

Qiushi Wu

Research Scientist, IBM Research | AI Supply Chain Security Group
qiushi@qiushiwu.com | [Website](#) | [Google Scholar](#) | [DBLP](#) | [ORCID](#)

Research Areas

Software and Systems Security; Program Analysis; AI for Security; Secure Software Supply Chains; Autonomous Vulnerability Discovery and Security of AI Agents.

I build program-analysis and AI systems that discover, prioritize, and mitigate vulnerabilities in large-scale software, from OS kernels and open-source infrastructure to AI-generated code and multi-agent AI systems.

Education

University of Minnesota, Twin Cities Minneapolis, MN Sep. 2018 – Oct. 2023
Ph.D. in Computer Science and Engineering
Advisor: Kangjie Lu
Dissertation: *Automated Code-Behavior and Semantic Understanding for Security*

University of Science and Technology of China Hefei, China Sep. 2013 – Jul. 2018
B.Eng. in Information Security

Research Appointments

IBM Research Yorktown Heights, NY
Research Scientist, AI Supply Chain Security Group Jul. 2023 – Present

- Conduct research on AI security, autonomous vulnerability detection, AI red-teaming, and secure infrastructure for multi-agent AI systems.
- Develop LLM- and agent-based techniques for vulnerability discovery, exploit validation, patch/report generation, secure code generation, and security-patch understanding.

University of Minnesota, Twin Cities Minneapolis, MN
Graduate Research Assistant; advised by Kangjie Lu Sep. 2018 – Oct. 2023

- Developed program-analysis techniques for software and systems security, with emphasis on OS kernels, vulnerability discovery, patch analysis, and code semantic understanding.

IBM Research Yorktown Heights, NY
Research Intern, System and Cloud Security Group May–Aug. 2021; May–Aug. 2022
Advised by Hani Jamjoom and Zhongshu Gu.

- Built precise operating-system kernel call graphs by resolving which functions indirect calls can reach, a precision prerequisite for whole-kernel security analysis.

Research Highlights

- 15 peer-reviewed conference and journal publications, including 11 at the top security venues USENIX Security, NDSS, and ACM CCS, plus papers at ICML, ACSAC, IEEE TDSC, and ESORICS. Google Scholar: 647 citations, h-index 14, i10-index 14 (July 2026).
- Research spans program analysis for OS kernels and foundational software, LLM-assisted security-patch understanding, secure code generation, and autonomous vulnerability discovery.
- **BugStone** (ICML 2026, first author) turns one patched bug into a reusable detection rule: 246 confirmed issues from a 400-case audit of 22,568 flagged Linux kernel candidates, and 20 exploit-validated command injections in top Python repositories, one fixed by the PaddleOCR maintainers.
- Current work extends this line to agentic end-to-end vulnerability discovery, with a manuscript in preparation.
- Security impact includes a granted kernel-analysis patent, public research artifacts, confirmed findings fixed in widely deployed open-source software, and AI-agent red-teaming workflows.

Publications

Peer-Reviewed Conference and Journal Publications

1. **Qiushi Wu**, Yue Xiao, Dhilung Kirat, Kevin Eykholt, Jiyong Jang, and Douglas Lee Schales. “One Bug, Hundreds Behind: LLMs for Large-Scale Bug Discovery.” *International Conference on Machine Learning (ICML)*, 2026.
2. Xingyu Li, Juefei Pu, Yifan Wu, Xiaochen Zou, Shitong Zhu, **Qiushi Wu**, Zheng Zhang, Joshua Hsu, Yue Dong, Zhiyun Qian, Kangjie Lu, Trent Jaeger, Michael De Lucia, and Srikanth V. Krishnamurthy. “What Do They Fix? LLM-Aided Categorization of Security Patches for Critical Memory Bugs.” *Network and Distributed System Security Symposium (NDSS)*, 2026.
3. Weiheng Bai, Keyang Xuan, Pengxiang Huang, **Qiushi Wu**, Jianing Wen, Jingjing Wu, and Kangjie Lu. “APILOT: Improving the Security and Usability of LLM Code Suggestions via Outdated API Mitigation.” *Annual Computer Security Applications Conference (ACSAC)*, 2025.
4. **Qiushi Wu**, Zhongshu Gu, Hani Jamjoom, and Kangjie Lu. “GNNIC: Finding Long-Lost Sibling Functions with Abstract Similarity.” *Network and Distributed System Security Symposium (NDSS)*, 2024. (Conference presenter.)
5. Jianhao Xu, Kangjie Lu, Zhengjie Du, Zhu Ding, Linke Li, **Qiushi Wu**, Mathias Payer, and Bing Mao. “Silent Bugs Matter: A Study of Compiler-Introduced Security Bugs.” *USENIX Security Symposium*, 2023.
6. Qingyang Zhou, **Qiushi Wu**, Dinghao Liu, Shouling Ji, and Kangjie Lu. “Non-Distinguishable Inconsistencies as a Deterministic Oracle for Detecting Security Bugs.” *ACM Conference on Computer and Communications Security (CCS)*, 2022.
7. **Qiushi Wu***, Yue Xiao*, Xiaojing Liao, and Kangjie Lu. “OS-Aware Vulnerability Prioritization via Differential Severity Analysis.” *USENIX Security Symposium*, 2022. (*Co-first authors.)
8. Wenjia Zhao, Kangjie Lu, **Qiushi Wu**, and Yong Qi. “Semantic-Informed Driver Fuzzing Without Both the Hardware Devices and the Emulators.” *Network and Distributed System Security Symposium (NDSS)*, 2022.
9. Dinghao Liu, **Qiushi Wu**, Shouling Ji, Kangjie Lu, Zhenguang Liu, Jianhai Chen, and Qinming He. “Detecting Missed Security Operations Through Differential Checking of Object-based Similar Paths.” *ACM Conference on Computer and Communications Security (CCS)*, 2021.
10. **Qiushi Wu**, Aditya Pakki, Navid Emamdoost, Stephen McCamant, and Kangjie Lu. “Detecting Disordered Error Handling with Precise Function Pairing.” *USENIX Security Symposium*, 2021. (Conference presenter.)
11. Navid Emamdoost, **Qiushi Wu**, Kangjie Lu, and Stephen McCamant. “Practically Detecting Kernel Memory Leaks in Specialized Modules and Beyond.” *Network and Distributed System Security Symposium (NDSS)*, 2021.
12. Bowen Wang*, Kangjie Lu*, **Qiushi Wu**, and Aditya Pakki. “Unleashing Fuzzing Through Comprehensive, Efficient, and Faithful Exploitable-Bug Exposing.” *IEEE Transactions on Dependable and Secure Computing (TDSC)*, 2021. (*Co-first authors.)
13. **Qiushi Wu**, Yang He, Stephen McCamant, and Kangjie Lu. “Precisely Characterizing Security Impact in a Flood of Patches via Symbolic Rule Comparison.” *Network and Distributed System Security Symposium (NDSS)*, 2020. (Conference presenter.)
14. Kangjie Lu, Aditya Pakki, and **Qiushi Wu**. “Detecting Missing-Check Bugs via Semantic- and Context-Aware Criticalness and Constraints Inferences.” *USENIX Security Symposium*, 2019.
15. Kangjie Lu, Aditya Pakki, and **Qiushi Wu**. “Automatically Identifying Security Checks for Detecting Kernel Semantic Bugs.” *European Symposium on Research in Computer Security (ESORICS)*, 2019.

Preprints, Workshops, and Other Publications

1. Kefu Wu, Weiheng Bai, Nanzi Yang, **Qiushi Wu**, and Kangjie Lu. “DesignSleuth: Detecting API Misdesigns via Expectation-Aligned Peer Comparison.” Under submission, 2026.
2. Weiheng Bai, Kefu Wu, **Qiushi Wu**, and Kangjie Lu. “AFLGopher: Accelerating Directed Fuzzing via Feasibility-Aware Guidance.” Under submission; arXiv preprint, 2025.
3. Weiheng Bai, Kefu Wu, **Qiushi Wu**, and Kangjie Lu. “Exploring the Influence of Prompts in LLMs for Security-Related Tasks.” *NDSS Workshop*, 2024.
4. Weiheng Bai, Kefu Wu, **Qiushi Wu**, and Kangjie Lu. “Guiding Directed Fuzzing with Feasibility.” *IEEE European Symposium on Security and Privacy Workshops (EuroS&P Workshops)*, 2023.
5. Weiheng Bai and **Qiushi Wu**. “Towards More Effective Responsible Disclosure for Vulnerability Research.” *EthiCS*, 2023.
6. **Qiushi Wu** and Kangjie Lu. “On the Feasibility of Stealthily Introducing Vulnerabilities in Open-Source Software via Hypocrite Commits.” Withdrawn manuscript, 2021. Account and principal primary sources: qiushiwu.com/#research-ethics.

Research Artifacts, Software, and Security Impact

- **Agentic end-to-end vulnerability discovery** (manuscript in preparation): a system that generalizes a confirmed bug’s pattern, validates candidates dynamically, and drafts patches and reports across large open-source codebases.
- **BugStoneBench** (first-authored, ICML 2026; github.com/code-genome/BugStoneBench): benchmark dataset and evaluation code accompanying “One Bug, Hundreds Behind,” for evaluating LLM-based large-scale bug discovery.
- **GNNIC** (first-authored, NDSS 2024): improves the precision of kernel call graphs by matching functions with similar behavior.
- **SID** (first-authored, NDSS 2020): determines the security impact of patches via symbolic rule comparison.
- **HERO** (first-authored, USENIX Security 2021): detects disordered error-handling bugs with precise function pairing.
- **DIFFCVSS** (co-first-authored, USENIX Security 2022; sites.google.com/view/diffcvss): prioritizes vulnerabilities via OS-aware differential severity analysis.
- **CRIX** (co-authored, USENIX Security 2019; github.com/umnsec/crix): detects missing-check bugs in OS kernels.
- **NDI** (co-authored, CCS 2022; github.com/umnsec/ndi): a deterministic oracle for detecting security bugs.
- **DualLM** (co-authored, NDSS 2026; github.com/seclab-ucr/DualLM): detects security-critical Linux kernel patches, focusing on use-after-free and out-of-bounds vulnerabilities.
- **APILOT** (co-authored, ACSAC 2025; github.com/Wayne-Bai/APILOT): mitigates outdated, vulnerable API usage in LLM code suggestions; awarded ACSAC’s Code Reviewed badge.
- **Automated vulnerability discovery**: contributed to program-analysis research whose findings led to fixes for hundreds of confirmed security bugs in widely deployed open-source software, including the Linux kernel and OpenSSL.
- **LLM security and secure code generation**: IBM Research work on outdated API mitigation, LLM-assisted vulnerability discovery, AI red-teaming, and secure infrastructure for multi-agent AI systems.
- **AI agents for security challenges** (claw-stack.com): public experiments applying AI agents to real CTF-style security tasks, alongside internal AI-agent red-teaming research.

Patents

- **Qiushi Wu**, Zhongshu Gu, and Hani Jamjoom. “Indirect Function Call Target Identification in Software.” U.S. Patent 11,853,751.
- **Qiushi Wu**, Zhongshu Gu, Enriquillo Valdez, and Hani Jamjoom. “Detecting Bugs Using Large Language Models.” U.S. Patent Application 2026/0017174 A1 (pending).

Teaching Experience

University of Minnesota, Twin Cities

Sep. 2022 – Jan. 2023

Graduate Teaching Assistant, CSCI 4271: Development of Secure Software Systems

- Sole graduate TA for a class of approximately 40–60 students; designed several hands-on labs, including a vulnerability-finding lab and an exploitation lab on prepared programs.
- Led lab sessions, held weekly office hours, and graded homework and exams.

University of Minnesota, Twin Cities

Sep. 2018 – Jan. 2019

Graduate Teaching Assistant, CSCI 2021: Machine Architecture and Organization

- One of two graduate TAs for 200+ enrolled students, coordinating a team of 5+ undergraduate TAs.
- Led lab sessions, held office hours, and graded homework and exams.

Mentoring

- **Weiheng Bai** (University of Minnesota): mentored him as a junior lab member through security research projects spanning responsible vulnerability disclosure, directed fuzzing, LLM prompting for security tasks, secure code generation, and API design analysis. We co-authored four publications and two manuscripts under submission, including a two-author EthICS 2023 paper. Now an Applied Scientist at Amazon; Ph.D. defense expected in 2026.
- **IBM Research Ph.D. intern** (Summer 2026, May–Aug.): weekly research mentoring on the security of agentic systems, guiding the research topic and approach for risks in coding/research agents and multi-agent workflows.
- **Collaborative mentoring with Yue Xiao (William & Mary)** (2026–present): help mentor two first-year Ph.D. students’ projects on LLM/agentic methods for CVE analysis and on how agentic coding affects open-source development, contributing vulnerability-research and systems-security direction.

Professional Service

- **Program committee:** IEEE Symposium on Security and Privacy (IEEE S&P), 2024; Network and Distributed System Security Symposium (NDSS), 2025.
- **Organizing committee:** Workshop on Artificial Intelligence System with Confidential Computing (AISCC), co-located with NDSS, 2024.
- **External reviewer:** ACM CCS 2019/2020; ICICS 2019; NDSS 2021.