



The Hong Kong University of Science and Technology

Department of Mathematics

PhD THESIS EXAMINATION

On the Architecture and Reasoning of Large Language Models

By

Mr. Tianhao CHEN

ABSTRACT

This thesis investigates two complementary axes for improving Large Language Models (LLMs): the architecture used during pretraining, and the reasoning behavior elicited during post-training.

The first part of the thesis addresses well-documented optimization pathologies of deep Transformer language models that limit the effective utilization of depth. We first present Gradient-Preserving Activation Scaling (GPAS), a residual-stream plugin that combats the exponential growth of activation variance — the root cause of residual-shortcut dominance over sub-layer outputs and the degraded learning of deeper layers, most acute in the Pre-LayerNorm (Pre-LN) Transformer — by scaling down intermediate activations while preserving their gradients. GPAS avoids the gradient-vanishing pathology of naive downscaling and yields consistent gains at model sizes from 71M to 1B parameters across Pre-LN, Post-LN, DeepNorm, Sandwich-LN, Mix-LN, and LayerNorm-Scaling. Building on this line of work, we then introduce Progressive Residual Warmup (ProRes), a depth-aware optimization curriculum that gradually warms up each layer’s residual contribution from zero to one, with deeper layers waiting longer than shallow ones; this “early layers learn first” inductive bias stabilizes optimization across initialization and normalization schemes and produces a distinctive trajectory that yields faster convergence, stronger generalization, and better downstream performance.

The second part of the thesis turns to reasoning. We present Double-Checker, a self-critical fine-tuning framework that teaches slow-thinking, long chain-of-thought LLMs to explicitly critique their own intermediate solutions and iteratively refine them until the model judges its answer correct under its own critique. With only 1,730 carefully curated self-critical instances, Double-Checker substantially improves the reasoning accuracy and self-correction reliability of strong slow-thinking baselines.

Together, these contributions improve modern LLMs along two complementary axes: macro-architectural interventions that improve training stability and depth utilization during pretraining, and lightweight self-critical fine-tuning that elicits more reliable reflective reasoning after pretraining.

Date : 14 July 2026, Tuesday

Time : 10:30 am

Venue : Rm 4503 (Lifts 25/26), 4/F Academic Building, HKUST

Thesis Examination Committee:

Chairman	:	Prof. Qi YING, ENVR/HKUST
Thesis Supervisor	:	Prof. Can YANG, MATH/HKUST
Member	:	Prof. Yang XIANG, MATH/HKUST
Member	:	Prof. Dong XIA, MATH/HKUST
Member	:	Prof. Hao CHEN, ECE/HKUST
External Examiner	:	Prof. Yixiang FANG, School of Data Science and School of Artificial Intelligence/ The Chinese University of Hong Kong, Shenzhen (via online mode) (Open to all faculty and students)

The student's thesis is now being displayed on the reception counter in the General Administration Office (Room 3461).