

Large Language Models in Computational Psycholinguistics

Large language models (LLMs) seem to understand and learn natural language like humans. However, despite their human-like performance, whether LLMs can serve as cognitive models of human language understanding and learning has been intensely debated. In this talk, I will review both promises and limitations of LLMs in computational psycholinguistics. Specifically, failures of "vanilla" LLMs demonstrate insufficiency (upper bound) of statistical learning, while successes of linguistically-informed LLMs suggest necessity (lower bound) of cognitive biases like syntax or memory. Overall, unlike the previous literature where human psycholinguistics is leveraged for AI scientists to evaluate LLMs, LLMs can also be employed for linguists to understand human psycholinguistics.