

Postgraduate Seminar Series

Venue: Graduate Seminar Room

Date & Time: August 29, 2026 at 2:30 PM

Speaker Information

Md. Sakif Khan (Std No. 0421052099) is a part time M.Sc. student in the department of CSE, BUET. He completed his undergraduate studies from Islamic University of Technology (IUT) in 2021. His research interest lies in the field(s) of Algorithms, Graph Theory, and Large Language Models. He is currently doing his postgraduate thesis under the supervision of Dr. Sadia Sharmin. He will be speaking about his ongoing research in this talk.



Agentic Graph Reasoning: Autonomous Knowledge Graph Navigation for Fact Verification and Hallucination Mitigation in Large Language Models

Large language models answer factual questions fluently. But, they tend to sound just as fluent when they don't hold the facts. To fix this, researchers introduced Retrieval-Augmented Generation (RAG), which connects the LLM to an external data source. This grounds the model's answers in retrieved text. But, the retrieval happens once and up front, before any reasoning starts. So, a question that spans several facts has to be answered from whatever that first query brought back. Systems that navigate a knowledge graph do better by interleaving retrieval with reasoning. Even they, however, send out whatever their final generation call produces, and nothing checks the answer's claims before it reaches the user.

To close that gap we propose Agentic Graph Reasoning (AGR): a knowledge-graph question-answering framework. Here, an agent plans sub-objectives and walks the graph through a constrained, deterministic tool API. Its distinguishing component is a Structural Verification Layer. Before the answer is emitted, it splits the draft into atomic claims and checks each one against the triples the agent actually traversed. Claims it cannot ground are re-explored or withdrawn. So, AGR is more likely to hedge than to assert. And it returns every answer paired with the triples that support it.

We evaluate AGR on a Freebase-derived environment of 2.59 million entities and 8.31 million triples, against four baselines under one shared backbone and budget, so any difference is architectural rather than a matter of model capacity. AGR beats every baseline on both benchmarks (Hits@1 of 0.755 on WebQSP, 0.522 on ComplexWebQuestions), using half as many model calls as the nearest agentic baseline. It also never asserts an entity it can't ground, versus a 22.1% ungrounded rate for the parametric control. And on ComplexWebQuestions, it's the only system whose accuracy rises with hop count.