

MultiHoax: A Dataset of Multi-hop False-premise questions

Anonymous ACL submission

Abstract

As Large Language Models are increasingly deployed in high-stakes domains, their ability to detect false assumptions and reason critically is crucial for ensuring reliable outputs. False-premise questions (FPQs) serve as an important evaluation method by exposing cases where flawed assumptions lead to incorrect responses. While existing benchmarks focus on single-hop FPQs, real-world reasoning often requires multi-hop inference, where models must verify consistency across multiple reasoning steps rather than relying on surface-level cues. To address this gap, we introduce MultiHoax, a benchmark for evaluating LLMs’ ability to handle false premises in complex, multi-step reasoning tasks. Our dataset spans seven countries and ten diverse knowledge categories, using Wikipedia as the primary knowledge source to enable cross-regional factual reasoning. Experiments reveal that state-of-the-art LLMs struggle to detect false premises across different countries, knowledge categories, and multi-hop reasoning types, highlighting the need for improved false premise detection and more robust multi-hop reasoning capabilities in LLMs.

1 Introduction

In recent years, the evolution of large language models (LLMs) has demonstrated their immense potential across a wide range of natural language processing tasks. However, despite their impressive successes in many domains, their ability to recognize and properly handle false premises in reasoning tasks remains a significant challenge (Hu et al., 2023; Yu et al., 2023). When engaging with false premises, an LLM should be able to identify and reject flawed assumptions rather than proceed as if they were valid.

False premises can appear in various forms, such as misleading statements, logically inconsistent claims, or factually incorrect contextual narratives. (Leite et al., 2023; Zhuang et al., 2023; Ghosh

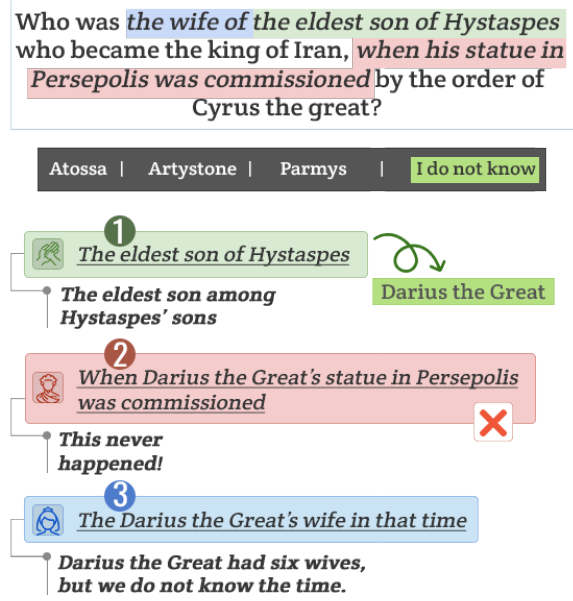


Figure 1: A sample MHFPQ from the history category related to Iran.

et al., 2024; Galitsky et al., 2024; Chen and Shu, 2023; Yamin et al., 2024; Zhang et al., 2024b; Csöngé, 2015) A common approach to evaluating an LLM’s ability to handle false premises is through question-answering (QA) tasks, where a question implicitly contains a misleading or incorrect assumption. (Yu et al., 2023; Hu et al., 2023; Kim et al., 2023) For example, a well-performing LLM should recognize that “Which year did Albert Einstein win the Nobel Prize in Chemistry?” contains a false premise—Einstein never won a Nobel Prize in Chemistry—and responds appropriately, rather than attempting to provide an incorrect or misleading answer.

Existing benchmarks focus on single-hop false premise detection, where the incorrect assumption can be identified within a single reasoning step (Yu et al., 2023; Hu et al., 2023; Kim et al., 2023). However, multi-hop reasoning presents a greater

challenge, as it requires models to connect multiple pieces of information across multiple inference steps before arriving at an answer. (Mavi et al., 2022) Unlike single-hop questions, multi-hop questions (MHQs) require models to derive intermediate conclusions, making them a critical area of research in the evaluation and improvement of LLM reasoning abilities (Mavi et al., 2022; Chen et al., 2024; Yang et al., 2024b; Tang and Yang, 2024; Chakraborty, 2024). The complexity of MHQs arises from the necessity of bridging multiple facts, understanding implicit dependencies between reasoning steps, and maintaining logical consistency throughout the inference process (Mavi et al., 2022).

While MHQs are already challenging, the problem becomes even more difficult when false premises are embedded within the reasoning chain, requiring models not only to answer questions correctly but also to detect and reject incorrect assumptions at intermediate steps. To evaluate LLMs’ ability to handle this challenge, we introduce a novel benchmark of Multi-Hop False-Premise Questions (MHFPQs), called **MultiHoax**, that combines the difficulty of false premise questions and multi-hop question answering. MHFPQs are designed to test whether LLMs can detect false premises that appear at one or more reasoning steps rather than simply providing an answer based on a flawed assumption.

Figure 1 shows an example of an MHFPQ from the history category related to Iran. Answering this question requires a structured reasoning process: first, the model must correctly identify that the eldest son of Hystaspes was Darius the Great. In the second reasoning step, the model must evaluate whether Darius the Great’s statue was commissioned in Persepolis by the order of Cyrus the Great, which is a false claim about an event never occurred. Since the time frame of the question depends on this fabricated event, the entire question becomes unanswerable. The third reasoning step then asks for Darius the Great’s wife during that supposed time period, but as the previous step is based on incorrect information, it is impossible to determine a valid answer.

To ensure realism and difficulty, we carefully designed the set of possible answer choices in the MHFPQ dataset. Each question includes believable distractors that align with historical facts, making the task particularly challenging for LLMs. For instance, in the example from Figure 1, the distractors were drawn from the actual wives of Darius

the Great, ensuring that incorrect answers remain historically plausible while still being contextually invalid. Without detecting the false premise, LLMs may be misled into selecting one of the plausible but incorrect answers, emphasizing the necessity for models to engage in deep contextual reasoning rather than relying on surface-level fact retrieval.

The dataset spans 7 countries and 10 distinct categories, ensuring broad diversity across historical, cultural, and geopolitical contexts. To generate the MHFPQs, we extracted relevant information from Wikipedia pages with the assistance of the Claude model¹. Each question is paired with three closely related distractors and a single “I do not know” option, which serves as the only valid response in the presence of a false premise. By structuring the benchmark in this way, we aim to assess LLMs’ ability to identify and reject incorrect assumptions embedded within multi-hop reasoning chains, rather than merely selecting the most superficially plausible response.

We conduct a comprehensive evaluation of six open-source and closed-source models using our MultiHoax² dataset of 700 carefully reviewed questions. Our results reveal suboptimal performance across all models, categories, and countries, indicating that the models struggle to comprehend knowledge at the country level and fail to detect falsehoods in country-specific topics. The dataset would be open-source and publicly available.

2 Related Work

False Information and False Premises. False information is a broad and multifaceted issue, encompassing explicitly misleading claims from social networks (Ma et al., 2022; Rode-Hasinger et al., 2022), blogs (Okazaki et al., 2013), news sources (Long et al., 2017; Qiao et al., 2020; Huang et al., 2023), and online forums (Yu et al., 2023). However, false premises are a distinct challenge since they are not always deliberate misinformation but rather incorrect assumptions that lead to flawed logical conclusions. These premises span various domains, including historical events (Gabba, 1981; Smith, 2001), religion (Kutlu, 2022), and sports (Dimov, 2021), making them particularly complex to detect.

False-Premise Questions. FPQs assess a model’s ability to reason over implicit false

¹<https://www.anthropic.com/>

²<https://anonymous.4open.science/r/MHFPQ-BF29>

assumptions about properties, actions, scope, existence, events, logic, or causality (Hu et al., 2023; Yu et al., 2023). For example, “How many eyes does the sun have?” wrongly assumes the sun has eyes (Hu et al., 2023). Since FPQs require models to identify and reject flawed premises rather than simply answering the question, LLMs often fail to recognize implicit falsehoods, leading to misleading or nonsensical responses (Yu et al., 2023). This inability to reject false premises poses risks for misinformation propagation and model alignment with human expectations.

Early studies on FPQs have focused on obvious falsehoods that humans can easily detect. For instance, Hu et al. (2023) introduced a dataset of simple FPQs, such as “What is the most common color of human’s wings?”. While earlier LLMs struggled with these, recent versions of the LLMs handle them easily, raising the need to evaluate models on more subtle false premises. Similarly, Kim et al. (2023) proposed a dataset evaluating models on questions containing false or unverifiable assumptions. However, their dataset includes unverifiable claims—statements that may become true over time. For example, “When is Steven Universe Season 5 coming to Hulu?” assumes the event has not yet occurred in order to be false, but this assumption could later become valid. This distinction makes their dataset less suitable for assessing persistent false premises. Another line of research has focused on false-presupposition questions extracted from online forums. Yu et al. (2023) introduced a dataset of FPQs sourced from Reddit, such as “How exactly is current stored in power plants?”—a misleading assumption since electric current is not stored. However, their dataset is primarily limited to scientific and technical topics, lacking the broad factual diversity required for a general evaluation of false premises.

Beyond FPQs, unanswerable questions have also been explored in related research. Zhao et al. (2024) focus on document-grounded unanswerable questions, where a question lacks supporting information within a given document. Their contribution is in evaluating models’ ability to reformulate unanswerable questions into answerable ones, rather than directly assessing how LLMs handle false premises in an open-domain setting. Additionally, Lin et al. (2022) study truthful question answering, focusing on questions that elicit misconceptions from humans. Relatedly, Zhang et al. (2024a) and Peng et al. (2024) examine how LLMs

handle questions beyond their knowledge, where the correct response should be “I do not know”. These studies, however, focus on uncertain or conflicting information rather than inherently false premises.

Multi-Hop False-Premise Reasoning. Most FPQ benchmarks focus on single-hop false premises, yet real-world misinformation often involves multi-step reasoning, making it significantly harder to detect. Multi-hop reasoning questions challenge LLMs by requiring them to connect multiple facts across inference steps. While extensive research has explored general MHQ understanding in LLMs (Yang et al., 2018; Rajabzadeh et al., 2023; Park et al., 2023; Chen et al., 2024), MHQs embedding false premises remain largely unstudied. Existing benchmarks assess multi-hop factual reasoning but do not evaluate whether LLMs can detect and reject false premises within reasoning chains. This gap is critical, as misinformation is rarely an isolated error—false premises are often interwoven across reasoning steps, making them subtle, plausible, and difficult to refute. Evaluating how LLMs handle multi-hop false premises is essential for enhancing their robustness against misinformation and ensuring logical consistency in complex reasoning tasks.

Daswani et al. (2024) attempt to address this with adversarial multi-hop false premise questions, modifying HotpotQA (Yang et al., 2018). Their approach replaces the title of a supporting document with a similar but unrelated distractor, selected based on shared Wikipedia categories. For example, Roger O. Egeberg and Steven K. Galson both fall under American public health doctors, making them interchangeable under this method. However, this technique relies on structured entity swaps rather than embedding implicit falsehoods, making it unclear whether a question contains a truly false premise or is merely unverifiable. Additionally, this approach can lead to unnatural question phrasing, limiting its applicability in assessing real-world false premise detection.

Overall, LLMs struggle with imperfect information, conflicting evidence, and questions beyond their knowledge scope (Lin et al., 2022; Zhang et al., 2024a; Peng et al., 2024; Kazemi et al., 2024; Payandeh et al., 2024; Shaier et al., 2024; Park and Lee, 2024; Longpre et al., 2021; Chen et al., 2022). They are also susceptible to distraction by irrelevant details, including false premises, often

leading to incorrect, misleading, or fabricated responses (Park and Lee, 2024; Yu et al., 2023; Kim et al., 2023; Hu et al., 2023; Asai and Choi, 2021). These findings underscore a critical gap in evaluating multi-step false premise reasoning, as prior benchmarks focus on single-hop FPQs and adversarial perturbations, failing to capture the complexity of false premises embedded within structured factual reasoning. To address this, we introduce **MultiHoax**, a new benchmark of multi-hop FPQs, designed to assess LLMs’ ability to detect and reject false premises within reasoning chains. By including implicit false factual claims across diverse country-related topics, our benchmark provides a more realistic and challenging evaluation of LLM robustness in handling misinformation at scale.

Multi-Regional and Multi-Cultural Resources. Our dataset aligns with multi-cultural and multi-regional NLP research, which examines how LLMs process knowledge across different geographic and cultural contexts. While much prior work has focused on subjective tasks like norms and values (Ziems et al., 2023; Cheng et al., 2023; Fung et al., 2024; Shi et al., 2024; Han et al., 2023), recent studies show that multicultural NLP also includes objective knowledge, such as region-specific facts, historical events, and socio-political contexts (Koto et al., 2024; Li et al., 2024; Koto et al., 2023; Son et al., 2024; Kim et al., 2024). This highlights the need to evaluate how LLMs handle factual reasoning across diverse regions, particularly in multi-hop tasks involving false premises.

Despite growing interest in multi-regional fact verification, most work focuses on binary misinformation detection rather than multi-hop false premise reasoning (Thaher et al., 2021; Sabzali et al., 2022; Sheng et al., 2022). A global benchmark for multi-hop false premise detection remains missing, even though factual inaccuracies vary significantly across regions. To address this, we introduce a multi-hop false premise benchmark spanning multiple factual domains across seven countries with varying NLP resource availability. Unlike previous datasets, which focus on single-domain errors or adversarial perturbations, our MultiHoax benchmark evaluates LLMs’ ability to reason over embedded falsehoods in diverse cultural and geographic contexts, providing a more realistic and comprehensive assessment of global factual reasoning.

3 MultiHoax dataset

This section describes the MHFPQ dataset, designed to evaluate multi-hop false premise reasoning across diverse countries, categories, and question types. The dataset structure enables a systematic analysis of how LLMs handle false premises within complex reasoning tasks.

3.1 Dataset Framework

Each MHFPQ instance consists of the following components:

- **Question and Answer Choices:** Each question contains at least a false premise and is paired with four answer choices, three plausible distractors, and one “I do not know”. The latter is positioned randomly to prevent models from exploiting positional bias.
- **False Premise Explanation:** A brief description clarifies why the assumption in the question is incorrect.
- **Country and Category:** The dataset spans seven countries (China, France, Germany, Italy, the United States, the United Kingdom, and Iran) and is categorized into ten knowledge domains, including food, sports, geography, education, history, entertainment, religion, science & technology, arts & literature, and holidays & leisure.
- **Types of False Premises and Answer Types:** Inspired by (Hu et al., 2023), each question has a field that shows how the false premise is introduced inside the question. The possible types of the false premise are Property, Event, Entity, and Scope. For example, “Which award did Venkatraman Ramakrishnan receive first: the Shaw Prize in Life Science and Medicine or the Lasker Award?” is from the event type as it implies an event never happened because he never won the mentioned awards. Inspired by (Hu et al., 2023), we annotate each question based on how it includes a false premise. The primary types of false premises include Property, Event, Entity, Scope, and Index. For instance, the question “Which award did Venkatraman Ramakrishnan receive first: the Shaw Prize in Life Science and Medicine or the Lasker Award?” falls under the Event type, as it falsely implies an event, that is winning these awards by him, while in reality, he did not.³

³Table 13 (in the appendix) presents the distribution of

Furthermore, following [Yang et al. \(2018\)](#), we classify answer types such as person, location, event, number, and common nouns. The answer type shows what the question is asking for.⁴

- **Multi-Hop Reasoning Type:** Following [Mavi et al. \(2022\)](#), we categorize why a question requires multi-step inference into five types: (1) Named Entity Reasoning: The question requires connecting two facts through an intermediate entity that links them logically, (2) Temporal Reasoning: An intermediate step involves identifying a specific time reference to answer the question correctly, (3) Geographical Reasoning: The reasoning process depends on understanding locations, spatial relationships, or geographic entities, (4) Intersection Reasoning: The answer is determined by an entity that satisfies multiple overlapping conditions, and (5) Comparison Reasoning: The question requires comparing attributes, facts, or values across multiple entities to arrive at the correct conclusion.⁵
- **Wikipedia Grounding:** Each question is linked to a relevant Wikipedia page for factual grounding.

3.2 Wikipedia Document Collection

We developed a pipeline to extract Wikipedia pages relevant to each country and category using ChatGPT-4o’s search tool. To select the set of pages, each page was evaluated based on three criteria: relevance to the category, association with the specified country, and existence. “Existence” ensures that the provided link leads to an actual Wikipedia page rather than just an important but undocumented topic. If a page was missing, the model was prompted to complete the list by suggesting alternatives. Ultimately, we collected 15 Wikipedia pages per country and category and extracted their content.

3.3 Question Generation

After collecting relevant documents, we developed a structured process for generating MHFPQs. First, we instructed Claude 3.5 Haiku to extract 15 facts

false premise types in the dataset, along with definitions of each type.

⁴Table 15 (in the appendix) provides a detailed breakdown of answer option types across the dataset, including an example for each type.

⁵Table 14 (in the appendix) presents the distribution of each multi-hop type in the dataset.

per document using the prompt in Appendix 6. Based on these facts, we then prompted the model (Appendix 6) to generate MHFPQs. Due to fundamental differences between Bridge Entity-based MHQs (named, geographical, and temporal entities) and other types (intersection and comparison), we used a separate prompt for each category. We requested three questions from the former and two from the latter but did not enforce specific subtypes (e.g., one intersection, one comparison), as not all documents contained relevant facts. Additionally, we avoided first generating MHQs and then falsifying them, as this could limit the variety of false information types ([Daswani et al., 2024](#)) and might not ensure incorrectness.

3.4 Curated Selection and Expert Review

After generating the questions for each category and country, an expert reviewer evaluated the generated questions and selected 10 MHFPQs for each combination of category and country. The selection was guided by a structured, multi-step approach to ensure the quality, validity, and plausibility of the questions.

The first step in the selection process was to verify whether a question adhered to the MH structure. Ensuring diversity in MH question types was a key objective. If a question was not initially formatted as an MHQ but could be converted into one, it passed this initial filter. For instance, the question “What was the name of the Achaemenid ruler who appointed Cyrus as governor of the Median Empire?” is inherently a single-hop reasoning question. However, by introducing an additional reasoning step, such as asking about the ruler’s son, it could be transformed into an MHQ: “Who was the son of the Achaemenid ruler who appointed Cyrus as governor of the Median Empire?”.

Next, the question needed to contain at least one universally false piece of information—meaning the falsehood had to be global rather than merely incorrect within the context of the associated document. To ensure this criterion was met, only questions that demonstrably contained globally false statements were selected. This verification process followed a rigorous three-step approach. First, the reviewer traced each question back to its corresponding fact or set of facts. Second, these facts were cross-checked against the information found in the relevant document. Finally, the reviewer assessed whether the question was indeed incorrect relative to both the established facts and the docu-

Question	Description	MH Type	Answer Type	FP Type
Which Bundesliga team, Bayern Munich or Borussia Dortmund, has the stadium with the largest seating capacity among clubs that have won the FIFA World Cup?	None of them. Spotify Camp Nou and Estadio Santiago Bernabéu have the largest capacities among clubs that have won the FIFA Club World Cup, not the FIFA World Cup.	Comparison	Group/Org	Event
Who both served as a served as a volunteer in the Illinois Militia during the Black Hawk War seeing lots of combat during his tour, and also was among the assassinated presidents of the US?	Abraham Lincoln served as a volunteer in the Illinois Militia April 21, 1832 – July 10, 1832, during the Black Hawk War. However, Lincoln never saw combat during his tour. He was assassinated as well.	Intersection	Person	Property
In which country, besides the U.S. and Italy, was the 1972 American epic gangster film directed by Francis Ford Coppola filmed?	The Godfather (1972) was filmed exclusively in locations around New York City and Sicily, with no scenes shot in other countries.	Named Entity	Location	Event

Table 1: Examples of MHFPQs from MultiHoax.

ment. If the model did not generate enough false information, the reviewer modified the question by introducing falsehoods aligned with the dataset’s predefined types of misinformation. However, if adding such falsehoods was not feasible, the question was discarded.

The third step involves explaining why the question is incorrect. To achieve this, the model’s explanation is evaluated based on the previous step’s results. If it fails to address the false information, the reviewer provides a more detailed clarification.

The final step involved verifying the plausibility of the answer choices. The reviewer ensured all options were contextually relevant by referring to the corresponding section of the document. If the model’s initial options were unrelated—either due to a change in the question’s focus or the model’s poor performance in generating relevant choices—the reviewer replaced them with more appropriate alternatives.

An example of this procedure is illustrated in Figure 1. The initial question generated by Claude 3.5 Haiku was: *Who was the Prime Minister of Iran when Darius the Great’s statue in Persepolis was commissioned?*. The model also provided Reza Shah Pahlavi, Mohammad Mosaddegh, and Hassan Rouhani as possible answer choices. However, the false premise was too obvious. As the model itself explained: *Darius the Great ruled the Achaemenid Empire in the 6th-5th centuries BCE, long before the modern nation of Iran and its government structures existed. There was no Prime Minister of Iran during Darius’s reign.* Due to this significant time gap, the question lacked subtle misinformation.

To introduce a more nuanced falsehood, the reviewer modified the question to: *Who was the wife of the eldest son of Hystaspes who became the king*

of Iran, when his statue in Persepolis was commissioned by the order of Cyrus the Great?. Here, the false premise—that Cyrus the Great commissioned Darius the Great’s statue—is less obvious because the time gap is reduced, and the context remains within the same historical period of ancient Iran. Additionally, to ensure the question met MH criteria, it asked for the name of Darius the Great’s wife. To provide plausible answer choices, three names were selected from his six known wives: Atossa, Artystone, Parmys, Phratagune, Phaedyia, and the daughter of Gobryas.

3.5 Secondary False Information Verification

To ensure the accuracy of falsified content, a second round of review was conducted. In this phase, a second reviewer independently examined the description field of each question against the corresponding Wikipedia page to verify the presence of false information. The reviewer categorized each question into one of three outcomes: “There is false information”, “There is no false information”, and “I cannot tell based on the provided information”.⁶ Questions confirmed to contain false information were directly included in the final dataset. Those labeled with the second option were double-checked and falsehood was added where required, while those detected with the last option were both improved in terms of clarity and added with false information. In the latter two cases, after the necessary modifications, the reviewer provided feedback to ensure that the question contained clear false information. The dataset was finalized upon completion of this verification process.⁷ Table 1 shows

⁶Table 16 (in the appendix) shows the distribution of the labels across categories.

⁷Table 12 (Appendix) provides the annotation guidelines.

[QUESTION]: 1. [FIRST OPTION] | 2. [SECOND OPTION] | 3. [THIRD OPTION] | 4. [FOURTH OPTION]

Please only provide the answer index.

Why did you choose “I do not know”?

1. You were uncertain about the question and did not have enough knowledge to answer.
2. You thought the question was wrong and contained false information.

Table 2: Evaluation prompts for multi-hop false premise reasoning. The top prompt assesses multiple-choice QA, where models may reject false premises by selecting “I do not know”. The bottom prompt evaluates whether models correctly justify this choice.

examples from our final resource from different types and countries.

4 Evaluation Setup

The **MultiHoax** dataset forms the foundation of our experiments, evaluating multi-hop false premise reasoning through two tasks: multiple-choice QA and justification. In the first task, models were presented with a question and four answer choices, including “I do not know”, allowing them to reject the question if they identified a false premise. The prompt used for this step is shown in the first part of Table 2.

In the second task (justification), models that selected “I do not know” were further prompted to explain their choice, as shown in the second part of Table 2. This step differentiates between cases where the model lacked knowledge and those where it explicitly identified a false premise. Only responses that both select “I do not know” in the first task and correctly justify it as a false premise in the second task can be considered successful detections of false premises. This justification step enhances the reliability of our evaluation, ensuring that refusal to answer stems from false premise detection rather than general uncertainty.

Models We tested 6 different proprietary and open-source LLMs in our experiments. The models include Claude Sonnet 3.5⁸, Gemini-2.0-pro-exp (Team et al., 2023), GPT-4o (Hurst et al., 2024), Qwen2.5-7B-Instruct (Yang et al., 2024a), Llama-3.1-8B-Instruct (Meta et al., 2024), and Deepseek-llm-7b-chat (Bi et al., 2024). All experiments used a zero temperature setting to ensure deterministic responses, with all data collected in February 2025.

5 Results

Table 3 presents model accuracy on the two tasks. The first task evaluates whether models correctly

Model	Accuracy	
	1st Task	2nd Task
Claude Sonnet 3.5	0.46	0.23
Gemini-2.0-pro-exp	0.29	0.26
GPT-4o-2024-11-20	0.23	0.25
Qwen2.5-7B-Instruct	0.19	0.03
Llama-3.1-8B-Instruct	0.13	0.01
Deepseek-llm-7b-chat	0.05	0.06

Table 3: Model performance on MultiHoax, evaluating multi-hop false premise reasoning in (1) multiple-choice QA, where models may reject false premises with “I do not know”, and (2) justification, assessing if they recognize the false premise.

reject false premises by selecting “I do not know”. While Claude, which was used during the question-generation phase, demonstrates higher accuracy compared to other models, the models generally struggle to recognize falsehoods in the first task, as they are unable to refuse to answer. Furthermore, although all models show poor performance on this task, open-source models generally exhibit lower accuracy than proprietary ones.

The second task analyzes whether models can justify their “I do not know” responses by correctly identifying the false premise. Here, Gemini slightly outperforms the generator model (Claude), but overall accuracy remains low across all models, indicating persistent difficulty in recognizing false premises.

Table 4 presents model performance across ten categories. Science and technology shows the highest accuracy, while art and literature ranks the lowest, indicating varying levels of LLMs’ factual knowledge across domains. Notably, proprietary models consistently outperform open-source models, though performance remains low across all categories. Interestingly, Claude’s dominance is most pronounced in high-information domains like education and food, while other models exhibit irregular performance trends. These results highlight significant disparities in model strengths across dif-

⁸<https://www.anthropic.com/>

Category/Model	Claude	Gemini	GPT	Qwen	Llama	Deepseek	Avg
Science and Technology	0.458	0.329	0.286	0.215	0.143	0.430	0.310
Entertainment	0.429	0.286	0.172	0.129	0.115	0.000	0.189
Education	0.529	0.343	0.200	0.215	0.129	0.072	0.248
Art and Literature	0.458	0.258	0.172	0.072	0.086	0.043	0.182
Food	0.529	0.300	0.186	0.200	0.158	0.072	0.241
Religion	0.543	0.200	0.200	0.158	0.100	0.058	0.210
Sports	0.443	0.415	0.315	0.300	0.186	0.058	0.286
Holiday, Celebrations, and Leisure	0.580	0.286	0.315	0.229	0.100	0.072	0.264
Geography	0.343	0.243	0.215	0.186	0.172	0.058	0.203
History	0.343	0.286	0.258	0.172	0.158	0.058	0.213
Avg	0.466	0.295	0.232	0.188	0.135	0.092	-

Table 4: Accuracy of the models across the different categories.

Country/Model	Claude	Gemini	GPT	Qwen	Llama	Deepseek	Avg
China	0.49	0.32	0.23	0.19	0.17	0.02	0.24
France	0.45	0.29	0.23	0.21	0.15	0.06	0.23
Germany	0.47	0.31	0.30	0.16	0.10	0.06	0.23
Iran	0.47	0.34	0.25	0.19	0.13	0.05	0.24
Italy	0.52	0.31	0.24	0.17	0.12	0.03	0.23
UK	0.47	0.31	0.22	0.25	0.17	0.08	0.25
USA	0.37	0.18	0.15	0.14	0.10	0.07	0.17
Avg	0.46	0.29	0.23	0.19	0.13	0.06	-

Table 5: Accuracy of the models across the different countries.

Multi-hop type	Claude	Gemini	GPT	Qwen	Llama	DeepSeek	Avg
Comparison	0.593	0.373	0.271	0.153	0.153	0.101	0.274
Geographical	0.439	0.367	0.235	0.224	0.173	0.082	0.253
Intersection	0.466	0.271	0.227	0.171	0.112	0.036	0.214
Named	0.445	0.283	0.191	0.197	0.133	0.035	0.214
Temporal	0.437	0.261	0.277	0.193	0.143	0.067	0.230
Avg	0.476	0.311	0.240	0.188	0.143	0.064	-

Table 6: Accuracy of the models across the different MH types.

ferent knowledge areas.

Table 5 analyzes accuracy across countries. The results indicate close accuracy levels, suggesting that models struggle to detect falsehoods regardless of the country. The lower accuracy in some countries can be a result of more challenging questions, which can suggest that the generator model was better able to design challenging questions for those countries.

Finally, Table 6 presents model accuracy across multi-hop reasoning types. Overall, models perform better on comparison-based questions, while struggling the most with intersection, named-entity, and temporal reasoning.⁹

6 Conclusion

We introduce a novel class of multi-hop false-premise questions (MHFPQs), combining the complexities of multi-hop reasoning and false-premise

detection. To support this research, we present MultiHoax, the first cross-regional FPQ resource, expanding beyond previous benchmarks to assess LLMs’ ability to navigate multi-step falsehoods. Our dataset provides a comprehensive evaluation framework, spanning seven countries and ten knowledge categories, allowing for a detailed analysis of how LLMs handle false premises across diverse topics and regions. Unlike prior FPQ datasets, MHFPQs require deeper reasoning, as falsehoods are not immediately apparent but emerge through multi-step inference. Beyond evaluating LLM reasoning capabilities, MultiHoax serves as a benchmark for advancing research in multi-hop question answering, cross-regional factual consistency, and robustness in complex reasoning tasks, providing insights into LLMs’ ability to handle factually inconsistent information across different domains and regions.

⁹Tables 17 and 18 (Appendix) provide model accuracy across different false premise and answer types.

Limitations

Our dataset has some limitations. While we have aimed to include a diverse range of countries with varying levels of resource availability, there are still opportunities to incorporate additional countries from other regions worldwide.

Second, our category set, which consists of 10 categories, could be expanded as scholars explore knowledge across various areas. While we have focused on a set of categories suitable for factual, objective questions, other potential categories could be included. Additionally, subjective questions, such as “Who first imported the most popular type of ingredient to China?” could be considered. This would be an MHQ, as it requires identifying both the most popular ingredient and the first person to import it. However, determining the most popular ingredient is subjective, and the question becomes MHFP if the ingredient was never imported to China. While subjective questions are feasible, reviewing them differs significantly from the process for factual objective questions.

Furthermore, while we include questions associated with a diverse set of countries, we do not have the translation of these questions in the local languages of these countries, except from the United States and the United Kingdom. Future research can be collecting questions in the local language of such countries or translating ours to those languages.

Finally, while our current resource presents a significant challenge to different LLMs, and even the best models struggle with our tasks, the rapid progress of LLMs may make our dataset less difficult over time. As models improve, we may need to update our resources to introduce more challenging tasks that better test LLMs’ reasoning abilities.

Ethical Considerations

We aim to provide a set of factually incorrect questions requiring multiple reasoning steps to challenge various models’ ability to detect falsehoods.

To compile our dataset, we utilized LLMs to generate potential questions, which greatly facilitated the process. However, this approach may introduce biases, as LLMs are more knowledgeable about certain countries than others. For instance, the USA’s lower accuracy could be attributed to LLMs having more in-depth knowledge of U.S. facts, allowing them to craft more challenging questions. Although we used Wikipedia documents to mitigate this bias,

it cannot be entirely eliminated.

In future iterations, we plan to diversify the question types, incorporating topics focused on values and norms rather than just factual knowledge, and we aim to minimize reliance on LLMs for question generation.

References

- Akari Asai and Eunsol Choi. 2021. [Challenges in information-seeking QA: Unanswerable questions and paragraph retrieval](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 1492–1504, Online. Association for Computational Linguistics.
- Xiao Bi, Deli Chen, Guanting Chen, Shanhuang Chen, Damai Dai, Chengqi Deng, Honghui Ding, Kai Dong, Qishi Du, Zhe Fu, et al. 2024. Deepseek llm: Scaling open-source language models with longtermism. *arXiv preprint arXiv:2401.02954*.
- Abir Chakraborty. 2024. Multi-hop question answering over knowledge graphs using large language models. *arXiv preprint arXiv:2404.19234*.
- Canyu Chen and Kai Shu. 2023. Can llm-generated misinformation be detected? *arXiv preprint arXiv:2309.13788*.
- Hung-Ting Chen, Michael Zhang, and Eunsol Choi. 2022. [Rich knowledge sources bring complex knowledge conflicts: Recalibrating models to reflect conflicting evidence](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 2292–2307, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Ruirui Chen, Weifeng Jiang, Chengwei Qin, Ishaan Singh Rawal, Cheston Tan, Dongkyu Choi, Bo Xiong, and Bo Ai. 2024. Llm-based multi-hop question answering with knowledge graph integration in evolving environments. *arXiv preprint arXiv:2408.15903*.
- Myra Cheng, Tiziano Piccardi, and Diyi Yang. 2023. Compost: Characterizing and evaluating caricature in llm simulations. *arXiv preprint arXiv:2310.11501*.
- Tamás Csöngé. 2015. Moving picture, lying image: Unreliable cinematic narratives. *Acta Universitatis Sapientiae, Film and Media Studies*, (10):89–104.
- Ashwin Daswani, Rohan Sawant, and Najoung Kim. 2024. Syn-qa2: Evaluating false assumptions in long-tail questions with synthetic qa datasets. *arXiv preprint arXiv:2403.12145*.
- Petko Dimov. 2021. Recognition of fake news in sports. , 29(4s):18–27.

732	Yi Fung, Ruining Zhao, Jae Doo, Chenkai Sun, and	Najoung Kim, Phu Mon Htut, Samuel R. Bowman, and	789
733	Heng Ji. 2024. Massively multi-cultural knowledge	Jackson Petty. 2023. (QA) ² : Question answering	790
734	acquisition & lm benchmarking. <i>arXiv preprint</i>	with questionable assumptions. In <i>Proceedings of the</i>	791
735	<i>arXiv:2402.09369</i> .	61st Annual Meeting of the Association for Compu-	792
736	Emilio Gabba. 1981. True history and false history in	tational Linguistics (Volume 1: Long Papers), pages	793
737	classical antiquity. <i>The Journal of Roman Studies</i> ,	8466–8487, Toronto, Canada. Association for Com-	794
738	71:50–62.	putational Linguistics.	795
739	Boris Galitsky, Anton Chernyavskiy, and Dmitry Il-	Fajri Koto, Nurul Aisyah, Haonan Li, and Timothy Bald-	796
740	vovsky. 2024. Truth-o-meter: Handling multiple	win. 2023. Large language models only pass primary	797
741	inconsistent sources repairing llm hallucinations. In	school exams in Indonesia: A comprehensive test on	798
742	<i>Proceedings of the 47th International ACM SIGIR</i>	IndoMMLU. In <i>Proceedings of the 2023 Conference</i>	799
743	<i>Conference on Research and Development in Infor-</i>	on Empirical Methods in Natural Language Process-	800
744	<i>mation Retrieval</i> , pages 2817–2821.	ing, pages 12359–12374, Singapore. Association for	801
745	Bishwamittra Ghosh, Sarah Hasan, Naheed Anjum	Computational Linguistics.	802
746	Arafat, and Arijit Khan. 2024. Logical consistency	Fajri Koto, Haonan Li, Sara Shatnawi, Jad Doughman,	803
747	of large language models in fact-checking. <i>arXiv</i>	Abdelrahman Sadallah, Aisha Alraeesi, Khalid Al-	804
748	<i>preprint arXiv:2412.16100</i> .	mubarak, Zaid Alyafeai, Neha Sengupta, Shady She-	805
749	Seungju Han, Junhyeok Kim, Jack Hessel, Liwei Jiang,	hata, Nizar Habash, Preslav Nakov, and Timothy	806
750	Jiwan Chung, Yejin Son, Yejin Choi, and Young-	Baldwin. 2024. ArabicMMLU: Assessing massive	807
751	jae Yu. 2023. Reading books is great, but not if	multitask language understanding in Arabic. In <i>Find-</i>	808
752	you are driving! visually grounded reasoning about	ings of the Association for Computational Linguistics:	809
753	defeasible commonsense norms. <i>arXiv preprint</i>	ACL 2024, pages 5622–5640, Bangkok, Thailand. As-	810
754	<i>arXiv:2310.10418</i> .	sociation for Computational Linguistics.	811
755	Shengding Hu, Yifan Luo, Huadong Wang, Xingyi	Mahmut Kutlu. 2022. Analysis of false religious posts	812
756	Cheng, Zhiyuan Liu, and Maosong Sun. 2023. Won't	on social media. <i>Ahi Evran Akademi</i> , 3(2):14–30.	813
757	get fooled again: Answering questions with false	João A Leite, Olesya Razuvayevskaya, Kalina	814
758	premises. In <i>Proceedings of the 61st Annual Meet-</i>	Bontcheva, and Carolina Scarton. 2023. Detect-	815
759	<i>ing of the Association for Computational Linguistics</i>	ing misinformation with llm-predicted credibility	816
760	(Volume 1: Long Papers), pages 5626–5643, Toronto,	signals and weak supervision. <i>arXiv preprint</i>	817
761	Canada. Association for Computational Linguistics.	<i>arXiv:2309.07601</i> .	818
762	Kung-Hsiang Huang, Kathleen McKeown, Preslav	Haonan Li, Yixuan Zhang, Fajri Koto, Yifei Yang, Hai	819
763	Nakov, Yejin Choi, and Heng Ji. 2023. Faking	Zhao, Yeyun Gong, Nan Duan, and Timothy Bald-	820
764	fake news for real fake news detection: Propaganda-	win. 2024. CMMLU: Measuring massive multitask	821
765	loaded training data generation. In <i>Proceedings</i>	language understanding in Chinese. In <i>Findings of</i>	822
766	<i>of the 61st Annual Meeting of the Association for</i>	<i>the Association for Computational Linguistics: ACL</i>	823
767	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,	2024, pages 11260–11285, Bangkok, Thailand. As-	824
768	pages 14571–14589, Toronto, Canada. Association	sociation for Computational Linguistics.	825
769	for Computational Linguistics.	Stephanie Lin, Jacob Hilton, and Owain Evans. 2022.	826
770	Aaron Hurst, Adam Lerer, Adam P Goucher, Adam	TruthfulQA: Measuring how models mimic human	827
771	Perelman, Aditya Ramesh, Aidan Clark, AJ Os-	falsehoods. In <i>Proceedings of the 60th Annual Meet-</i>	828
772	trow, Akila Welihinda, Alan Hayes, Alec Radford,	<i>ing of the Association for Computational Linguistics</i>	829
773	et al. 2024. Gpt-4o system card. <i>arXiv preprint</i>	(Volume 1: Long Papers), pages 3214–3252, Dublin,	830
774	<i>arXiv:2410.21276</i> .	Ireland. Association for Computational Linguistics.	831
775	Mehran Kazemi, Quan Yuan, Deepti Bhatia, Najoung	Yunfei Long, Qin Lu, Rong Xiang, Minglei Li, and	832
776	Kim, Xin Xu, Vaiva Imbrasaite, and Deepak Ra-	Chu-Ren Huang. 2017. Fake news detection through	833
777	machandran. 2024. Boardgameqa: A dataset for	multi-perspective speaker profiles. In <i>Proceedings of</i>	834
778	natural language reasoning with contradictory infor-	<i>the Eighth International Joint Conference on Natu-</i>	835
779	mation. <i>Advances in Neural Information Processing</i>	<i>ral Language Processing (Volume 2: Short Papers)</i> ,	836
780	<i>Systems</i> , 36.	pages 252–256, Taipei, Taiwan. Asian Federation of	837
781	Eunsu Kim, Juyoung Suk, Philhoon Oh, Haneul Yoo,	Natural Language Processing.	838
782	James Thorne, and Alice Oh. 2024. CLICk: A bench-	Shayne Longpre, Kartik Perisetla, Anthony Chen,	839
783	mark dataset of cultural and linguistic intelligence	Nikhil Ramesh, Chris DuBois, and Sameer Singh.	840
784	in Korean. In <i>Proceedings of the 2024 Joint In-</i>	2021. Entity-based knowledge conflicts in question	841
785	<i>ternational Conference on Computational Linguistics,</i>	answering. In <i>Proceedings of the 2021 Conference</i>	842
786	<i>Language Resources and Evaluation (LREC-</i>	<i>on Empirical Methods in Natural Language Process-</i>	843
787	<i>COLING 2024)</i> , pages 3335–3346, Torino, Italia.	ing, pages 7052–7063, Online and Punta Cana, Do-	844
788	ELRA and ICCL.	minican Republic. Association for Computational	845
		Linguistics.	846

847	Guanghui Ma, Chunming Hu, Ling Ge, and Hong	Samyo Rode-Hasinger, Anna Kruspe, and Xiao Xiang	902
848	Zhang. 2022. Open-topic false information detec-	Zhu. 2022. True or false? detecting false informa-	903
849	tion on social networks with contrastive adversarial	tion on social media using graph neural networks.	904
850	learning. In <i>Proceedings of the 2022 Conference on</i>	In <i>Proceedings of the Eighth Workshop on Noisy</i>	905
851	<i>Empirical Methods in Natural Language Processing,</i>	<i>User-generated Text (W-NUT 2022)</i> , pages 222–229,	906
852	pages 2911–2923, Abu Dhabi, United Arab Emirates.	Gyeongju, Republic of Korea. Association for Com-	907
853	Association for Computational Linguistics.	putational Linguistics.	908
854	Vaibhav Mavi, Anubhav Jangra, and Adam Jatowt. 2022.	Maryam Sabzali, Majid Sarfi, Mostafa Zohouri, Tahereh	909
855	A survey on multi-hop question answering and gener-	Sarfi, and Morteza Darvishi. 2022. Fake news and	910
856	ation. <i>arXiv preprint arXiv:2204.09140</i> .	freedom of expression: An Iranian perspective. <i>Jour-</i>	911
857	AI Meta, Abhinav Jauhri, Abhinav Pandey, Abhishek	<i>nal of Cyberspace Studies</i> , 6(2):205–218.	912
858	Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil	Sagi Shaier, Ari Kobren, and Philip V. Ogren. 2024.	913
859	Mathur, Alan Schelten, Amy Yang, Angela Fan, et al.	Adaptive question answering: Enhancing language	914
860	2024. The llama 3 herd of models. <i>arXiv preprint</i>	model proficiency for addressing knowledge con-	915
861	<i>arXiv:2407.21783</i> , 2.	flicts with source citations. In <i>Proceedings of the</i>	916
862	Naoaki Okazaki, Keita Nabeshima, Kento Watanabe,	<i>2024 Conference on Empirical Methods in Natural</i>	917
863	Junta Mizuno, and Kentaro Inui. 2013. Extracting	<i>Language Processing</i> , pages 17226–17239, Miami,	918
864	and aggregating false information from microblogs.	Florida, USA. Association for Computational Lin-	919
865	In <i>Proceedings of the Workshop on Language Pro-</i>	guistics.	920
866	<i>cessing and Crisis Information 2013</i> , pages 36–43,	Qiang Sheng, Juan Cao, H Russell Bernard, Kai Shu,	921
867	Nagoya, Japan. Asian Federation of Natural Lan-	Jintao Li, and Huan Liu. 2022. Characterizing multi-	922
868	guage Processing.	domain false news and underlying user effects on	923
869	Jinyoung Park, Ameen Patel, Omar Zia Khan, Hyun-	Chinese Weibo. <i>Information Processing & Manage-</i>	924
870	woo J Kim, and Joo-Kyung Kim. 2023. Graph-	<i>ment</i> , 59(4):102959.	925
871	guided reasoning for multi-hop question answer-	Weiyan Shi, Ryan Li, Yutong Zhang, Caleb Ziems, Raya	926
872	ing in large language models. <i>arXiv preprint</i>	Horeh, Rogério Abreu de Paula, Diyi Yang, et al.	927
873	<i>arXiv:2311.09762</i> .	2024. Culturebank: An online community-driven	928
874	Seong-Il Park and Jay-Yoon Lee. 2024. Toward robust	knowledge base towards culturally aware language	929
875	RALMs: Revealing the impact of imperfect retrieval	technologies. <i>arXiv preprint arXiv:2404.15238</i> .	930
876	on retrieval-augmented language models. <i>Transac-</i>	Andrew F Smith. 2001. False memories: The invention	931
877	<i>tions of the Association for Computational Linguis-</i>	of culinary fakelore and food fallacies. In <i>Food and</i>	932
878	<i>tics</i> , 12:1686–1702.	<i>the Memory: Proceedings of the Oxford Symposium</i>	933
879	Amirreza Payandeh, Dan Pluth, Jordan Hosier, Xuesu	<i>on Food and Cookery</i> , pages 254–260.	934
880	Xiao, and Vijay K. Gurbani. 2024. How susceptible	Guijin Son, Hanwool Lee, Sungdong Kim, Seungone	935
881	are LLMs to logical fallacies? In <i>Proceedings of</i>	Kim, Niklas Muennighoff, Taekyoon Choi, Cheon-	936
882	<i>the 2024 Joint International Conference on Compu-</i>	bok Park, Kang Min Yoo, and Stella Biderman.	937
883	<i>tational Linguistics, Language Resources and Eval-</i>	2024. Kmmlu: Measuring massive multitask lan-	938
884	<i>uation (LREC-COLING 2024)</i> , pages 8276–8286,	guage understanding in Korean. <i>arXiv preprint</i>	939
885	Torino, Italia. ELRA and ICCL.	<i>arXiv:2402.11548</i> .	940
886	Zhiyuan Peng, Jinming Nian, Alexandre Evfimievski,	Yixuan Tang and Yi Yang. 2024. Multihop-rag: Bench-	941
887	and Yi Fang. 2024. Scopeqa: A framework for gener-	marking retrieval-augmented generation for multi-	942
888	ating out-of-scope questions for rag. <i>arXiv preprint</i>	hop queries. <i>arXiv preprint arXiv:2401.15391</i> .	943
889	<i>arXiv:2410.14567</i> .	Gemini Team, Rohan Anil, Sebastian Borgeaud, Jean-	944
890	Yu Qiao, Daniel Wiechmann, and Elma Kerz. 2020.	Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan	945
891	A language-based approach to fake news detection	Schalkwyk, Andrew M Dai, Anja Hauth, Katie	946
892	through interpretable features and BRNN. In <i>Pro-</i>	Millican, et al. 2023. Gemini: a family of	947
893	<i>ceedings of the 3rd International Workshop on Ru-</i>	highly capable multimodal models. <i>arXiv preprint</i>	948
894	<i>mours and Deception in Social Media (RDSM)</i> , pages	<i>arXiv:2312.11805</i> .	949
895	14–31, Barcelona, Spain (Online). Association for	Thaer Thaher, Mahmoud Saheb, Hamza Turabieh, and	950
896	Computational Linguistics.	Hamouda Chantar. 2021. Intelligent detection of	951
897	Hossein Rajabzadeh, Suyuchen Wang, Hyock Ju Kwon,	false information in Arabic tweets utilizing hybrid	952
898	and Bang Liu. 2023. Multimodal multi-hop question	Harris Hawks based feature selection and machine	953
899	answering through a conversation between tools and	learning models. <i>Symmetry</i> , 13(4):556.	954
900	efficiently finetuned large language models. <i>arXiv</i>	Khurram Yamin, Shantanu Gupta, Gaurav R Ghosal,	955
901	<i>preprint arXiv:2309.08922</i> .	Zachary C Lipton, and Bryan Wilder. 2024. Failure	956
		modes of LLMs for causal reasoning on narratives.	957
		<i>arXiv preprint arXiv:2410.23884</i> .	958

959	An Yang, Baosong Yang, Binyuan Hui, Bo Zheng,	Wikipedia Document Retrieval Prompt	1015
960	Bowen Yu, Chang Zhou, Chengpeng Li, Chengyuan	In this section, we first define each category. Then,	1016
961	Li, Dayiheng Liu, Fei Huang, et al. 2024a. Qwen2	we them to retrieve relevant documents for each cat-	1017
962	technical report. <i>arXiv preprint arXiv:2407.10671</i> .	egory using ChatGPT-4o’s search using the prompt	1018
963	Sohee Yang, Elena Gribovskaya, Nora Kassner, Mor	in Table 7	1019
964	Geva, and Sebastian Riedel. 2024b. Do large lan-	Food	1020
965	guage models latently perform multi-hop reasoning?	• Cuisine: Signature dishes, cooking styles, tra-	1021
966	<i>arXiv preprint arXiv:2402.16837</i> .	ditional meals.	1022
967	Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio,	• Ingredients: Locally grown spices, crops, and	1023
968	William Cohen, Ruslan Salakhutdinov, and Christo-	special ingredients.	1024
969	pher D. Manning. 2018. HotpotQA: A dataset for	• Drinks: Popular beverages, traditional teas,	1025
970	diverse, explainable multi-hop question answering .	or alcoholic drinks.	1026
971	In <i>Proceedings of the 2018 Conference on Empiri-</i>	Sports	1027
972	<i>cal Methods in Natural Language Processing</i> , pages	• National and Popular Sports: Widely	1028
973	2369–2380, Brussels, Belgium. Association for Com-	played or watched sports in the country and	1029
974	putational Linguistics.	Official sports of a country.	1030
975	Xinyan Yu, Sewon Min, Luke Zettlemoyer, and Han-	• Athletes: Famous sportspeople or Olympic	1031
976	naneh Hajishirzi. 2023. CREPE: Open-domain ques-	medalists.	1032
977	tion answering with false presuppositions . In <i>Pro-</i>	• Tournaments and Sports Venues: Major	1033
978	<i>ceedings of the 61st Annual Meeting of the Associa-</i>	leagues, championships, or cups, as well as	1034
979	<i>tion for Computational Linguistics (Volume 1: Long</i>	iconic stadiums, arenas, or tracks.	1035
980	<i>Papers)</i> , pages 10457–10480, Toronto, Canada. As-	Education	1036
981	sociation for Computational Linguistics.	• Education System AND Literacy: Structure	1037
982	Hanning Zhang, Shizhe Diao, Yong Lin, Yi Fung, Qing	(primary, secondary, higher education) AND	1038
983	Lian, Xingyao Wang, Yangyi Chen, Heng Ji, and	Efforts to promote literacy or improve access	1039
984	Tong Zhang. 2024a. R-tuning: Instructing large lan-	to education.	1040
985	guage models to say ‘I don’t know’ . In <i>Proceedings</i>	• Schools and Universities AND Curriculum:	1041
986	<i>of the 2024 Conference of the North American Chap-</i>	Prestigious or historic institutions AND Sub-	1042
987	<i>ter of the Association for Computational Linguistics:</i>	jects emphasized or unique courses.	1043
988	<i>Human Language Technologies (Volume 1: Long</i>	• Famous Educators: Scholars, reformers, or	1044
989	<i>Papers)</i> , pages 7113–7139, Mexico City, Mexico. As-	pioneers in education.	1045
990	sociation for Computational Linguistics.	Holidays/Celebrations/Leisure	1046
991	Wenyuan Zhang, Jiawei Sheng, Shuaiyi Nie, Zefeng	• National Holidays: Independence days, con-	1047
992	Zhang, Xinghua Zhang, Yongquan He, and Tingwen	stitution days, or memorials.	1048
993	Liu. 2024b. Revealing the challenge of detecting	• Festivals: Cultural, religious, or seasonal fes-	1049
994	character knowledge errors in llm role-playing. <i>arXiv</i>	tivals.	1050
995	<i>preprint arXiv:2409.11726</i> .	• Others: Other topics related to Holi-	1051
996	Wenting Zhao, Ge Gao, Claire Cardie, and Alexander M	days/Celebrations/Leisure.	1052
997	Rush. 2024. I could’ve asked that: Reformulating		
998	unanswerable questions . In <i>Proceedings of the 2024</i>		
999	<i>Conference on Empirical Methods in Natural Lan-</i>		
1000	<i>guage Processing</i> , pages 4207–4220, Miami, Florida,		
1001	USA. Association for Computational Linguistics.		
1002	Yuchen Zhuang, Yue Yu, Kuan Wang, Haotian Sun, and		
1003	Chao Zhang. 2023. Toolqa: A dataset for llm ques-		
1004	tion answering with external tools . In <i>Advances in</i>		
1005	<i>Neural Information Processing Systems</i> , volume 36,		
1006	pages 50117–50143. Curran Associates, Inc.		
1007	Caleb Ziems, Jane Dwivedi-Yu, Yi-Chia Wang, Alon		
1008	Halevy, and Diyi Yang. 2023. Normbank: A knowl-		
1009	edge bank of situational social norms. <i>arXiv preprint</i>		
1010	<i>arXiv:2305.17008</i> .		
1011	Appendix A: Prompts		
1012	In this section, we provide the different prompts		
1013	that we leveraged to collect the Wikipedia pages		
1014	and also prompts to generate the questions.		

We are developing a dataset of factual questions and aiming to collect questions across various categories and countries. Currently, we are identifying relevant important Wikipedia pages for each category and country. Please gather Wikipedia pages related to [CATEGORY] for the following countries:

- China
- France
- Germany
- Italy
- United States
- United Kingdom
- Iran

Please provide **15** distinct Wikipedia pages related to [CATEGORY] for the mentioned countries, based on the following category definition below. DEFINITION OF THE CATEGORY>DEFINITION OF THE CATEGORY
ONLY PROVIDE THE PAGE NAMES LINKED TO THE PAGE WITHOUT ANY EXPLANATION.

Table 7: The prompt for searching for related documents using ChatGPT4o.

1053	History	• Others: Other related topics to Science and Technology like Medical Breakthroughs, Research Centers, Computing Pioneers, Green Technology, Digital Platforms, and Communications.	1073
1054	• Historical Figures: Leaders, revolutionaries, empires and kingdoms, or intellectuals.		1074
1055			1075
1056	• Important Events: Battles, treaties, or turning points in history.		1076
1057		Arts and Literature	1077
1058	• Landmarks: Historical monuments or UNESCO heritage sites.	• Writers: Prominent authors, poets, or playwrights.	1079
1059			1080
1060	Geography	• Books: National epics, famous novels, or historical documents.	1081
1061	• Natural Features AND Resources: Mountains, rivers, lakes, and deserts AND Natural resources, agriculture, or energy production.	• Artists: Prominent artists.	1082
1062		Religion	1083
1063	• Cities AND Regions: Capitals, major cities, or urban landmarks AND Administrative divisions or cultural regions.	• Religions and Religious Practices: Popular religions and Worship styles or Religious rituals.	1084
1064			1085
1065	• Geopolitics: Borders, neighbors, or disputed territories.	• Holy Sites: Temples, churches, mosques, or pilgrimage locations.	1086
1066			1087
1067		• Others: Religious Leaders, Religious Festivals, or Sacred Texts.	1088
1068			1089
1069	Science and Technology		1090
1070	• Scientists: Modern renowned scientists.	Entertainment	1091
1071	• Engineering: Famous modern constructions, bridges, or technology.	• Cinema and TV: National cinema, famous directors, popular movies, or actors.	1092
1072			1093
			1094

1095	• Music: Traditional music styles, Musicians, or iconic bands.	Appendix B: Annotation Guideline	1123
1096		Annotation Guideline	1124
1097	• Others: Other topics related to entertainment like theater, gaming, festivals related to entertainment, or media.	The annotation guideline, shown in the Table 12	1125
1098		was used to familiarize false information reviewers with the dataset structure and their tasks. The guideline first defines what an FPQ is, then explains how questions are structured in our dataset, and finally outlines the reviewing task.	1126
1099			1127
1100	These category definitions were used in the prompt in the Table 7 to get 15 related pages for each country.		1128
1101			1129
1102			1130
1103	Fact Extraction Prompt	Annotators details	1131
1104	To extract facts from each document, we leveraged a simple brief prompt. In the fact extraction prompt, we define the number of facts to be extracted, the name of the document, and the document content.	Following the initial round of annotation by the authors, the dataset was divided into three parts for the second review phase. We then recruited three university student annotators—one female and two male—each paid £12.21 per hour for ten hours of work.	1132
1105	Extract 15 facts from the document. Document Content: DOCUMENT CONTENT		1133
1106			1134
1107			1135
1108			1136
1109			1137
1110	Question Generation Prompt		
1111	For question generation, we designed two distinct prompts corresponding to the two major types of MHQs: Entity-based and Intersection & Comparison Combination. Each prompt consists of three parts: an introduction to the task, an explanation of the specific MHQ type, and a set of generation rules. The first section, shown in the Table 8, and the last section, shown in the Table 9 are the same in both generation prompts, but the second part, the definition of the specific MHQ types, is different.		
1112	The tables 10 and 11 show the the two different second sections.		
1113			
1114			
1115			
1116			
1117			
1118			
1119			
1120			
1121			
1122			

We are developing a dataset of counterfactual multi-hop reasoning false-premise questions (MHFPQs). Multi-hop reasoning questions require retrieving and connecting multiple pieces of information across two or more logical steps to derive the final answer. We have seen what a multi-hop question is. Now, let's focus on creating questions with false premises—where in the question there exists a false assumption that makes the question wrong and impossible to answer correctly. MHFPQs have similar types like Multi-hop but with false-premises. Here are some examples: Note that these are only examples, DO NOT include them in your answers.

Table 8: The first part of the generation prompt.

Your task is to extract false premise multi-hop questions FROM THE FACTS PROVIDED. Here are the instructions:

- The structure of the question should be like the given structures, but the content can be different.
- False premises are implicitly embedded within the questions. Also, false premises must not be obvious.
- Provide 3 relevant, engaging, and realistic options.
- Questions should be based on the provided FACTS from the specific country.
- Focus is exclusively on verifiable factual claims, avoiding cultural norms or subjective topics.
- For each question, a clear explanation must be provided:
 - Identifying the false premise.
 - Clarifying the actual truth.
- If it is not possible to design questions from all the types, you can only focus on most probable ones.

Return each question in the following format:

<false_premise_multi_hop_question> | <first_option> | <second_option> | <third_option> | <description>
| <reasoning_steps> | <type_of_multi_hop>

Table 9: The third part of the generation prompt.

- **Temporal Multi-Hop Reasoning**

Example: Who was the president of Iran in the year in which Ehsan Rouzbahani won the Olympics Bronze medal in Tokyo?

Description: Ehsan Rouzbahani did not compete in the Tokyo Olympics, making it impossible for him to win a (bronze) medal.

Reasoning Steps:

1. The year when Ehsan Rouzbahani won the Olympic bronze medal in Tokyo.
2. The president of Iran at that time.

- **Geographical Multi-Hop Reasoning**

Example: Which football team with the most championships in the territory Alexander the Great conquered before turning 18?

Description: Alexander the Great did not conquer any territory before turning 18.

Reasoning Steps:

1. The territory Alexander the Great conquered before turning 18.
2. Football team with the most championships in that territory.

- **Named Entity Multi-Hop Reasoning**

Example: What is the camera brand used by Spielberg when filming his Academy Award-winning student film at USC?

Description: Spielberg never attended USC and didn't win an Academy Award as a student.

Reasoning Steps:

1. The Spielberg's Academy Award-winning student film at USC.
 2. The camera brand used for that film.
-

Table 10: The second part of the generation prompt for the first group.

- **Intersection-Type Multi-Hop Reasoning**

Example: Which architect both designed the golden-domed Old Basilica and incorporated Aztec symbols in its facade in 1695?

Description: The Old Basilica did not have Aztec symbols. It only had yellow and blue Talavera mosaics.

Reasoning Steps:

1. The architect who designed the golden-domed Old Basilica.
2. The architect who incorporated Aztec symbols in the Old Basilica’s facade in 1695.

- **Comparison-Type Multi-Hop Reasoning**

Example: Which of the Chinese and the Germans first invented sauerkraut in the 18th century?

Description: Sauerkraut was not invented in the 18th century, and it existed much earlier.

Reasoning Steps:

1. Did the Chinese invented sauerkraut in the 18th century?
2. Did Germans invented sauerkraut in the 18th century?

Note: In Comparison type, multi-hop reasoning questions, none of the entities satisfies the condition.

Table 11: The second part of the generation prompt for the second group.

Appendix C: Statistics

In this section, we provide further statistical information regarding the dataset.

Table 13 provides the distribution of the various types of MHQs. The majority of the resource consists of bridge entity-based questions, while the intersection type is largely included.

Table 14 shows the distribution of the various types of false information included in the questions. We include four types in our resources. The description of types is also provided in Table 14.

Moreover, Table 15 provides the details of answer type occurrence over the resource the person type as the major type.

Regarding the annotation details, 16 contains the number of times the second reviewer chose each possible label for the questions.

The table 17 provides the detailed accuracy of the models across different types of answers. As shown, event, location, and number have the lowest accuracies suggesting that the models find these types more challenging than the rest while they generally show poor performance.

Tables 6 and 18 also show the detailed accuracies for different types of MHQs and FPQs.

False-premise Questions: We have a set of false-premise questions. A false premise question is a question with at least one piece of false information. A simple example of false-premise questions can be “How many eyes does the sun have?” Such simple questions include false information that is easily detectable by humans. More challenging false-premise questions, which are the target of our experiment, have false pieces of information that are difficult for non-expert humans to detect.

Question: During which time, 1985 to 1990 or 1995 to 2010, did CERN’s affiliates win more Nobel prizes in physics?

1. 1985 to 1990
2. 1995 to 2010
3. In both durations, CERN’s affiliates won only 1 Nobel prize
4. I don’t know

Explanation: In none of the durations, CERN’s affiliates won a Nobel prize. 1984, 1992, and 2013 are the years when CERN’s affiliates won the award.

As you can see, such false information types are not detectable unless the person knows about the history of the mentioned institute, which is not the case for non-experts.

Format of Data: In the dataset, there are a number of fields, like the following example.

Question: What is the name of the new Humanistic Buddhist organization that was established in Beijing in the 2000s to promote the revival of Vajrayana Buddhism in China? Options:

1. Cǐjì
2. Huácáng Zōngmén
3. Zhēnfó Zōng
4. I don’t know

Description: The question contains a false premise that a new Humanistic Buddhist organization was established in Beijing in the 2000s to promote the revival of Vajrayana Buddhism. According to the facts, the Humanistic Buddhist movement in China is associated with organizations like Cǐjì, which has been working in mainland China since 1991, not a new organization focused on Vajrayana Buddhism.

Wikipedia: https://en.wikipedia.org/wiki/Buddhism_in_China

Your task: You are supposed to read the question and options, and then check the provided description, which is the description of why the question includes false information. Afterward, you are supposed to label these questions after checking if there is any false information in them or not. “*There is false information*”, “*There is no false information*”, and “*I cannot tell based on the provided information*” are the possible options. You need to visit the Wikipedia page related to each question and check false information based on that **Wikipedia** page.

- If you choose “*There is no false information*”, then you are supposed to explain why you have chosen this. For example, in the above case, if you choose “*There is no false information*”, then an example explanation can be “CERN’s affiliates won the Nobel prize in 1986 making the question true”.
- If you choose “*I cannot tell based on the provided information*”, you must also explain the ambiguity or the problem you have in verifying the question in the explain column.
- If you choose “*There is false information*”, then there is no need to explain.

Keep the explanation clear, simple, and concise.

Table 12: Reviewing guidelines.

Type	Percentage (%)
Intersection	36
Named-entity	25
Temporal-entity	17
Geographical-entity	14
Comparison	8

Table 13: Distribution of MH types in the dataset.

Type	Description	Percentage (%)
Event	The event didn't happen in history.	49
Property	The entity does not have the property.	36
Scope	A fact is not valid in the scope.	13
Entity	The entity cannot exist.	2

Table 14: Distribution of FP types in the dataset.

Type	Percentage (%)
Person	44
Group/Org	13
Location	8
Number	7
Language or Nationality or Country	6
Date	6
Common noun	5
Food	4
Artwork	3
Event	2
Other proper noun	2

Table 15: Distribution of Answer types in the dataset.

File	False Information	Cannot Tell	No False Information
Science and technology	49	12	9
Entertainment	66	3	1
Education	65	5	0
Art and Literature	59	11	0
Food	62	6	2
Religion	59	6	5
Sports	56	8	6
Holiday, Celebrations, and Leisure	52	4	14
Geography	51	3	16
History	57	0	13
Sum	576	57	67

Table 16: Analysis of false information inclusion across different categories based on the second phase review.

Answer Type	Claude	Gemini	GPT	Qwen	Llama	DeepSeek	Avg
Artwork	0.48	0.36	0.28	0.16	0.08	0.04	0.233
Common Noun	0.459	0.351	0.243	0.243	0.162	0.162	0.270
Date or Time Period	0.6	0.45	0.375	0.275	0.2	0.025	0.321
Event	0.429	0.143	0.214	0.143	0.071	0.0	0.167
Food	0.577	0.308	0.192	0.192	0.038	0.077	0.231
Group or Org	0.494	0.205	0.241	0.181	0.108	0.060	0.215
Language/Nationality/Country	0.595	0.238	0.238	0.167	0.167	0.095	0.250
Location	0.339	0.271	0.186	0.136	0.136	0.034	0.184
Number	0.333	0.333	0.208	0.146	0.167	0.0	0.198
Other Proper Noun	0.417	0.25	0.167	0.167	0.167	0.167	0.223
Person	0.453	0.296	0.219	0.193	0.135	0.045	0.223
Avg	0.477	0.278	0.226	0.185	0.125	0.066	0.226

Table 17: Accuracy of the models across different answer types.

FP Type	Claude	Gemini-2.0	GPT-4	Qwen2.5	Llama-3.1	DeepSeek-7B	Avg
Entity	0.778	0.510	0.333	0.278	0.278	0.111	0.381
Event	0.451	0.309	0.209	0.197	0.151	0.062	0.230
Property	0.478	0.251	0.247	0.183	0.112	0.048	0.220
Scope	0.400	0.311	0.244	0.144	0.111	0.022	0.205
Avg	0.527	0.345	0.258	0.201	0.163	0.061	0.259

Table 18: Accuracy of the models across different FP types.