

Gaozhuo Liu

(+86) 1870-853-5142 | 3255946590@bupt.edu.cn

EDUCATION

• Beijing University of Posts and Telecommunications

Sep 2023 – July 2027

Bachelor of Science (Engineering)

Major: Telecommunication Engineering

Average Mark: 87.93/100; GPA: 3.59/4.0; Rank in 20%

PUBLICATIONS

- [1] ICML 2026 Lightweight Federated Incremental Learning via Decoupled Replay.
Xiuying Wang, Yichen Li, Hang Su, **Gaozhuo Liu**, Shiwei Li, Chuang Zhao, Jiangming Shi, Imran Razzak

UNDER REVIEW

- [2] NeurIPS 2026 Heterogeneity-aware Distillation for Federated Continual Learning.
Scores: 5 / 4 / 4 **Gaozhuo Liu**, Yichen Li, Xiuying Wang, Yulong Li, Chuang Zhao, Yankai Jiang, Imran Razzak

- [3] NeurIPS 2026 Personalized Safety in Federated Fine-Tuning of Large Language Models.
Scores: 5 / 5 / 4 Tianzhe Xiao, **Gaozhuo Liu**, Yichen Li, Haozhao Wang, Hanlin Cai, Ruixuan Li, Ozgur B. Akan

- [4] AAAI 2027 Invisible Channels, Visible Risks: Attacking and Defending Latent Communication in Multimodal Federated Multi-Agent Systems.

Tianzhe Xiao, **Gaozhuo Liu**, Yichen Li, Haozhao Wang, Xiong Yongfu, Daiming Kuang, Yi Wang, YI LIU, Ruixuan Li

- [5] AAAI 2027 SafeFedSkill: Poison-Resilient Federated Skill Evolution through Structured Consensus.
Tianzhe Xiao, **Gaozhuo Liu**, Yankai Jiang, Bo Liu, Yichen Li, Imran Razzak

RESEARCH EXPERIENCE

• Intelligent and Distributed Computing Lab, HUST

Mar 2024 – Present

Working closely with [Yichen Li](#)

Wuhan, China

HUST Ph.D. Candidate; Visiting Student at the University of Cambridge and MBZUAI

- Heterogeneity-aware Federated Continual Learning [2]: led research on Ha-FCL, a distillation-based framework for robust federated continual learning under coupled data, task, and model heterogeneity.
- Lightweight Federated Incremental Learning [1]: contributed to Li-FIL, a privacy-preserving feature-replay framework that mitigates catastrophic forgetting while reducing memory and communication overhead.
- Personalized Safety in Federated LLMs [3]: translated a policy-conditioned safety framework into a complete federated LLM fine-tuning pipeline for client-specific safety alignment and safe-helpful evaluation.
- Poison-Resilient Agent Skill Evolution [5]: engineered and benchmarked SafeFedSkill, a structured consensus system for poison-resilient skill sharing, evidence tracking, and safe local rebinding in tool agents.
- Multimodal Multi-Agent Security [4]: built a latent KV-cache security testbed for multimodal agents to characterize representation-level attacks and validate integrity-gating and multi-view safety defenses.

RESEARCH INTERESTS

- **Federated & Distributed Learning:** Heterogeneous and personalized federated learning, privacy-preserving distributed learning, federated continual/incremental learning, and safety-aware federated LLM adaptation.
- **Trustworthy LLMs & Agent Systems:** Personalized safety alignment, safe model adaptation, tool-agent security, memory and skill evolution, provenance, and rollback.
- **Multimodal Learning & Agent Security:** Vision-language models, multimodal reasoning, latent communication, KV-cache security, and trustworthy multi-agent collaboration.
- **Knowledge Distillation & Representation Learning:** Feature/logit distillation, representation alignment, model adaptation, and efficient knowledge transfer across heterogeneous models and agents.

SKILLS

- **Programming:** Python, C/C++, MATLAB, Bash, LaTeX
- **Deep Learning & LLMs:** PyTorch, Transformers, LLMs, LoRA/PEFT fine-tuning, representation learning
- **Tools:** Linux, Git, Conda, CUDA, Docker, distributed training
- **Sports:** Good at Volleyball, Table Tennis, and Swimming

LANGUAGE

- **English:** CEFR B2
- **Chinese:** Native