

Hongcan Guo

hongcanguo@gmail.com

+86 18871755323

[Homepage](#) (Not updated in time)

[Google Scholar](#)



EDUCATION

Beijing University of Posts and Telecommunications (BUPT), B.Sc. in Artificial Intelligence 2022 – 2026

Average score: 86.5

INTERNSHIP EXPERIENCE

Institute of Artificial Intelligence, China Telecom, Research Intern Oct. 2024 – May 2025

- *Description:* Engaged in advancing Large Language Model (LLM) reasoning capabilities, focusing on replicating o1 and r1 pipelines and developing efficient reinforcement learning (RL) algorithms.
- *Responsibilities:* Initially responsible for Chain-of-Thought (CoT) SFT and LLM distillation strategies. Since 2025, focused on RL algorithm design for distilled models, aiming to push reasoning frontiers.
- *Achievements:* Delivered T1 development for TeleAI in Dec. 2024 by implementing CoT and distillation pipelines. Reproduced deepseek-R1-Zero (7B/32B) using the GRPO algorithm in Apr. 2025. In May 2025, significantly improved a 7B distilled model via a self-improved DAPO framework; currently preparing a related publication.

Seed Vision, ByteDance, Research Intern Jun. 2025 – Present

- *Description:* Investigating novel Diffusion LLM architectures to overcome efficiency bottlenecks of existing auto-regressive LLMs.
- *Responsibilities:* Due to confidentiality, further details cannot be disclosed. My work involves constructing a complete training and evaluation pipeline, analyzing the problem from a mathematical perspective, and implementing the improvements and code.
- *Achievements:* Currently preparing a paper and will open-source the entire project, with a target submission to ICLR 2026.

ACADEMIC PUBLICATIONS

Advancing Expert Specialization for Better MoE 2025

- *Description:* To address the conflict between load balancing and performance in fine-tuning Mixture-of-Experts (MoE) models for downstream tasks, we propose a meticulously designed auxiliary loss function, $\mathcal{L}_{balance}$. This approach decouples the conflict by specializing experts and routers, significantly boosting performance while maintaining load balance.
- *Status:* **First author, submitted to NeurIPS 2025 (CCF-A)**. The paper has garnered attention from several leading research institutions including ByteDance, Alibaba, Baidu, and Gaoling at RUC. See <https://arxiv.org/abs/2505.22323>.

Two Is Better Than One: Rotations Scale LoRAs 2025

- *Description:* We identify a generalization bottleneck in existing architectures that combine multiple LoRA modules. To address this, we propose *RadarGate*, which introduces rotational degrees of freedom to expand the expressiveness of LoRA's convex cones, enhancing fitting capability and scalability.
- *Status:* **First author, submitted to NeurIPS 2025 (CCF-A)**. See <https://arxiv.org/abs/2505.23184>.

VideoMiner: Iteratively Grounding Key Frames of Hour-Long Videos via Tree-based Group Relative Policy Optimization 2025

- *Description:* We propose a complete keyframe extraction framework for long video understanding and introduce T-GRPO, an on-policy RL method to train its policy model. Our approach effectively identifies keyframes in hour-long videos, improving spatio-temporal comprehension.
- *Status:* **First author, ICCV 2025 (CCF-A), Accepted.**

Advancing Compositional LLM Reasoning with Structured Task Relations in Interactive Multimodal Communications 2025

- *Description:* To enhance the compositional reasoning ability of multimodal LLMs, we propose ContextLoRA, which leverages task dependency graphs to guide structured LoRA partitioning and fine-tuning. This enables latent inter-task relational modeling for more adaptive multimodal reasoning.
- *Status:* **First author, JSAC (CCF-A, SCI Q1, IF=13.6) Accepted.**

Exploring LLM-Based Multi-Agent Situation Awareness for Zero-Trust Space-Air-Ground Integrated Network 2025

- *Description:* To address key challenges in SAGINs—threat perception, adaptive security assessment, and decision-making, we develop the SAG-Attack simulator and LLM-SA framework. By leveraging LLM-based multi-agent collaboration, our approach enables zero-trust situational awareness.
- *Status:* **Third author, JSAC (CCF-A, SCI Q1, IF=13.6) published.** See <https://ieeexplore.ieee.org/document/10963886>.

SKILLS

- **LLM Pre-training:** Familiar with the LLM pre-training pipeline, with experience in building pre-training codebases and training on large-scale computational resources.

- **LLM Post-training:** Proficient in the complete LLM post-training pipeline, including SFT and RL stages. Experienced in data construction, parameter tuning, algorithm design, and practical implementation with strong innovation capability.
- **Reinforcement Learning:** Deep understanding of RL theory, algorithm derivations, and optimization logic. Skilled in implementing and customizing mainstream algorithms such as PPO and GRPO.
- **Generative Model:** Familiar with the foundational architectures of generative models such as Diffusion and VAEs, proficient in their mathematical principles for theoretical analysis and innovation. Also familiar with the source code of models like DDPM, VAE, and Flow Matching.
- **Machine Learning Theory:** Strong foundation in ML fundamentals, including generalization theory and core algorithm.
- **Mathematical Foundations:** Proficient in combinatorics, probability theory, matrix analysis, and convex optimization. Adept at theoretical modeling for algorithmic improvements.

RESEARCH INTERESTS

- **LLM Pre-training:** Techniques for LLM pre-training, including the complete training pipeline, data composition strategies, model architecture, algorithms, and key observational metrics.
- **LLM Post-training:** Full-stack optimization of LLMs via SFT, RL, and efficient fine-tuning.
- **Novel LLM Architecture Design:** Focus on novel architectures like Mixture-of-Experts (MoE) and Diffusion LLMs to push the boundaries of foundation models in scale, performance, and generalization.
- **Reasoning Enhancement:** Reinforcement learning-based methods for enhancing reasoning in multimodal LLMs, committed to achieving the goal of Artificial Super Intelligence (ASI).
- **Unified Architecture for Generation and Understanding:** Explore unified architectures that go beyond full-modal understanding and generation, committed to achieving the goal of Artificial General Intelligence (AGI).
- **Autonomous Learning Frameworks:** Novel learning paradigms beyond traditional supervised/RL settings, including multi-agent collaboration and unsupervised automation.