

Yixiao Wang

 yixiao.site |  yixiaowang@sjtu.edu.cn |  +8613390606996

ACADEMIC BACKGROUND

Shanghai Jiao Tong University

09/2022 – 06/2026

B.E. in Software Engineering, School of Computer Science

GPA: 3.92 / 4.3 (Ranking: 6 / 88)

RESEARCH EXPERIENCE

ELORA: Efficient LoRA and KV Cache Management for Multi-LoRA LLM Serving [\[Paper\]](#)

- ELORA is a high-performance cache management system for multi-LoRA LLM serving. By unifying LoRA adapter and KV cache management, it reduces TTFT and TPOT while improving throughput and addressing ineffective cache usage and memory fragmentation in existing systems.
- The system introduces a Dependency-aware Cache Manager and a Performance-driven Cache Swapper, which dynamically manage LoRA and KV caches based on usage dependencies. This enables smarter allocation and utilization of HBM resources, thereby reducing latency and improving throughput.

MOBA: Multifaceted Memory-Enhanced Adaptive Planning for Efficient Mobile Task Automation [\[Paper\]](#)

- MOBA is a memory-enhanced planning framework for efficient mobile task automation. It improves long-horizon task execution by addressing limitations of existing mobile agents, including insufficient contextual understanding, weak task decomposition, and limited adaptability in dynamic environments.
- The system leverages multifaceted memory to capture historical interactions, task experiences, and environment-aware information. Using these memories, MOBA refines planning strategies, enabling reliable action selection and improving completion efficiency across mobile automation scenarios.

BuddyMoE: Exploiting Expert Redundancy to Accelerate Memory-Constrained Mixture-of-Experts Inference [\[Paper\]](#)

- BuddyMoE is an efficient inference system for memory-constrained Mixture-of-Experts models. It targets the latency bottleneck caused by expert offloading, where prefetch failures require expensive expert transfers from CPU memory to GPU memory and significantly slow down inference.
- The system exploits redundancy among MoE experts by replacing missing experts with semantically similar “buddy” experts resident in GPU memory, reducing CPU-GPU transfer stalls while preserving model accuracy, thereby improving large-scale MoE inference efficiency under limited GPU memory.

HONORS & AWARDS

- | | |
|--|--------------------|
| – Ye Jun & Neil Shen Zhiyuan Outstanding Scholarship | June 2026 |
| – Outstanding Undergraduate Thesis Nomination of SCS | June 2026 |
| – Outstanding Graduate of Shanghai | May 2026 |
| – Zhiyuan Honor Scholarship, Shanghai Jiao Tong University | 2022–2025, 4 times |
| – Class C Scholarship, Shanghai Jiao Tong University | 2023–2025, 3 times |

ACADEMIC SERVICE

- | | |
|---|------------------|
| – Teaching Assistant: Programming Design and Methodology (Honor, C++) | Fall 2024–2025 |
| – Teaching Assistant: Probability and Statistics (Honor) | Spring 2025–2026 |