



COMP RESEARCH STUDENT SEMINAR

Date : 26 Nov 2025 (Wed)
Time : 10:45 am - 12:15 pm
Venue : PQ305 (Face-to-face)

Cracking the Code of Juxtaposition: Can AI Models Understand the Humorous Contradictions

Abstract

Recent advancements in large multimodal language models have demonstrated remarkable proficiency across a wide range of tasks. Yet, these models still struggle with understanding the nuances of human humor through juxtaposition, particularly when it involves nonlinear narratives that underpin many jokes and humor cues. This paper investigates this challenge by focusing on comics with contradictory narratives, where each comic consists of two panels that create a humorous contradiction. We introduce the YesBut benchmark, which comprises tasks of varying difficulty aimed at assessing AI's capabilities in recognizing and interpreting these comics, ranging from literal content comprehension to deep narrative reasoning. Through extensive experimentation and analysis of recent commercial or open-sourced large (vision) language models, we assess their capability to comprehend the complex interplay of the narrative humor inherent in these comics. Our results show that even state-of-the-art models still lag behind human performance on this task. Our findings offer insights into the current limitations and potential improvements for AI in understanding human creative expressions.



Mr HU Zhe

PhD candidate
Department of Computing

About the Speaker

Mr HU Zhe is a PhD candidate in the Department of Computing at The Hong Kong Polytechnic University. His research focuses on human-centered reasoning and decision-making in large models. He also works on developing methods for text generation and summarization.

Merlin's Whisper: Enabling Efficient Reasoning in LLMs via Black-box Adversarial Prompting



Mr XIA Heming

PhD candidate
Department of Computing

About the Speaker

Mr XIA Heming received his master's degree from the MOE Key Lab of Computational Linguistics at Peking University. Prior to that, he completed his bachelor's degree from the School of Physics at Peking University. He is now a PhD student in the NLP Group at The Hong Kong Polytechnic University, supervised by Prof. LI Wenjie. His research interest focuses on efficient and effective NLP, with the goal of making LLMs faster, more scalable, and broadly applicable. Specifically, his work centers on speculative decoding and efficient reasoning.

Abstract

Large reasoning models (LRMs) have demonstrated remarkable proficiency in tackling complex reasoning tasks through step-by-step thinking. However, such a lengthy reasoning process incurs substantial computational and latency overheads, hindering the practical deployment of these models. In this work, we present a new perspective on mitigating overthinking in LRMs via black-box adversarial prompting. By treating both open-source LRMs and closed-source APIs as black-box communicators, we investigate how to elicit concise responses without sacrificing accuracy. We introduce AdvPrompt, an iterative refinement framework that generates high-quality adversarial prompts from diverse perspectives. Experiments across multiple benchmarks demonstrate that AdvPrompt consistently reduces token usage while preserving performance. Notably, AdvPrompt achieves a 3x reduction in average response length on simple GSM8K questions for the Qwen3 model series, and delivers an average ~40% token reduction across four benchmarks. For closed-source APIs, AdvPrompt reduces token usage on MATH-500 by 35% for Claude-3.7 and 47% for Gemini-2.5. Further analysis reveals the generalizability of AdvPrompt across various model scales and families, underscoring the potential of black-box prompting as a practical and effective strategy for enhancing LRM efficiency.



COMP RESEARCH STUDENT SEMINAR

Date : 26 November 2025 (Wed)

Time : 10:45 am - 12:15 pm

Venue : PQ305 (Face-to-face)

How do large language models understand graph patterns? a benchmark for graph pattern comprehension

Abstract

Benchmarking the capabilities and limitations of large language models (LLMs) in graph-related tasks is becoming an increasingly popular and crucial area of research. Recent studies have shown that LLMs exhibit a preliminary ability to understand graph structures and node features. However, the potential of LLMs in graph pattern mining remains largely unexplored. This is a key component in fields such as computational chemistry, biology, and social network analysis. To bridge this gap, this work introduces a comprehensive benchmark to assess LLMs' capabilities in graph pattern tasks. We have developed a benchmark that evaluates whether LLMs can understand graph patterns based on either terminological or topological descriptions. Additionally, our benchmark tests the LLMs' capacity to autonomously discover graph patterns from data. The benchmark encompasses both synthetic and real datasets, and a variety of models, with a total of 11 tasks and 7 models. Our experimental framework is designed for easy expansion to accommodate new models and datasets. Our findings reveal that: (1) LLMs have preliminary abilities to understand graph patterns, with O1-mini outperforming in the majority of tasks; (2) Formatting input data to align with the knowledge acquired during pretraining can enhance performance; (3) The strategies employed by LLMs may differ from those used in conventional algorithms.



Mr QU Haohao

PhD candidate

Department of Computing

About the Speaker

QU Haohao is currently a PhD candidate of the Department of Computing (COMP), The Hong Kong Polytechnic University (PolyU), under the supervision of Dr. Wenqi Fan and Prof. Qing Li. Before joining the PolyU, he received both a Master's degree and a Bachelor's degree from Sun Yat-sen University in 2022 and 2019, respectively. His research interest covers Recommender Systems and Large Language Models. He has authored and co-authored around 20 innovative papers, most of them have been published in top journals and conferences, such as IEEE TKDE, ACM TOIS, ICLR, and KDD. For more information, please visit <https://quhaoh233.github.io/page>.