

韩槿一

☎ 15224975562 ✉ jinyihan099@gmail.com 🏠 <https://jinyihan99.github.io>

📖 研究兴趣

主要聚焦在大模型深度思考/推理能力增强以及大模型元认知能力增强

深度思考/推理增强: 主要研究如何通过算法设计、训练范式创新和推理优化, 来赋予大模型更强的逻辑推理、多步思考和动态规划等能力, 使其在处理复杂任务时能够进行更具条理、深度和高效的推理过程

大模型的元认知能力增强: 包括对其自身知识边界的识别、自我纠错机制的构建、自我完善策略的制定以及基于生成反馈的持续进化能力。这些类人的元认知能力有利于推动大模型从被动响应向主动认知转变, 提升其在复杂任务中的表现以及可靠性

🔍 科研经历

增强模型自我纠错能力

- 提出一种内在自我纠错 (Intrinsic Self-Correction, ISC) 方法, 通过监督微调将纠错机制内化为模型能力, 使模型能够在生成过程中自主完成答案验证与修正, 从而摆脱对复杂提示工程的依赖
- 设计自我纠错数据构建流程, 并引入部分答案掩码 (Partial Answer Masking, PAM) 技术进行训练, 以减少错误答案对初始生成准确率的干扰, 使参数规模在 6B ~ 13B 的开源模型也具备内在自我纠错能力
- 实验结果表明, 采用 ISC 方法后, ChatGLM 在 OpenBookQA 数据集上的准确率提升 5.6%

增强模型自知之明能力

- 提出细粒度置信度校验方法(FineCE),能够在生成过程中预测模型生成正确答案的概率
- 通过蒙特卡洛采样构建置信度训练数据, 并用于微调模型, 从而赋予其表达置信度的能力。进一步提出三种策略以识别执行置信度估计的最佳位置, 并在推理阶段引入后向置信度融合机制, 实现对当前生成内容的全局置信度建模
- 在数学推理、常识问答等多个任务上进行评估, FineCE均优于现有的置信度估计方法, 且在生成过程约三分之一处, 即可提前预测最终输出的正确概率, 有效提升决策效率与可靠性

增强模型主动完善能力

- 提出一种结合强化学习 (GRPO) 的训练框架, 赋予模型主动完善能力(PASR), 使其在生成过程中能够自主判断是否、何时以及如何执行完善操作
- 设计包含格式、准确率与完善效果在内的三种奖励信号, 并针对不同类型的奖励制定特定引导策略, 促进模型在生成过程中进行高质量的主动完善
- 在 Qwen3-8B 模型上应用 PASR 方法后, 在常识推理与复杂推理等多项任务中实现显著提升: 平均 token 消耗减少 45%, 同时准确率提升 8.2%

基于深度思考的大模型自主工具学习方法

- 利用大模型的深度思考能力, 赋予其快速掌握并自主使用外部工具的能力, 使其在生成过程中能够根据任务需求主动选择工具以辅助推理
- 在计算器与 Python 编程工具上开展初步实验, 使模型在数学推理中具备自动调用与集成工具的能力, 从而有效提升复杂问题的解决效率与准确性
- 在超过 9 位数的复杂数学问题上, 相较于基准模型 (Qwen2.5-7B-Instruct), 实现 16.7% 的准确率提升

📄 论文列表

[*: 导师第一作者, 学生为第二作者; †: 通讯作者; 横线加粗: 本人(曾用名: 韩海霞)]

- Small Language Model Can Self-correct. Haixia Han, Jiaqing Liang, Jie Shi, Qianyu He, Yanghua Xiao. 2024, AAAI.

- Think Thrice Before You Act: Progressive Thought Refinement in Large Language Models. Chengyu Du, Jinyi Han, Yizhou Ying, Aili Chen, Qianyu He, Haokun Zhao, Haoran Guo, Sirui Xia, Jiaqing Liang, Zulong Chen, Liangyue Li, Yanghua Xiao. 2025, ICLR.
- CDS: Data Synthesis Method Guided by Cognitive Diagnosis Theory. Haokun Zhao, Jinyi Han, Jiaqing Liang, Xiaojun Meng, Jiansheng Wei, Yanghua Xiao. 2025, ACL.
- Mind the Generation Process: Fine-Grained Confidence Estimation During LLM Generation. Jinyi Han, tingyun li, Shisong Chen, Jie Shi, Xinyi Wang, Guanglei Yue, Jiaqing Liang, Xin Alex Lin, liqian wen, Zulong Chen, Yanghua Xiao. 2025, EMNLP (Under review).
- A Stitch in Time Saves Nine: Proactive Self-Refinement for Language Models. Jinyi Han, Xinyi Wang, Haiquan Zhao, tingyun li, Zishang Jiang, Sihang Jiang, Jiaqing Liang, Xin Alex Lin, Weikang Zhou, Zeye Sun, Fei Yu, Yanghua Xiao. 2025, NIPS (Under review).
- SED: Self-evaluation decoding enhances large language models for better generation. Ziqin Luo, Haixia Han, Haokun Zhao, Guochao Jiang, Chengyu Du, Tingyun Li, Jiaqing Liang, Deqing Yang, Yanghua Xiao. 2024.
- CEM: A Data-Efficient Method for Large Language Models to Continue Evolving From Mistakes. Haokun Zhao, Jinyi Han, Jie Shi, Chengyu Du, Jiaqing Liang and Yanghua Xiao. 2025 (Under review).
- Dynamic Clustering Based Contextual Combinatorial Multi-armed Bandit for Online Recommendation. Cairong Yan, Haixia Han*, Yanting Zhang, Dandan Zhu, Yongquan Wan. 2022, Knowledge-Based Systems.
- CoCoB: Adaptive Collaborative Combinatorial Bandits for Online Recommendation. Cairong Yan, Jinyi Han[†], Jin Ju, Yanting Zhang, Zijian Wang, Xuan Shao. 2025, Dasfaa.
- Thompson Sampling with Time-Varying Reward for Contextual Bandits. Cairong Yan, Hualu Xu, Haixia Han, Yanting Zhang, Zijian Wang. 2023, Dasfaa.

 项目经历

基于GRPO的大模型深度思考能力复现 [开源项目](#) 2025.03 – 2025.04
主要贡献人 (1.1K Star) [开源地址](#)

- 基于 GRPO 方法高效复现 R1-zero 的自发反思能力，整体实现仅约 200 行代码，且仅依赖基础深度学习库，具有良好的可复现性与轻量性
- 通过5步优化，使Qwen2.5-3B准确率稳定在60%以上，峰值可达约70%；30步后格式遵循能力接近100%

大模型自主工具调用 [开源项目](#) 2025.04 – 至 今
主要贡献人 [开源项目地址](#)

- 利用大模型的深度思考能力，开源实现其自主调用工具的机制
- 设计并训练支持“边想边干”和“专业分工”两种工具使用模式的系统，使模型能灵活、自主地多次调用工具，并在生成过程中自然融合调用结果以优化后续推理与决策

基于认知诊断的大模型持续进化和知识更新研究 华为 2024.09 – 2025.02
项目技术核心人员

- 现有大模型持续训练的方法主要通过“评估-合成-训练”的流程，但是由于评估粒度较粗难以衡量模型真实的能力。引入知识点组件来刻画模型细粒度能力，以精确定位模型错误的根因，并据此构建有针对性的合成指令数据集来提升模型性能
- 在数学推理、代码补全与综合学科测试等多个基准上，取得优于现有合成数据训练方法的性能表现

基于大模型的代码智能化技术 华为 2023.11 – 2024.11
项目模块核心人员

- 大模型应具备类似人类程序员的“检索—复用—改写”能力，以更高效地完成代码生成任务。首先通过合成低资源数据，提升模型在低资源场景下的代码生成能力；随后结合智能检索，获取与用户需求相关的高质量代码片段，以增强模型的上下文理解与复用能力；最后通过风格一致性微调，引导模型输出符合特定编码规

范的代码，从而实现高质量、可控性强的代码生成过程

- 从变量命名、API 调用和代码结构等多个维度定义代码风格，并构建风格判别器，用于识别生成代码与目标风格之间的一致性。当检测到风格不一致时，触发代码风格一致性迁移，自动调整生成内容的风格

大模型中的数据科学 华为

项目模块核心人员

2024.12 – 至 今

- 研究复杂专业思维数据的合成与优化方法
- 定义一系列思维算子，并在模型生成过程的关键位置插入相应算子，有效提升模型回答的准确率

🎓 教育经历

华东师范大学 (复旦大学联合培养) 学术博士 计算机科学与技术学院	2023.09 – 至 今
东华大学 学硕硕士 计算机科学与技术学院	2020.09 – 2023.03
河南财经政法大学 本科 计算机科学与技术学院	2016.09 – 2020.06

♥️ 获奖情况

上海市优秀毕业生	2023.03
东华大学优秀硕士毕业论文	2023.03
学业奖学金一等	2020.10-2022.10
河南省优秀学位论文	2021.07