

BOMIN WEI

📍 Los Angeles, CA | ✉ davidwei23@g.ucla.edu | 🔗 LinkedIn | 🐙 GitHub | 🎓 Google Scholar

Research Interests

Multimodal Reasoning and World Models; Visual Grounding and Continuous Representations; Trustworthy AI and Foundation Model Safety.

Education

University of California, Los Angeles
B.S. in Computer Science | GPA: 3.8/4.0

Expected June 2027
Los Angeles, CA

Selected Publications

- Y. Qin, **B. Wei**, et al. “Chain-of-Visual-Thought: Teaching VLMs to See and Think Better with Continuous Visual Tokens”. In *ECCV 2026*. (Accepted Oral)
- Y. Li, F. H. Shezan, **B. Wei**, G. Wang, and Y. Tian. “SoK: Towards Effective Automated Vulnerability Repair”. In *USENIX Security Symposium 2025*. (Published)
- B. Wei**, Y. Zhang, and X. Gong. “DeepLPI: a novel deep learning-based model for protein–ligand interaction prediction for drug repurposing.” *Scientific Reports*, 2022. 50+ citations. (Published)

Research Experience

Berkeley AI Research (BAIR) – Trevor Darrell Group

2024–Present

Undergraduate Researcher – Multimodal Reasoning and Visual Grounding Mentor: Xudong Wang (now at Physical Intelligence)

- Co-developed **Chain-of-Visual-Thought (CoVT)**, a continuous visual-reasoning framework for VLMs; contributed to problem formulation, methodology design, and large-scale experiments. Accepted to ECCV 2026.
- Identified loss of fine-grained visual evidence in language-dominant VLM reasoning and proposed supervising intermediate hidden states in continuous representation space with external visual encoders.
- Designed the initial DINO-based alignment mechanism and two-stage training strategy for generating and using continuous visual tokens; evaluated alternative conditioning designs through targeted sanity checks.

UCLA Security Lab – Prof. Yuan Tian

2023–Present

Undergraduate Researcher – AI Security, Privacy, and Software Reliability

- Built evaluation pipelines for VulRepair, VulMaster, and VRPilot; showed sharp degradation without oracle bug localization, supporting a USENIX Security SoK on barriers to real-world automated vulnerability repair.
- Studied prompt-injection vulnerabilities in multilingual translation LLMs, hidden-weight-projection fingerprinting, and membership-inference risks in diffusion models.

Realm Labs

Summer 2026

Research Intern – Foundation Model Safety and Observability

- Developed a white-box, decoding-time detector for unsafe multimodal inputs and MLLM outputs, monitoring internal representations for PII leakage and hate speech with over 95% classification accuracy.

Peking University VCL Lab

Summer 2024

Research Assistant – 3D Vision and Representation Learning

- Explored LLM-assisted representation learning and standardization for large-scale 3D visual datasets.

Honors & Awards

UCLA Engineering Scholarship (2025) • UCLA Undergraduate Research Fellows Program (2024) • USACO Gold • S.-T. Yau National High School Science Award, 2nd Place • Mercer Science & Engineering Fair, 1st Place (ISEF Affiliated)