J. Set-of-Mark Prompting Unleashes Extraordinary Visual Grounding in GPT-4V, 2023, [[arXiv:cs.CV/2310.11441](https://arxiv.org/abs/cs/2310.11441)].
14. Madaan, A.; Tandon, N.; Gupta, P.; Hallinan, S.; Gao, L.; Wiegrefe, S.; Alon, U.; Dziri, N.; Prabhunoye, S.; Yang, Y.; et al. Self-Refine: Iterative Refinement with Self-Feedback, 2023, [[arXiv:cs.CL/2303.17651](https://arxiv.org/abs/cs/2303.17651)].
15. Shinn, N.; Cassano, F.; Berman, E.; Gopinath, A.; Narasimhan, K.; Yao, S. Reflexion: Language Agents with Verbal Reinforcement Learning, 2023, [[arXiv:cs.AI/2303.11366](https://arxiv.org/abs/cs/2303.11366)].
16. Yao, S.; Yu, D.; Zhao, J.; Shafran, I.; Griffiths, T.L.; Cao, Y.; Narasimhan, K. Tree of Thoughts: Deliberate Problem Solving with Large Language Models, 2023, [[arXiv:cs.CL/2305.10601](https://arxiv.org/abs/cs/2305.10601)].
17. Koh, J.Y.; McAleer, S.; Fried, D.; Salakhutdinov, R. Tree Search for Language Model Agents, 2024, [[arXiv:cs.AI/2407.01476](https://arxiv.org/abs/cs/2407.01476)].
18. Yu, X.; Chen, M.; Yu, Z. Prompt-Based Monte-Carlo Tree Search for Goal-Oriented Dialogue Policy Planning, 2023, [[arXiv:cs.CL/2305.13660](https://arxiv.org/abs/cs/2305.13660)].
19. Yu, X.; Peng, B.; Vajipey, V.; Cheng, H.; Galley, M.; Gao, J.; Yu, Z. ExACT: Teaching AI Agents to Explore with Reflective-MCTS and Exploratory Learning, 2025, [[arXiv:cs.CL/2410.02052](https://arxiv.org/abs/cs/2410.02052)].
20. Android Developers. Monkey. <https://developer.android.com>, 2022. Accessed: 2025-06-25.
21. Mesbah, A.; Van Deursen, A.; Lensenlink, S. Crawling Ajax-based web applications through dynamic analysis of user interface state changes. *ACM Transactions on the Web (TWEB)* 2012, 6, 1–30.
22. Stocco, A.; Willi, A.; Starace, L.L.L.; Biagiola, M.; Tonella, P. Neural embeddings for web testing. *arXiv preprint arXiv:2306.07400* 2023.
23. Liu, C.; Wang, J.; Yang, W.; Zhang, Y.; Xie, T. Judge: Effective State Abstraction for Guiding Automated Web GUI Testing. *ACM Transactions on Software Engineering and Methodology* 2025.
24. Zheng, Y.; Liu, Y.; Xie, X.; Liu, Y.; Ma, L.; Hao, J.; Liu, Y. Automatic web testing using curiosity-driven reinforcement learning. In Proceedings of the 2021 IEEE/ACM 43rd International Conference on Software Engineering (ICSE). IEEE, 2021, pp. 423–435.
25. Sherin, S.; Muqet, A.; Khan, M.U.; Iqbal, M.Z. QExplore: An exploration strategy for dynamic web applications using guided search. *Journal of Systems and Software* 2023, 195, 111512.
26. Liu, Z.; Chen, C.; Wang, J.; Che, X.; Huang, Y.; Hu, J.; Wang, Q. Fill in the blank: Context-aware automated text input generation for mobile gui testing. In Proceedings of the 2023 IEEE/ACM 45th International Conference on Software Engineering (ICSE). IEEE, 2023, pp. 1355–1367.
27. Liu, Z.; Chen, C.; Wang, J.; Chen, M.; Wu, B.; Tian, Z.; Huang, Y.; Hu, J.; Wang, Q. Testing the limits: Unusual text inputs generation for mobile app crash detection with large language model. In Proceedings of the Proceedings of the IEEE/ACM 46th international conference on software engineering, 2024, pp. 1–12.
28. Wang, S.; Wang, S.; Fan, Y.; Li, X.; Liu, Y. Leveraging Large Vision-Language Model for Better Automatic Web GUI Testing. In Proceedings of the 2024 IEEE International Conference on Software Maintenance and Evolution (ICSME). IEEE, 2024, pp. 125–137.
29. Alian, P.; Nashid, N.; Shahbandeh, M.; Shabani, T.; Mesbah, A. Feature-Driven End-To-End Test Generation. In Proceedings of the 2025 IEEE/ACM 47th International Conference on Software Engineering (ICSE). IEEE Computer Society, 2025, pp. 678–678.
30. Shahbandeh, M.; Alian, P.; Nashid, N.; Mesbah, A. Navigate: Functionality-guided web application navigation. *arXiv preprint arXiv:2409.10741* 2024.

31. Liu, C.; Gu, Z.; Wu, G.; Zhang, Y.; Wei, J.; Xie, T. Temac: Multi-Agent Collaboration for Automated Web GUI Testing. *arXiv preprint arXiv:2506.00520* **2025**.
32. Liu, Z.; Chen, C.; Wang, J.; Chen, M.; Wu, B.; Che, X.; Wang, D.; Wang, Q. Make llm a testing expert: Bringing human-like interaction to mobile gui testing via functionality-aware decisions. In Proceedings of the Proceedings of the IEEE/ACM 46th International Conference on Software Engineering, 2024, pp. 1–13.
33. Yoon, J.; Feldt, R.; Yoo, S. Autonomous large language model agents enabling intent-driven mobile gui testing. *arXiv preprint arXiv:2311.08649* **2023**.
34. Lee, S.; Choi, J.; Lee, J.; Wasi, M.H.; Choi, H.; Ko, S.; Oh, S.; Shin, I. Mobilegpt: Augmenting llm with human-like app memory for mobile task automation. In Proceedings of the Proceedings of the 30th Annual International Conference on Mobile Computing and Networking, 2024, pp. 1119–1133.
35. Wen, H.; Tian, S.; Pavlov, B.; Du, W.; Li, Y.; Chang, G.; Zhao, S.; Liu, J.; Liu, Y.; Zhang, Y.Q.; et al. AutoDroid-V2: Boosting SLM-based GUI Agents via Code Generation. *arXiv preprint arXiv:2412.18116* **2024**.
36. Chen, M.; Liu, Z.; Chen, C.; Wang, J.; Wu, B.; Hu, J.; Wang, Q. Standing on the Shoulders of Giants: Bug-Aware Automated GUI Testing via Retrieval Augmentation. *Proceedings of the ACM on Software Engineering* **2025**, 2, 825–846.
37. Müller, M.; Žunič, G. Browser Use: Enable AI to control your browser, 2024.
38. Anthropic, A. Claude 3.7 and Claude Code, 2025.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.