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

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 Automated web testing plays a critical role in ensuring high-quality user experi-
2 ences and delivering business value. Traditional approaches primarily focus on
3 code coverage and load testing, but often fall short of capturing complex user behav-
4 iors, leaving many usability issues undetected. The emergence of large language
5 models (LLM) and AI agents opens new possibilities for web testing by enabling
6 human-like interaction with websites and a general awareness of common usability
7 problems. In this work, we present WebProber, a prototype AI agent-based web
8 testing framework. Given a URL, WebProber autonomously explores the website,
9 simulating real user interactions, identifying bugs and usability issues, and pro-
10 ducing a human-readable report. We evaluate WebProber through a case study of
11 120 academic personal websites, where it uncovered 29 usability issues—many of
12 which were missed by traditional tools. Our findings highlight agent-based testing
13 as a promising direction while outlining directions for developing next-generation,
14 user-centered testing frameworks.

15 The modern web hosts billions of websites (Chakarov, 2023), offering rich services and content that
16 span nearly every aspect of daily life. Common web applications include e-commerce websites such
17 as Amazon, social media platforms like Facebook, information portals like Wikipedia along with a
18 vast number of personal websites.

19 To ensure the quality and reliability of these web applications, automated web testing has become a
20 critical component of modern web development cycles. Traditional web testing approaches, such as
21 static and dynamic analysis (Ricca & Tonella, 2001), have been crucial in mitigating common issues
22 and vulnerabilities such as layout and functional bugs. These traditional approaches mainly rely on
23 verifying code paths, automating scripted UI interactions, and measuring load performance using
24 established tools like Cypress, Puppeteer, and JMeter (Cypress.io, 2025; Google Chrome Developers,
25 2025; Apache Software Foundation, 2025). Despite these efforts, such approaches face significant
26 challenges in detecting real-world, user-facing issues. Since real users’ actions are highly diverse
27 and context-dependent, software-based methods often fail to cover test-cases that capture the full
28 spectrum of user behavior. This results in many undetected bugs and missing features that degrade
29 user experience (see Figure 1 for real-world examples).

30 We introduce **WebProber**, a highly extensible web testing framework that leverages AI agents to
31 simulate complex human behaviors on the web. Unlike existing approaches (Le et al., 2025; Wang
32 et al., 2025; Lu et al., 2025) that use large language models to generate test cases or interact with
33 post-processed HTML files, WebProber employs powerful visual language models (VLMs) (Bordes
34 et al., 2024) to interact directly with visual webpages like human testers. Given a URL, WebProber
35 explores the webpage for common user-side bugs by performing actions such as clicking, typing,
36 and scrolling. It generates a comprehensive report of unexpected website behaviors based on its
37 interaction history. We illustrate this workflow in Figure 2, which consists of three stages. (1) a
38 proposal module that suggests error-prone features to investigate, guided by a bug database, (2)
39 an interaction module that simulates user experience guided by VLMs, and (3) a report generation

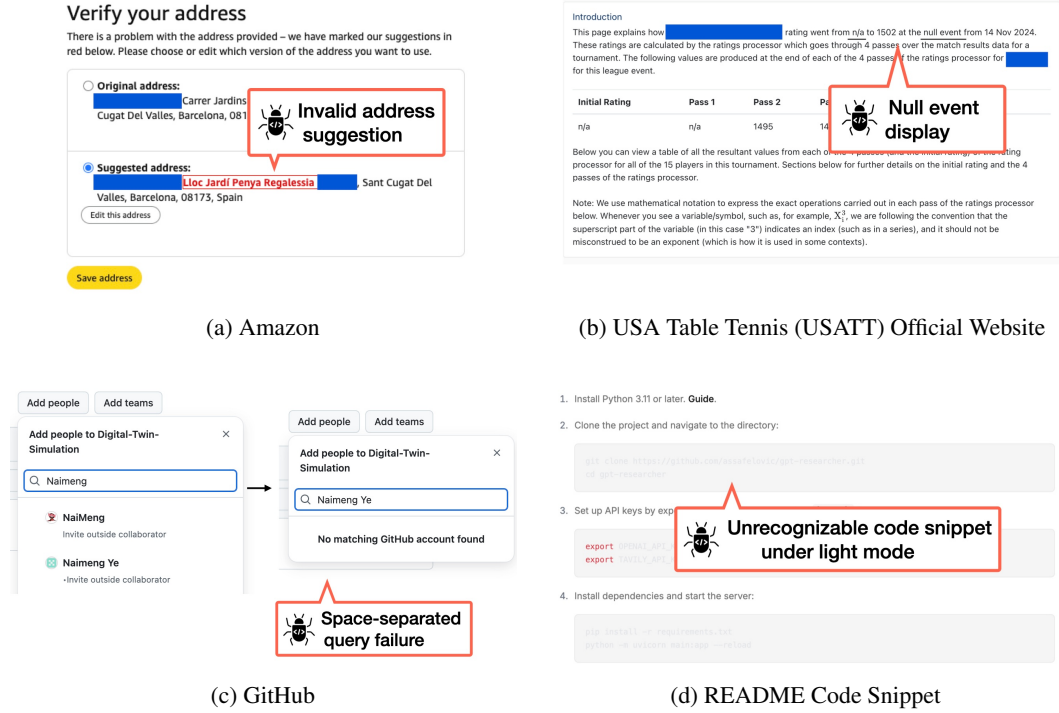


Figure 1: Website usability bugs that are not easily detected by traditional web testing techniques. (a) On the Amazon Spain website, during a purchase, the system suggests a non-existent and unclickable address. (b) The USATT website displays null event text for a league event. (c) In a GitHub organization repository, the user search function does not support queries with spaces when adding users. (d) On certain MCP server pages, code snippets in the README file are illegible in light mode due to poor color contrast.

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

42 As a case study, we deployed our framework on 120 personal websites in the wild and found that
43 our framework is able to identify 29 usability issues that impact user experience. Many of these
44 issues—such as textual errors and misdirected links—were not detected by traditional automated
45 testing tools, highlighting the unique strengths of agent-based testing in uncovering subtle, human-
46 centric problems. Our empirical study on personal websites presents a first step towards building
47 a scalable web testing framework based on AI agents, and we hope that this work can serve as a
48 foundation for future research in this direction.

49 In summary, our contributions are:

- 50 1. We introduce WebProber, a highly extensible web testing framework that leverages AI
51 agents to simulate human behavior on the web.
- 52 2. We present a case study on 120 personal websites in the wild, on which WebProber found
53 29 usability issues.
- 54 3. We release our code and our human-annotated bug database for future research.

55 1 Related Work

56 **Browser-Use Agents** With the advent of visual language models (VLMs), many works have
57 explored the use of powerful models like GPT-4o and Claude-3.7 for web navigation tasks (Liu
58 et al., 2023a; Zhou et al., 2024; Koh et al., 2024a). Earlier efforts used the accessibility tree or a

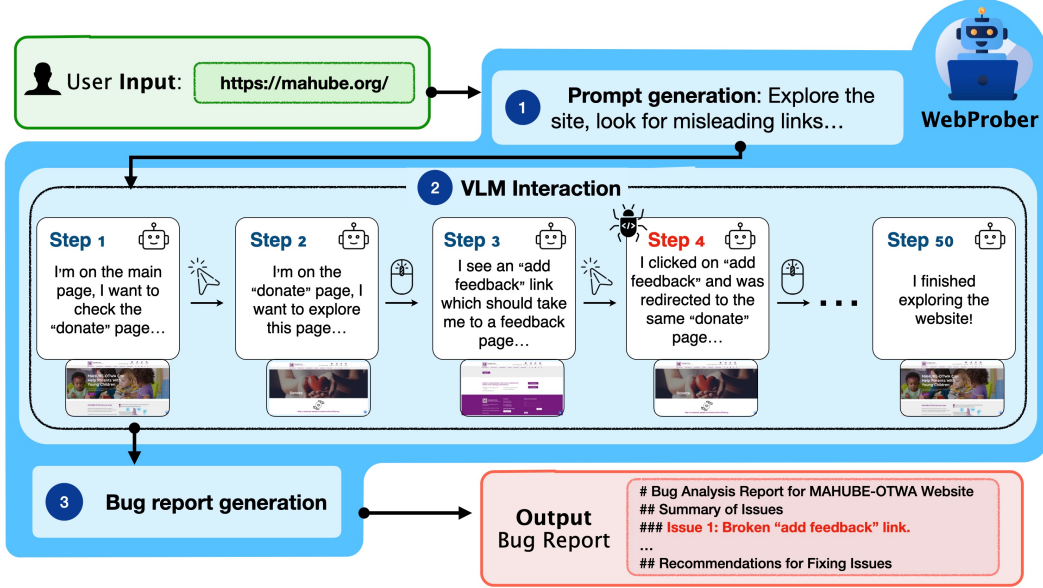


Figure 2: Workflow of WebProber. Given a user-provided URL, the agent generates a comprehensive bug report through three stages: (1) testing prompt generation, (2) VLM-guided interaction, and (3) bug report generation. For more details, refer to section 2.

screenshot of the webpage as input to the VLM, and prompted it to generate actions such as clicking and typing (Yang et al., 2023; Koh et al., 2024a). 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 (Madaan et al., 2023; Shinn et al., 2023), or augmenting the agent’s decision process using search algorithms such as breadth- or depth-first search (Yao et al., 2023), best-first search (Koh et al., 2024b), and Monte Carlo tree search (Yu et al., 2023, 2025). 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 (Android Developers (2022)); (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 (Mesbah et al. (2012); Stocco et al. (2023); Liu et al. (2025b)); and (3) techniques based on reinforcement learning, which generate action sequences while maximizing a reward signal (Zheng et al. (2021); Sherin et al. (2023)).

More recently, the field has begun to incorporate LLMs in automated web testing. They are used to expand the test action space (Liu et al. (2023b, 2024b); Wang et al. (2024), and to guide navigation and interaction (Alian et al. (2025); Shahbandeh et al. (2024); Liu et al. (2025a). In parallel, similar trends have emerged in mobile app testing (Liu et al. (2024a); Yoon et al. (2023); Lee et al. (2024); Wen et al. (2024); Chen et al. (2025)). 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.

2 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 (Müller & Žunič, 2024), 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 (Anthropic (2025)) 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.

3 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 3.1.

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.

3.1 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.

137 • **Logical inconsistencies** Finally, we find our WebProber is also able to detect logical
 138 inconsistencies in website contents, typically resulting from typographical errors. These
 139 errors sometimes lead to factual inaccuracies or user confusion. For example, in Figure 3b,
 140 the agent identified a spring course syllabus (determined by calendar dates) that incorrectly
 141 scheduled a "Fall break" week.

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

144 3.2 Discussion

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

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

156 **Undetected bugs** On our representative subset of 80 websites, we manually identified 32 bugs,
 157 of which WebProber successfully detected 19, achieving a coverage of 59.4%¹. The undetected
 158 bugs fell in two primary failure modes. First, and most frequently, bugs were often located deep
 159 within the website hierarchy, requiring navigation through multiple pages. Since we executed our
 160 pipeline only once per website, the agent often terminated exploration before encountering these
 161 deeply embedded issues. Second, certain pages containing bugs were inaccessible due to dynamic
 162 content rendering issues that our current implementation cannot handle effectively. These results
 163 suggest that effective bug detection requires improved exploration strategies capable of performing
 164 systematic, long-horizon traversals of website hierarchies and handling dynamic content. We defer
 165 these enhancements to future work.



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."

166 4 Conclusion and Future Work

167 We introduced WebProber, an agent-based web testing framework. Applied to 120 aca-
 168 demic personal websites, WebProber uncovered 29 usability issues—many missed by traditional

¹Since the authors may not have found all possible bugs, the actual coverage may be lower.

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.

169 tools—demonstrating the potential of agent-driven testing. This case study also revealed several
170 challenges and future directions:

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

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

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

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

184 References

185 Parsa Alian, Noor Nashid, Mobina Shahbandeh, Taha Shabani, and Ali Mesbah. Feature-driven end-
186 to-end test generation. In *2025 IEEE/ACM 47th International Conference on Software Engineering*
187 *(ICSE)*, pp. 678–678. IEEE Computer Society, 2025.

188 Android Developers. Monkey. <https://developer.android.com>, 2022. Accessed: 2025-
189 06-25.

190 AI Anthropic. Claude 3.7 and claude code, 2025. URL [https://www.anthropic.com/
191 news/claude-3-7-sonnet](https://www.anthropic.com/news/claude-3-7-sonnet).

192 Apache Software Foundation. Apache jmeter: Load testing for web applications. [https://
193 jmeter.apache.org/](https://jmeter.apache.org/), 2025. Accessed: 2025-06-26.

194 Florian Bordes, Richard Yuanzhe Pang, Anurag Ajay, Alexander C Li, Adrien Bardes, Suzanne
195 Petryk, Oscar Mañas, Zhiqiu Lin, Anas Mahmoud, Bargav Jayaraman, et al. An introduction to
196 vision-language modeling. *arXiv preprint arXiv:2405.17247*, 2024.

197 Radoslave Chakarov. How many websites are there? how many are active in 2023? [https:
198 //webtribunal.net/blog/how-many-websites](https://webtribunal.net/blog/how-many-websites), 2023.

199 Mengzhuo Chen, Zhe Liu, Chunyang Chen, Junjie Wang, Boyu Wu, Jun Hu, and Qing Wang.
200 Standing on the shoulders of giants: Bug-aware automated gui testing via retrieval augmentation.
201 *Proceedings of the ACM on Software Engineering*, 2(FSE):825–846, 2025.

202 Cypress.io. Cypress: Testing frameworks for javascript. <https://www.cypress.io/>, 2025.
203 Accessed: 2025-06-26.

204 Google Chrome Developers. Puppeteer: Headless chrome node.js api. <https://pptr.dev/>,
205 2025. Accessed: 2025-06-26.

206 Jing Yu Koh, Robert Lo, Lawrence Jang, Vikram Duvvur, Ming Chong Lim, Po-Yu Huang, Graham
207 Neubig, Shuyan Zhou, Ruslan Salakhutdinov, and Daniel Fried. Visualwebarena: Evaluating
208 multimodal agents on realistic visual web tasks, 2024a. URL [https://arxiv.org/abs/](https://arxiv.org/abs/2401.13649)
209 [2401.13649](https://arxiv.org/abs/2401.13649).

210 Jing Yu Koh, Stephen McAleer, Daniel Fried, and Ruslan Salakhutdinov. Tree search for language
211 model agents, 2024b. URL <https://arxiv.org/abs/2407.01476>.

212 Nguyen-Khang Le, Quan Minh Bui, Minh Ngoc Nguyen, Hiep Nguyen, Trung Vo, Son T. Luu,
213 Shoshin Nomura, and Minh Le Nguyen. Automated web application testing: End-to-end test
214 case generation with large language models and screen transition graphs, 2025. URL [https:](https://arxiv.org/abs/2506.02529)
215 [//arxiv.org/abs/2506.02529](https://arxiv.org/abs/2506.02529).

216 Sunjae Lee, Junyoung Choi, Jungjae Lee, Munim Hasan Wasi, Hojun Choi, Steve Ko, Sangeun
217 Oh, and Insik Shin. Mobilegpt: Augmenting llm with human-like app memory for mobile task
218 automation. In *Proceedings of the 30th Annual International Conference on Mobile Computing*
219 *and Networking*, pp. 1119–1133, 2024.

220 Chenxu Liu, Zhiyu Gu, Guoquan Wu, Ying Zhang, Jun Wei, and Tao Xie. Temac: Multi-agent
221 collaboration for automated web gui testing. *arXiv preprint arXiv:2506.00520*, 2025a.

222 Chenxu Liu, Junheng Wang, Wei Yang, Ying Zhang, and Tao Xie. Judge: Effective state abstrac-
223 tion for guiding automated web gui testing. *ACM Transactions on Software Engineering and*
224 *Methodology*, 2025b.

225 Xiao Liu, Hao Yu, Hanchen Zhang, Yifan Xu, Xuanyu Lei, Hanyu Lai, Yu Gu, Hangliang Ding,
226 Kaiwen Men, Kejuan Yang, Shudan Zhang, Xiang Deng, Aohan Zeng, Zhengxiao Du, Chenhui
227 Zhang, Sheng Shen, Tianjun Zhang, Yu Su, Huan Sun, Minlie Huang, Yuxiao Dong, and Jie
228 Tang. Agentbench: Evaluating llms as agents, 2023a. URL [https://arxiv.org/abs/](https://arxiv.org/abs/2308.03688)
229 [2308.03688](https://arxiv.org/abs/2308.03688).

230 Zhe Liu, Chunyang Chen, Junjie Wang, Xing Che, Yuekai Huang, Jun Hu, and Qing Wang. Fill in the
231 blank: Context-aware automated text input generation for mobile gui testing. In *2023 IEEE/ACM*
232 *45th International Conference on Software Engineering (ICSE)*, pp. 1355–1367. IEEE, 2023b.

233 Zhe Liu, Chunyang Chen, Junjie Wang, Mengzhuo Chen, Boyu Wu, Xing Che, Dandan Wang, and
234 Qing Wang. Make llm a testing expert: Bringing human-like interaction to mobile gui testing via
235 functionality-aware decisions. In *Proceedings of the IEEE/ACM 46th International Conference on*
236 *Software Engineering*, pp. 1–13, 2024a.

237 Zhe Liu, Chunyang Chen, Junjie Wang, Mengzhuo Chen, Boyu Wu, Zhilin Tian, Yuekai Huang,
238 Jun Hu, and Qing Wang. Testing the limits: Unusual text inputs generation for mobile app
239 crash detection with large language model. In *Proceedings of the IEEE/ACM 46th international*
240 *conference on software engineering*, pp. 1–12, 2024b.

241 Yuxuan Lu, Bingsheng Yao, Hansu Gu, Jing Huang, Zheshen Jessie Wang, Yang Li, Jiri Gesi, Qi He,
242 Toby Jia-Jun Li, and Dakuo Wang. Uxagent: An llm agent-based usability testing framework for
243 web design. In *Proceedings of the Extended Abstracts of the CHI Conference on Human Factors*
244 *in Computing Systems*, pp. 1–12, 2025.

245 Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Gao, Sarah Wiegrefe, Uri Alon,
246 Nouha Dziri, Shrimai Prabhumoye, Yiming Yang, Shashank Gupta, Bodhisattwa Prasad Majumder,
247 Katherine Hermann, Sean Welleck, Amir Yazdanbakhsh, and Peter Clark. Self-refine: Iterative
248 refinement with self-feedback, 2023. URL <https://arxiv.org/abs/2303.17651>.

249 Ali Mesbah, Arie Van Deursen, and Stefan Lenselink. Crawling ajax-based web applications through
250 dynamic analysis of user interface state changes. *ACM Transactions on the Web (TWEB)*, 6(1):
251 1–30, 2012.

252 Magnus Müller and Gregor Žunič. Browser use: Enable ai to control your browser, 2024. URL
253 <https://github.com/browser-use/browser-use>.

254 F. Ricca and P. Tonella. Analysis and testing of web applications. In *Proceedings of the 23rd*
255 *International Conference on Software Engineering. ICSE 2001*, pp. 25–34, 2001. doi: 10.1109/
256 ICSE.2001.919078.

257 Mobina Shahbandeh, Parsa Alian, Noor Nashid, and Ali Mesbah. Navigate: Functionality-guided
258 web application navigation. *arXiv preprint arXiv:2409.10741*, 2024.

259 Salman Sherin, Asmar Muqet, Muhammad Uzair Khan, and Muhammad Zohaib Iqbal. Qexplore:
260 An exploration strategy for dynamic web applications using guided search. *Journal of Systems and*
261 *Software*, 195:111512, 2023.

262 Noah Shinn, Federico Cassano, Edward Berman, Ashwin Gopinath, Karthik Narasimhan, and
263 Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning, 2023. URL
264 <https://arxiv.org/abs/2303.11366>.

265 Andrea Stocco, Alexandra Willi, Luigi Libero Lucio Starace, Matteo Biagiola, and Paolo Tonella.
266 Neural embeddings for web testing. *arXiv preprint arXiv:2306.07400*, 2023.

267 Dakuo Wang, Ting-Yao Hsu, Yuxuan Lu, Hansu Gu, Limeng Cui, Yaochen Xie, William Headean,
268 Bingsheng Yao, Akash Veeragouni, Jiapeng Liu, Sreyashi Nag, and Jessie Wang. Agenta/b:
269 Automated and scalable web a/btesting with interactive llm agents, 2025. URL [https://](https://arxiv.org/abs/2504.09723)
270 arxiv.org/abs/2504.09723.

271 Siyi Wang, Sinan Wang, Yujia Fan, Xiaolei Li, and Yepang Liu. Leveraging large vision-language
272 model for better automatic web gui testing. In *2024 IEEE International Conference on Software*
273 *Maintenance and Evolution (ICSME)*, pp. 125–137. IEEE, 2024.

274 Hao Wen, Shizuo Tian, Borislav Pavlov, Wenjie Du, Yixuan Li, Ge Chang, Shanhui Zhao, Jiacheng
275 Liu, Yunxin Liu, Ya-Qin Zhang, et al. Autodroid-v2: Boosting slm-based gui agents via code
276 generation. *arXiv preprint arXiv:2412.18116*, 2024.

277 Jianwei Yang, Hao Zhang, Feng Li, Xueyan Zou, Chunyuan Li, and Jianfeng Gao. Set-of-mark
278 prompting unleashes extraordinary visual grounding in gpt-4v, 2023. URL [https://arxiv.](https://arxiv.org/abs/2310.11441)
279 [org/abs/2310.11441](https://arxiv.org/abs/2310.11441).

280 Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Thomas L. Griffiths, Yuan Cao, and Karthik
281 Narasimhan. Tree of thoughts: Deliberate problem solving with large language models, 2023.
282 URL <https://arxiv.org/abs/2305.10601>.

283 Juyeon Yoon, Robert Feldt, and Shin Yoo. Autonomous large language model agents enabling
284 intent-driven mobile gui testing. *arXiv preprint arXiv:2311.08649*, 2023.

285 Xiao Yu, Maximillian Chen, and Zhou Yu. Prompt-based monte-carlo tree search for goal-oriented
286 dialogue policy planning, 2023. URL <https://arxiv.org/abs/2305.13660>.

287 Xiao Yu, Baolin Peng, Vineeth Vajipey, Hao Cheng, Michel Galley, Jianfeng Gao, and Zhou Yu.
288 Exact: Teaching ai agents to explore with reflective-mcts and exploratory learning, 2025. URL
289 <https://arxiv.org/abs/2410.02052>.

290 Yan Zheng, Yi Liu, Xiaofei Xie, Yepang Liu, Lei Ma, Jianye Hao, and Yang Liu. Automatic web
291 testing using curiosity-driven reinforcement learning. In *2021 IEEE/ACM 43rd International*
292 *Conference on Software Engineering (ICSE)*, pp. 423–435. IEEE, 2021.

293 Shuyan Zhou, Frank F. Xu, Hao Zhu, Xuhui Zhou, Robert Lo, Abishek Sridhar, Xianyi Cheng,
294 Tianyue Ou, Yonatan Bisk, Daniel Fried, Uri Alon, and Graham Neubig. Webarena: A realistic
295 web environment for building autonomous agents, 2024. URL [https://arxiv.org/abs/](https://arxiv.org/abs/2307.13854)
296 [2307.13854](https://arxiv.org/abs/2307.13854).