

PhantomWiki: On-Demand Datasets for Reasoning & Retrieval Evaluation



Takeaways

- PhantomWiki stress tests LLMs on multi-hop question-answering tasks, while ensuring:
1. LLMs cannot rely on factual knowledge from training
 2. Factually consistent questions and answers
 3. Fast and fully-automated dataset generation: 1M-sized dataset with 4-hop questions takes < 30 CPU hours!

Motivation

- LLMs must answer questions with up-to-date knowledge, but evaluation is a challenge:
- Static datasets are prone to **data leakage**
 - Disentangling a model’s **internal knowledge**, **reasoning**, and **retrieval** capabilities is difficult with datasets curated from public data (e.g., Wikipedia).

Many Evaluation Scenarios

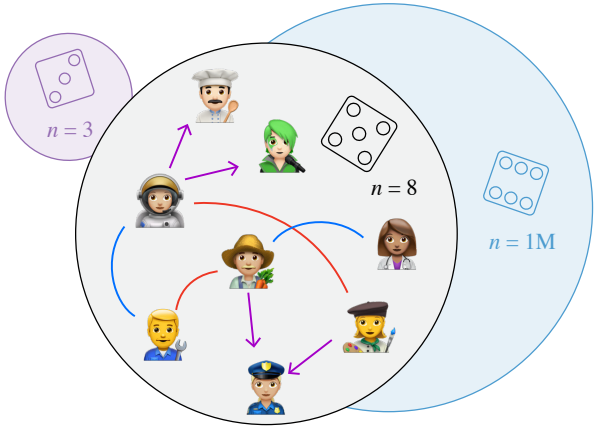
1. **In-Context**: document corpus fits in LLM context window [NIAH, LongBench, RULER, ∞ -Bench, HELMET]
2. **RAG**: relevant documents are retrieved using an external retriever [MultiHop-RAG, BRIGHT, ARES]
3. **Agentic**: LLM uses external tools to obtain relevant context [ToolQA]



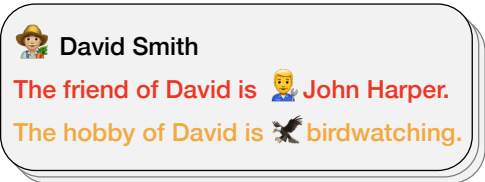
Prolog + Context-Free Grammar

Fact = predicate + constants	Query = predicate + variables
The mother of Alice is Charlotte. <code>mother("Alice", "Charlotte").</code>	Who is the mother of Alice? <code>mother("Alice", X).</code>
$S \rightarrow \text{Who is } R \text{ ?}$ $R \rightarrow \text{the } \langle \text{relation} \rangle \text{ of } R'$ $R' \rightarrow R \mid \langle \text{name} \rangle$	

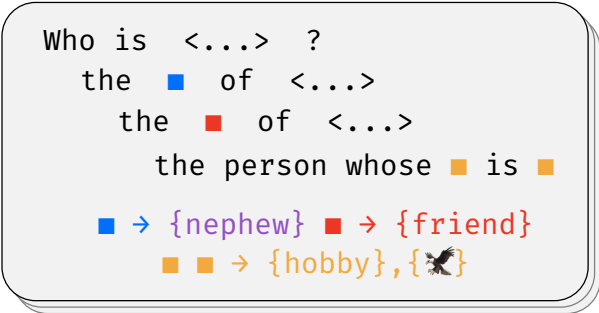
PhantomWiki Pipeline



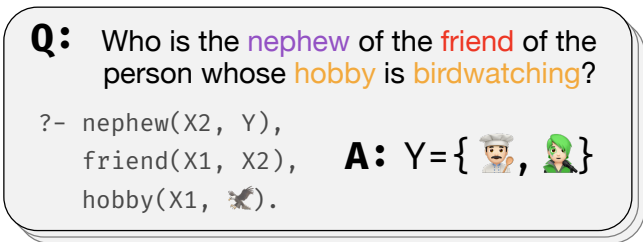
(1) Generate a random universe of size n



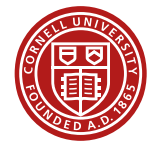
(2) Generate document corpus for the universe



(3) Generate questions using a context-free grammar



(4) Use Prolog to deduce ground-truth answers

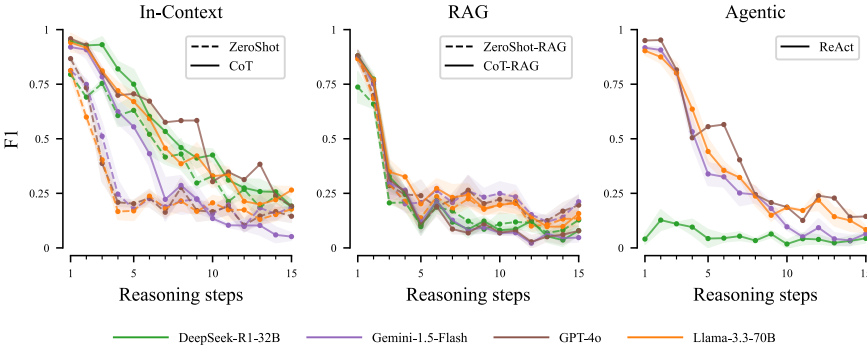


arxiv.org/abs/2502.20377



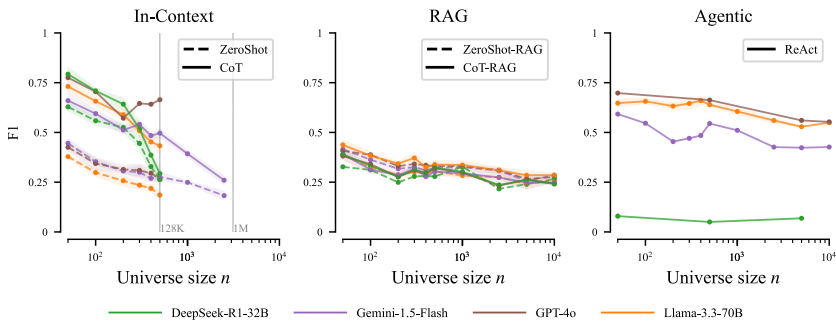
github.com/kilian-group/phantom-wiki

PhantomWiki for reasoning evaluation

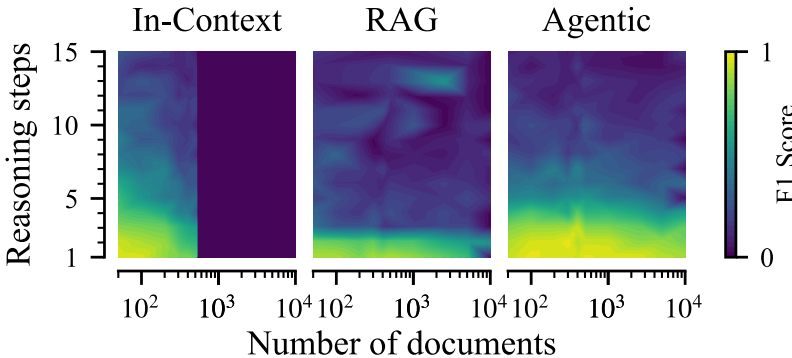


Reasoning steps = CFG depth + relation difficulty

PhantomWiki for retrieval evaluation

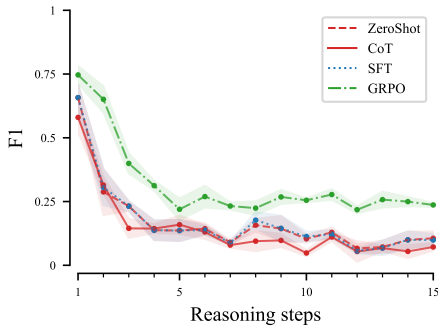


Retrieval & Reasoning (w/ Llama 3.3-70B)

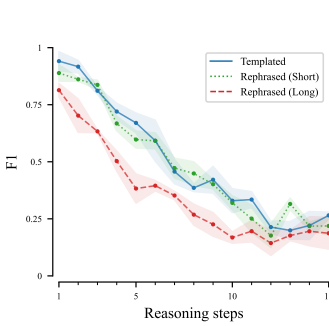


Does training improve reasoning ability?

Does article style affect reasoning ability?



LoRA on Qwen2.5-3B



Templated vs LLM-written