

Hanmin Li

Ph.D. Candidate in Computer Science | E-MAIL: hanmin.li@kaust.edu.sa | [Google Scholar](#) | [Website](#) | [GitHub](#) | [LinkedIn](#)

Optimization for ML | Distributed Training | Model Compression | LLM efficiency | LLM Agents | LLM Training

PROFILE

Ph.D. candidate in CS at [King Abdullah University of Science and Technology](#) (KAUST), supervised by [Prof. Peter Richtárik](#).

- **Optimization Foundations:** convex and non-convex optimization, first order methods, proximal and ball-proximal methods, trust-region type methods, linear minimization oracles (LMOs), constrained minimization, Frank–Wolfe type methods, matrix-valued and geometry-aware optimization.
- **Scalable learning systems:** LLM optimizer design, communication-efficient systems, and federated optimization, PyTorch distributed training, and large-scale distributed training and inference for dense and Mixture-of-Experts (MoE) models with hybrid parallelism (Data / Tensor / Pipeline / Expert / Context).
- **Agentic AI, post-training, and model efficiency:** Supervised Fine-Tuning (SFT), Reinforcement Learning from Human Feedback (RLHF), with Verifiable Rewards (RLVR), policy optimization methods (PPO, GRPO) long horizon agents, RL with evolving rubrics for DeepResearch and DeepWork, Multi-Turn Agent Systems and Custom Harness Design, LLM-as-a-judge evaluation, vLLM-based serving, activation steering, model compression, and MoE pruning.

SKILLS

- **LLM Training & Post-Training:** Pretraining, SFT, RLHF/RLVR, preference optimization (DPO), policy optimization (PPO, GRPO, DAPO, GSPO), on-policy distillation, and reward/evaluation design.
- **Agent Systems & Evaluation:** Long-horizon multi-turn DeepWork agents, tool/runtime orchestration, evolving rubrics, custom harnesses, and LLM-as-a-judge.
- **Distributed & MoE Systems:** PyTorch DDP/FSDP, DeepSpeed ZeRO, Hugging Face Transformers/Accelerate, Slurm, vLLM, and Megatron DP/TP/PP/EP/CP.
- **Model Efficiency:** MoE routing telemetry, expert pruning, model compression, steering, checkpoint surgery.
- **Optimization:** Convex/non-convex optimization, proximal and trust-region methods, LMO, Frank-Wolfe, matrix stepsizes, gradient compression, and federated learning.
- **Programming & ML Systems:** Python, PyTorch, NumPy, pandas, Matplotlib, C++, Bash, Git, Linux, CUDA, VeRL, Weights & Biases, Jupyter, and LaTeX.

WORK EXPERIENCE

Applied Scientist Intern, Microsoft AI |

June 2026 - Present

Focus: Long-horizon LLM agents, MoE pruning, activation steering, and post-training.

- **Agent platform:** Architected an end-to-end DeepWork tool-use platform on GDPval for Qwen3.6-35B-A3B across 8 x A100 nodes, with a custom harness, trajectory capture, artifact validation, and LLM-as-a-judge evaluation.
- **MoE observability and pruning:** Built exact token-level expert telemetry and deployable pruned checkpoints; reduced routed experts from 256 to 160 per layer, removed about 12.1B parameters, and cut BF16 model-weight memory by 34% while preserving shared experts and top-k routing.
- **Specialization analysis:** Quantified task-dependent expert specialization across sectors and modalities using divergence, dispersion, and layerwise decomposition; specialized expert banks outperformed one-size-fits-all pruning.
- **Post-training interventions:** Developed and Evaluated on policy distillation and activation steering from successful teacher rollouts and pruned-model failures, targeting repetitive behavior loops with layerwise interventions.

EDUCATION

➤ **King Abdullah University of Science and Technology | Thuwal, Saudi Arabia**
Ph.D. Candidate in Computer Science

Advisor: Prof. Peter Richtárik

(Expected: Spring 2027)

King Abdullah University of Science and Technology | Thuwal, Saudi Arabia
M.Eng. in Computer Science - GPA: 3.86/4.00

Dec 2022

- **University of Science and Technology of China | Hefei, China**
B.Eng. in Computer Science and Technology, School of the Gifted Young - GPA: 3.64/4.30
- **The University of Texas at Austin | Austin, USA**
Summer School in Software Engineering

July 2021

Aug 2018

PROJECTS & PUBLICATIONS

Efficient Layerwise Optimizers for Scalable LLM Training (Ongoing Research)

Relevant paper (1): [Local LMO \(Preprint, 2026\)](#)

Techs: PyTorch, FSDP, DeepSpeed ZeRO, Hybrid Parallelisms, LLM pretraining, SFT, distributed training, optimizer design

- **Built** scalable Muon-style layerwise optimizers with row- and column-wise block decomposition, top-k block selection, and magnitude-based update sparsification to improve the efficiency and stability of large-model training. **Established** theoretical connections between Muon variants and local linear minimization oracles, showing how layerwise optimizer updates can be interpreted through constrained and projection-free optimization.
- Benchmarked AdamW, and Muon variants on pretraining and SFT workloads spanning GPT-2 124M, Qwen2-7B, and built multi-GPU training pipelines with ZeRO-3, gradient checkpointing, mixed precision.

Layerwise Matrix Stepsizes for Compressed Distributed Training

Relevant papers (2): [Det-CGD \(ICLR 2024\)](#) | [Variance-Reduced Matrix Stepsizes \(NeurIPS FL@FM 2023\)](#)

Techs: PyTorch, efficient LLM training, gradient compression, distributed training, matrix-valued step size.

- **Introduced** det-CGD with layerwise Bernoulli-style gradient compression and matrix-valued step sizes, reducing communication with no additional computational overhead and establishing free-compression regimes without degraded iteration complexity. **Extended** the framework by incorporating momentum-based variance reduction.
- **Designed** an adaptive layerwise optimizer for NanoGPT 124M pretraining, estimating smoothness from warm-up trajectories and dynamically updating matrix-valued stepsizes throughout training to reflect layer-specific optimization geometry, consistently outperforming standard SGD.

Server Extrapolation for Communication-Efficient Distributed Training

Relevant papers (2): [The Power of Extrapolation \(NeurIPS 2024\)](#) | [FedProx with Inexact Prox \(NeurIPS OPT-ML 2024\)](#)

Techs: PyTorch, distributed training, efficient training, proximal methods, server-side acceleration, inexact local solvers

- **Proposed** constant and adaptive server extrapolation to accelerate distributed training under heterogeneous local objectives and extended the framework to inexact local solves while preserving exact convergence under relative-accuracy conditions across convex and weakly convex regimes.
- **Benchmarked** the framework on CIFAR and Shakespeare benchmarks, achieving up to 2× faster convergence than non-extrapolated baselines with no additional runtime overhead, while remaining robust to coarse local approximations.

Ball Proximal and Trust-Region Methods for Robust Layerwise Non-Convex Optimization

Relevant papers (3): [Ball-Proximal Point Method \(2025\)](#) | [Stabilized PPM \(2026\)](#) | [Broximal Alignment \(2026\)](#)

Techs: PyTorch, LLM pretraining, distributed training, optimizers, proximal methods, trust region methods

- **Developed** a family of ball-constrained proximal methods that prevent vanishing updates and strengthen convergence from stationarity to global-minimizer guarantees under alignment conditions, while clarifying connections to trust-region methods, proximal algorithms, SGD-style updates, linear minimization oracles, and Muon-style optimizers.
- **Validated** the framework on NanoGPT (124M) pretraining, demonstrating fast and stable convergence and using trajectory smoothness as a quantitative diagnostic of optimizer stability.

HONORS & AWARDS

- CEMSE Dean's List Award - KAUST, 2026 - top 20%
- Outstanding Ph.D. Annual Evaluation - KAUST, 2023 - 2024.
- Invited as a visiting scholar to [ELLIIT Focused Period Lund 2026 Optimization for Learning](#) - 2026, Lund, Sweden.
- Invited as a visiting scholar to [EUROPT 2024](#) - 2024, Lund Sweden.
- KAUST Graduate Student Fellowship – KAUST, 2021 – 2026
- Scholarship for Outstanding Students - School of the Gifted Young, USTC, 2018-2019 - top 20%.