

Peirong Zheng

PHD CANDIDATE IN COMPUTING · EFFICIENT VISION-LANGUAGE-ACTION SYSTEMS

Hong Kong SAR

+852 9283 3476 | peirong.zheng@connect.polyu.hk

Education

The Hong Kong Polytechnic University (PolyU)

PHD IN COMPUTING

- Best Poster Award, COMP 50th Anniversary Research Student Conference, 2024.

Hong Kong SAR

Jan. 2024 - Present

Nanyang Technological University

MSC IN COMMUNICATIONS ENGINEERING

Singapore

Aug. 2022 - Dec. 2023

University of Electronic Science and Technology of China

BENG IN COMMUNICATION ENGINEERING | CGPA: 3.65/4.00 | YINGCAI HONORS PROGRAM

Chengdu, China

Sep. 2018 - Jul. 2022

University of Cambridge and University of Oxford

VISITING STUDENT | SUMMER INSTITUTE

Cambridge and Oxford, UK

Aug. 2019

Research Experience

AR-VLA: Conditional Trajectory Reuse for Efficient Autoregressive Vision-Language-Action Inference

Hong Kong SAR

CURRENT RESEARCH PROJECT

2026 - Present

- Developed a training-free conditional trajectory-reuse framework for frozen autoregressive VLA policies, combining per-task dual-view VLA-native memory, threshold-gated top-2 future-action aggregation, and periodic model anchors.
- Designed kinematic descent and grasp-transition guards to truncate unsafe reuse; served **64.1%** of actions from memory and achieved **3.1%** better average success rate on LIBERO tasks with a **2.08×** end-to-end speedup.

Efficient LLM Systems Research

Hong Kong SAR

PHD RESEARCHER | ADVISOR: PROF. WENCHAO XU, CURRENTLY AT HKUST

Jan. 2024 - Present

- Achieved up to **3.72×** **decoding speedup** over state-of-the-art methods while maintaining nearly lossless accuracy.
- HALO and Coda (INFOCOM 2026; MobiCom Poster 2025):** Developed semantic- and context-aware distributed LLM inference that relaxes synchronization under lossy edge networks.
- Implemented and evaluated LLaMA-series inference on Raspberry Pi clusters; HALO delivered up to **3.41×** **end-to-end speedup** under packet loss, while Coda achieved up to **5.62×** **speedup**.

Technical Skills

Programming

Codex, PyTorch, C++, MATLAB

Systems

OpenVLA, LIBERO, trajectory retrieval, distributed inference, GPU/CPU expert offloading

Research Areas

Vision-language-action models, efficient embodied AI, LLM inference acceleration

Publications

Self-Speculative Decoding for On-device MoE Acceleration

WWW '26

PEIRONG ZHENG, WENCHAO XU, AND HAOZHAO WANG

DOI

HALO: Semantic-Aware Distributed LLM Inference in Lossy Edge Network

INFOCOM '26

PEIRONG ZHENG, WENCHAO XU, HAOZHAO WANG, JINYU CHEN, AND XUEMIN SHERMAN SHEN

Poster: Coda: Context-aware Acceleration for Distributed LLM Decoding in Edge

MobiCom '25

PEIRONG ZHENG, WENCHAO XU, HAOZHAO WANG, AND JINYU CHEN

DOI

Deploying Foundation Model Powered Agent Services: A Survey

IEEE CST

WENCHAO XU, JINYU CHEN, PEIRONG ZHENG, XIAOQUAN YI, TIANYI TIAN, WENHUI ZHU, QUAN WAN, HAOZHAO WANG, YUNFENG FAN, QINLIANG SU, AND XUEMIN SHEN

DOI

FedNLR: Federated Learning with Neuron-wise Learning Rates

KDD '24

HAOZHAO WANG*, PEIRONG ZHENG*, XINGSHUO HAN, WENCHAO XU, RUIXUAN LI, AND TIANWEI ZHANG

DOI

A Semi-supervised Algorithm for Atrial Fibrillation Attack Prediction

EMBC '23

YUCHEN JIANG*, PEIRONG ZHENG*, AND DAKUN LAI | *EQUAL CONTRIBUTION

DOI