

Hanmin Li 李瀚民

Ph.D. Candidate in Computer Science

Optimization · Distributed Training · Model Compression · Efficient LLM · LLM Agents · LLM Training

hanmin.li@kaust.edu.sa | [Website](#) | [Google Scholar](#) | [GitHub](#) | [LinkedIn](#)

ABOUT ME

CURRENT POSITION Ph.D. Candidate in Computer Science, Center of Excellence for Generative AI;
Supervised by [Prof. Peter Richtárik](#), King Abdullah University of Science and Technology (KAUST)

Current research themes

- **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 Parallelism)
- **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, DPO, GRPO) long-horizon tool-using 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.

EDUCATION

PRESENT	Ph.D. Candidate in Computer Science KAUST , Thuwal, Saudi Arabia. Advisor: Prof. Peter Richtárik
DEC. 2022	Master of Engineering in Computer Science KAUST , Thuwal, Saudi Arabia. GPA: 3.86/4.00.
JULY 2021	Bachelor of Engineering in Computer Science and Technology School of the Gifted Young (少年班学院), University of Science and Technology of China (USTC), Hefei, China. GPA: 3.64/4.30.
AUG. 2018	Summer School, Software Engineering The University of Texas at Austin, Austin, Texas, USA.

PUBLICATIONS AND PREPRINTS

PREPRINT 2026	Local LMO: Constrained Gradient Optimization via a Local Linear Minimization Oracle Peter Richtárik, Kaja Grunkowska, Hanmin Li arXiv:2605.08850 <i>Contribution:</i> Develops a local linear minimization oracle (LMO) that extends the classical LMO framework and bridges gradient descent with projection-free constrained optimization, while retaining projected-gradient-type convergence guarantees across a broad range of settings. <i>Keywords:</i> projection-free methods; local Linear Minimization Oracles; Frank-Wolfe methods; first-order methods; constrained optimization.
PREPRINT 2026	Broximal Alignment for Global Non-Convex Optimization Kaja Grunkowska, Hanmin Li, Xun Qian, Peter Richtárik arXiv:2604.13483 <i>Contribution:</i> Introduces Broximal Alignment, a structural framework establishing global convergence of the Ball-Proximal Method to a global minimizer rather than merely a stationary point, without requiring convexity, smoothness, or Lipschitz continuity. <i>Keywords:</i> global non-convex optimization; beyond stationarity; non-convex alignment; ball proximal methods; global convergence
PREPRINT 2026	Stabilized Proximal Point Method via Trust Region Control Hanmin Li, Kaja Grunkowska, Peter Richtárik arXiv:2604.02943 <i>Contribution:</i> Introduces a trust-region-stabilized Proximal Point Method that prevents step collapse and establishes linear descent outside any prescribed neighborhood, without smoothness or strong convexity. <i>Keywords:</i> trust-region methods; proximal algorithms; nonsmooth optimization; modern optimizer dynamics

PREPRINT
2025

The Ball-Proximal (“Broximal”) Point Method: A New Algorithm, Convergence Theory, and Applications

Kaja Gruntkowska, Hanmin Li, Aadi Rane, Peter Richtárik

[arXiv:2502.02002](#)

Contribution: Introduces the Ball-Proximal Point Method, replacing quadratic proximal regularization with ball constraints and establishing linear and finite-step convergence in nonsmooth convex optimization, with extensions through ball-convexity.

Keywords: ball-proximal methods; trust regions; nonsmooth optimization; adaptive stepsizes

CONFERENCE
2024

The Power of Extrapolation in Federated Learning

Hanmin Li, Kirill Acharya, Peter Richtárik

[NeurIPS 2024](#)

Contribution: Introduces constant and smoothness-adaptive server extrapolation for FedProx, establishing improved convergence guarantees under gradient diversity and demonstrating consistent acceleration across distributed training and federated learning benchmarks.

Keywords: proximal point method; server extrapolation; adaptive stepsizes; distributed training; efficient training

WORKSHOP
2024

On the Convergence of FedProx with Extrapolation and Inexact Prox

Hanmin Li, Peter Richtárik

[NeurIPS OPT-ML Workshop poster](#);

Contribution: Establishes convergence of FedExProx with inexact client-side proximal solves, showing that relative accuracy preserves exact convergence and extrapolation benefits across strongly convex, convex, and weakly convex regimes.

Keywords: inexact proximal methods; distributed optimization; server extrapolation; bias reduction; biased compression; efficient training

CONFERENCE
2024

Det-CGD: Compressed Gradient Descent with Matrix Stepsizes for Non-Convex Optimization

Hanmin Li, Avetik Karagulyan, Peter Richtárik

[ICLR 2024](#)

Contribution: Introduces compressed gradient methods with layerwise matrix-valued stepsizes for communication efficient distributed training, establishing free-compression regimes (layerwise Bernoulli-style compressor) and outperforming scalar-step baselines in training complexity.

Keywords: matrix preconditioning; gradient compression; layerwise optimization; communication efficient distributed training

WORKSHOP
2023

Variance Reduced Distributed Non-Convex Optimization Using Matrix Stepsizes

Hanmin Li, Avetik Karagulyan, Peter Richtárik

[NeurIPS FL@FM Workshop poster](#)

Contribution: Introduces det-MARINA and det-DASHA, combining matrix-valued stepsizes and gradient compression with MARINA-style and momentum-based variance reduction to improve training efficiency.

Keywords: distributed training; momentum-based variance reduction; matrix-valued stepsizes; gradient compression; efficient training

JOURNAL
2022

SD²: Spatially Resolved Transcriptomics Deconvolution Through Integration of Dropout and Spatial Information

Haoyang Li, Hanmin Li, Juexiao Zhou, Xin Gao

[Bioinformatics 38\(21\), 4878–4884](#)

Contribution: Builds an autoencoder and graph-convolutional pipeline that combines dropout-derived features with spatial relationships to estimate cell-type composition in spatial transcriptomics.

Keywords: graph convolutional networks; spatial omics; representation learning; computational biology

WORK EXPERIENCE

JUNE 2026
-
PRESENT

Applied Scientist Intern, Microsoft AI

Long-horizon LLM agents, Mixture-of-Experts model pruning, activation steering, and post-training.

- Architected an end-to-end platform for DeepWork long-horizon tool-use agents on the GDPval benchmark, serving Qwen3.6-35B-A3B on nodes of 8 A100 GPUs each and integrating a customized agent harness, trajectory capture, artifact validation, and LLM-as-a-judge evaluation.
- Built exact token-level MoE observability across complete agent trajectories and converted routing statistics into deployable pruned checkpoints, reducing routed-expert banks from 256 to 160 experts per layer, eliminating approximately 12.1 billion parameters and cutting BF16 model-weight memory by 34% while preserving shared experts, top-*k* execution, and benchmark performance.
- Discovered and quantified task-dependent expert specialization across industry sectors and output modalities, using routing-distribution divergence, dispersion analysis, and layerwise decomposition to construct specialized expert banks that outperform one-size-fits-all pruning strategies.
- Developed and evaluated on-policy distillation (OPD) and contrastive activation steering as complementary

methods for recovering capabilities lost under aggressive MoE pruning, using successful teacher rollouts and pruned-model failure trajectories to address repetitive undesirable behavior loops and derive targeted layerwise interventions.

INVITED TALKS AND SELECTED PRESENTATIONS

12 MAY 2026 Lund, Sweden	Stabilizing PPM: Trust Regions, Linear Descent, and Connections to Modern ML Optimizers Invited talk, ELLIIT Focus Period: Optimization for Learning
15 DEC. 2024 Vancouver, Canada	On the Convergence of FedProx with Extrapolation and Inexact Prox Poster, NeurIPS OPT-ML Workshop
11 DEC. 2024 Vancouver, Canada	The Power of Extrapolation in Federated Learning Poster, NeurIPS 2024
26 JUNE 2024 Lund, Sweden	Compressed Gradient Descent with Matrix Stepsizes for Non-Convex Optimization Invited talk, EUROPT 2024
7 MAY 2024 Vienna, Austria	Det-CGD: Compressed Gradient Descent with Matrix Stepsizes for Non-Convex Optimization Poster, ICLR 2024
16 DEC. 2023 New Orleans, USA	MARINA Meets Matrix Stepsizes: Variance-Reduced Distributed Non-Convex Optimization Poster, FL@FM-NeurIPS 2023
16 DEC. 2023 New Orleans, USA	Det-CGD: Compressed Gradient Descent with Matrix Stepsizes for Non-Convex Optimization Poster, NeurIPS 2023 OPT Workshop

ACADEMIC SERVICE

CONFERENCES	Reviewer: NeurIPS (2024–2026), ICLR (2025–2026), and ICML (2025–2026).
JOURNALS	Reviewer: Transactions on Machine Learning Research (TMLR), Journal of Machine Learning Research (JMLR), IEEE Transactions on Neural Networks and Learning Systems, IEEE Transactions on Signal Processing, and Optimization Methods and Software.
WORKSHOPS	Reviewer: NeurIPS OPT-ML Workshop (2024).

HONORS AND AWARDS

2026	CEMSE Dean's List Award King Abdullah University of Science and Technology; top 20%.
2023 - 2024	Outstanding Ph.D. Annual Evaluation King Abdullah University of Science and Technology.
2021 - 2026	KAUST Graduate Student Fellowship King Abdullah University of Science and Technology.
2018 - 2019	Scholarship for Outstanding Students School of the Gifted Young, USTC; top 20%.
2017	Shitsan Pai Class Scholarship University of Science and Technology of China; top 10%.

TECHNICAL SKILLS AND LANGUAGES

LLM TRAINING, AGENTS, AND EFFICIENCY	Pretraining and supervised fine-tuning; RLHF and RLVR; preference optimization (DPO); policy optimization (PPO, GRPO, DAPO, GSPO); long-horizon multi-turn agents; agent runtime design, and tool orchestration; rubric-driven training and LLM-as-a-judge evaluation; vLLM serving; MoE routing analysis, expert pruning, model compression, and activation steering.
DISTRIBUTED TRAINING AND INFRA	PyTorch DDP and FSDP; DeepSpeed ZeRO; Hugging Face Transformers and Accelerate; Slurm-based multi-GPU and multi-node workflows; Megatron-style hybrid parallelism for dense and Mixture-of-Experts models (DP, TP, PP, EP, CP).
OPTIMIZATION	Convex and non-convex optimization; proximal and trust-region methods; linear minimization oracles and Frank–Wolfe methods; matrix stepsizes; gradient compression; federated optimization.
PROGRAMMING AND ML SYSTEMS	Python (PyTorch, NumPy, pandas, Matplotlib), C++, Bash, Git, Linux, and CUDA; distributed LLM post-training with VeRL; inference with vLLM; custom agent evaluation and execution harnesses; experiment tracking with Weights & Biases; Jupyter and LaTeX/Overleaf.
LANGUAGES	Mandarin Chinese (native), English (fluent; TOEFL 110, GRE 333), French (intermediate).

ADDITIONAL INTERESTS

PADI Advanced Open Water Diver; regular hiking and international travel.