

# ZELEI SHAO

+1 2177661047 | [zelei2@illinois.edu](mailto:zelei2@illinois.edu) | <https://github.com/zelci-hub>

## EDUCATION

|                                                                                          |                                    |                   |
|------------------------------------------------------------------------------------------|------------------------------------|-------------------|
| <b>University of Illinois Urbana-Champaign</b><br><i>Master of Science in Statistics</i> | College of Liberal Arts & Sciences | 08/2024 - now     |
| <b>Peking University</b><br><i>Major in Computational Mathematics</i>                    | School of Mathematical Sciences    | 09/2019 - 07/2023 |

## PUBLICATION

|                                                                                                                                                                                                                                                                                                      |             |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------|
| <b>Beat the long tail: Distribution-Aware Speculative Decoding for RL Training</b><br>Zelei Shao*, Vikranth Srivatsa*, Sanjana Srivastava, Qingyang Wu, Alpaya Ariyak, Xiaoxia Wu, Ameen Patel, Jue Wang, Percy Liang, Tri Dao, Ce Zhang, Yiyang Zhang, Ben Athiwaratkun, Chenfeng Xu, Junxiong Wang | MLSys 2026  |
| <b>MiLo: Effective Quantized MoE Inference with Mixture of Low-Rank Compensators</b><br>Beichen Huang*, Yueming Yuan*, Zelei Shao*, Minjia Zhang                                                                                                                                                     | MLSys 2025  |
| <b>Scaling Speculative Decoding with Lookahead Reasoning</b><br>Yichao Fu, Rui Ge, Zelei Shao, Zhijie Deng, Hao Zhang                                                                                                                                                                                | NeurIPS2025 |

## INTERNSHIP

|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |              |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------|
| <b>Together AI</b>   <i>Research Intern in Turbo team</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               | 06/2025- Now |
| <ul style="list-style-type: none"><li>Designed and implemented <b>DAS</b> (Distribution-Aware Speculative Decoding), a systems-ML co-designed framework that <b>accelerates RL rollout</b> by <b>30–50%</b> while preserving identical training curves.</li><li>Built a self-evolving suffix-tree-based drafter with <b>incremental online updates</b>, enabling high-acceptance draft tokens aligned with evolving RL policies.</li><li>Developed a <b>length-aware draft budget allocator</b>, optimizing compute utilization for long-tail trajectories and reducing straggler makespan during RL rollouts.</li><li>Integrated DAS into <b>VeRL</b> and <b>vLLM</b>, scaling experiments to up to six 8×H100 nodes.</li><li>Co-authored a paper accepted to MLSys 2026, presenting a unified analytical model for optimal speculative budget assignment and suffix-tree-based speculation.</li></ul> |              |

## RESEARCH

|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |                   |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------|
| <b>Efficient Low-Bit Quantization and Inference System for MoE</b><br><i>Group Research, Supervised by Prof.Minjia Zhang, University of Illinois Urbana-Champaign</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        | 07/2024 - 01/2025 |
| <ul style="list-style-type: none"><li>Implemented a highly efficient <b>W3A16 mixed-precision GeMM CUDA kernel</b> based on Marlin, achieving a <b>4.7x speedup</b> v.s. PyTorch on A100</li><li>Designed a zero-bit-waste 3-bit weights packing technique, maximizing memory utilization</li><li>Developed an efficient I2F (INT3-to-FP16) dequantization algorithm leveraging <b>register-level parallelism</b></li><li>Integrated optimizations into MiLo inference backend, achieving <b>3× speedup</b> v.s. Pytorch</li></ul>                                                                                                                                                                                           |                   |
| <b>Scaling Speculative Decoding with Lookahead Reasoning</b><br><i>Group Research, Supervised by Prof.Hao Zhang, University of California San Diego</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | 03/2025 - 05/2025 |
| <ul style="list-style-type: none"><li>Co-designed <b>Lookahead Reasoning</b>, a novel step-level speculation method that scales speculative decoding (SD) along a new dimension orthogonal to token-level approaches</li><li>Provided a <b>theoretical analysis</b> of the technique both in isolation and in combination with token-level SD, showing that under finite concurrency, peak performance is achieved only when both are used together</li><li>Designed and implemented ablation studies on the impact of draft tree width and depth</li><li>Integrated Lookahead Reasoning into <b>vLLM</b> and benchmarked the system, demonstrating up to <b>1.5× additional speedup</b> over token-level SD alone</li></ul> |                   |

## SKILL

**Programming Language:** Python, CUDA, C++  
**Tools&Framework:** Pytorch, vLLM, VeRL

## AWARDS AND HONORS

|                                                                              |      |
|------------------------------------------------------------------------------|------|
| • Merit Student, Peking University                                           | 2021 |
| • Half-Dream Scholarship, School of Mathematical Sciences, Peking University | 2021 |
| • <b>1st Prize</b> in the National High School Mathematics Competition       | 2018 |
| • <b>2nd Prize</b> in Chinese Mathematical Olympiad                          | 2018 |