

I Could've Asked That: Reformulating Unanswerable Questions

Wenting Zhao Ge Gao Claire Cardie Alexander M. Rush

Cornell University

{wz346, gg462, ctc9, arush}@cornell.edu

Abstract

When seeking information from unfamiliar documents, users frequently pose questions that cannot be answered by the documents. While existing large language models (LLMs) identify these unanswerable questions, they do not assist users in reformulating their questions, thereby reducing their overall utility. We curate **COULDASK**, an evaluation benchmark composed of existing and new datasets for document-grounded question answering, specifically designed to study reformulating unanswerable questions. We evaluate state-of-the-art open-source and proprietary LLMs on **COULDASK**. The results demonstrate the limited capabilities of these models in reformulating questions. Specifically, GPT-4 and Llama2-7B successfully reformulate questions only 26% and 12% of the time, respectively. Error analysis shows that 62% of the unsuccessful reformulations stem from the models merely rephrasing the questions or even generating identical questions. We publicly release the benchmark¹ and the code to reproduce the experiments².

1 Introduction

Applying large language models (LLMs) to perform question answering (QA) over documents, such as legal and medical texts, has become increasingly popular (Agrawal et al., 2022; Guha et al., 2023). However, users' limited knowledge of these documents often results in the formulation of unanswerable questions, whose assumptions either conflict with or cannot be verified with the information available in the documents. We will refer to these assumptions as *presupposition errors*.³ Gao et al. (2023) and Yu et al. (2023) found that around 30% of information-seeking questions written by

¹<https://huggingface.co/datasets/wentingzhao/couldask>

²<https://github.com/wenting-zhao/couldask>

³Kim et al. (2023) refer to these assumptions as questionable assumptions.

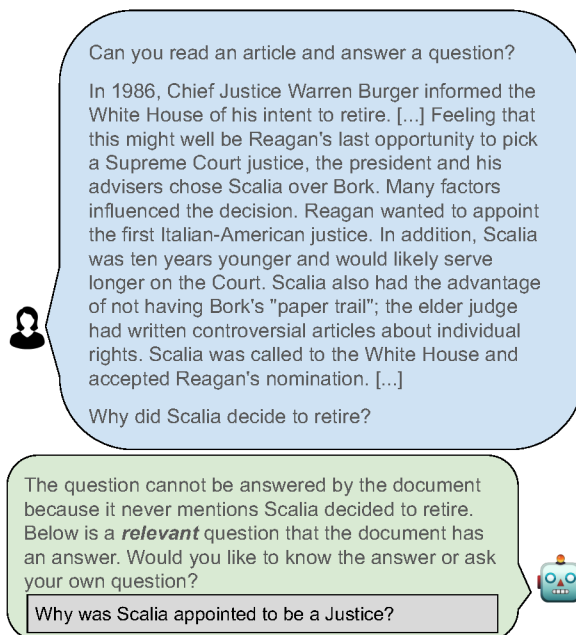


Figure 1: An example of an LLM suggesting an alternative relevant question the user could have asked whose answers can be found in the document, besides only informing users with the presupposition errors.

users include presupposition errors. Research in the field has primarily focused on the detection of unanswerable questions (Rajpurkar et al., 2018; Tran et al., 2023; Hu et al., 2023) and providing explanations for why such questions cannot be answered (Yu et al., 2023; Kim et al., 2023). However, this goal is insufficient for fostering an effective interaction between users and LLMs. Identifying unanswerable questions only serves as a starting point in question reformulation; without additional guidance or feedback on how to rephrase the question, users, especially those unfamiliar with the content, might find themselves caught in a repetitive cycle of formulating questions. In a large-scale industrial experiment, Faustini et al. (2023) have shown that the practice of rewriting unanswerable questions users ask virtual assistants significantly

enhances the experience for millions of users.

This work aims to improve the utility of QA systems by introducing a new task that requires both detecting unanswerable questions and generating questions closely related to the initial query and grounded in the document. Opting to generate a relevant question rather than a summary of related information emphasizes a user-centered approach. While producing a summary is more document-focused, formulating a relevant question targets understanding and predicting the user’s intent, aligning the interaction more closely with the user’s specific needs and queries. We provide an example of how to suggest such questions in Figure 1. Although generating relevant and grounded questions conditioned on initial queries offers greater utility to users, it remains a difficult task even for the best models in a few-shot setting.

We first characterize human-reformulated questions and describe several different strategies for updating questions to remove presupposition errors. Motivated by these strategies, we curate COUL-DASK, an evaluation benchmark for document-grounded QA that consists of a combination of existing and new datasets to study question reformulation in the presence of presupposition errors. We evaluate several prompting methods such as few-shot prompting and chain-of-thought prompting, employing state-of-the-art open-source and proprietary models. The results illustrate the limitations of existing models and prompting techniques in accurately detecting unanswerable questions and reformulating questions: the F1 scores for identifying unanswerable questions range from 41.16% to 67.82%, and success rates for reformulating questions range from 7.13% to 26.21%, depending on base models. Analysis shows that most of the unsuccessful reformulation come from rephrasing or repeating the original questions and that LLMs are worse at reformulating questions necessitating global edits compared to those solely requiring local edits.

2 Related Work

Several datasets have been proposed to study unanswerable questions. [Rajpurkar et al. \(2018\)](#) curate the first document-grounded QA dataset that features unanswerable questions. More recently, [Yu et al. \(2023\)](#) and [Kim et al. \(2023\)](#) have collected questions with presupposition errors from Google user queries and Reddit posts, respectively.

How to identify unanswerable questions, especially with off-the-shelf LLMs, has remained understudied. [Kim et al. \(2021\)](#) proposed to first extract presuppositions from a question and then perform natural language inference (NLI) to check for presupposition violations. However, this pipeline requires supervision. In practice, supervision is often not available; [Kim et al. \(2023\)](#) thus explore prompting large language models in the chain-of-thought style to identify unanswerable questions. However, the results remain unsatisfying; using GPT-3 only yields detection accuracy that is only slightly better than random guesses.

[Faustini et al. \(2023\)](#) investigate unanswerable questions and their reformulations in the domain of spoken QA, focusing on issues stemming from disfluencies, grammatical errors, and awkward phrasing. We, however, study unanswerable questions arising from presupposition errors, which require a more profound semantic comprehension of both contexts and questions.

Unanswerable questions are closely related to ambiguous questions ([Min et al., 2020](#)). While there has been extensive research into reformulating ambiguous questions ([Rao and Daumé III, 2018](#); [White et al., 2021](#); [Pyatkin et al., 2023](#); [Majumder et al., 2021](#)), the problem of reformulating unanswerable questions receives little attention. Strategies for rephrasing ambiguous questions often involve making questions more specific by mentioning precise entities or events ([Min et al., 2020](#)). In contrast, as we will show in Section 3, reformulating unanswerable questions necessitates a wide range of strategies.

Finally, we discuss the connection between document-grounded QA and open-domain QA ([Kim et al., 2023](#)). Document-grounded QA is essentially open-domain QA with the correct document retrieved. Reformulating questions based on the identified document is a separate skill from retrieving the document. Therefore, question reformulation is an interesting task in itself.

3 Task: Question Reformulation

To assist users with question reformulation when reading unfamiliar documents, we define the following task. Given a document and a user question, the system must determine if the question is unanswerable. Upon identifying the unanswerable question, it must reformulate the question such that the new question is answerable by the document while remaining relevant to the original question.

Issue	Strategy	%	Original question	Revised question
Contradictory	Correct	17	What other foods are adulterated with turmeric? Which major cities are located at great lakes?	What other substances are commonly adulterated in turmeric? Which U.S. states have jurisdictions that extend into the Great Lakes?
Unverifiable	Generalize	20	When was the Twenty-third Amendment to the United States Constitution passed into law? How many children does Jennifer Lawrence have?	Has the Twenty-third Amendment to the United States Constitution been proposed? Is Jennifer Lawrence expecting her first child?
Unverifiable	Nearest Match	44	How many arms does Krishna have? How many schools are in CPS?	How is Krishna typically depicted in terms of arms? What kind of schools are included in CPS?
Unverifiable	Specify	19	While working at Edison, what inventions did Nikola Tesla make? How many children did Toni Morrison have?	What specific aspects of electrical engineering did Nikola Tesla work on and further develop while working at Edison? How many children did Toni Morrison have with Harold Morrison before their divorce?

Table 1: Strategies from a human user reformulating questions on 100 examples.

	Total	Unans%	Document Length	Question Length	# Question Entities	Source	Domain
SQuADv2	1000	50.70	143.83 $\pm$ 59.97	11.21 $\pm$ 3.54	2.29 $\pm$ 0.94	Human	Wikipedia
QA ²	506	48.81	818.28 $\pm$ 684.61	7.74 $\pm$ 1.39	2.05 $\pm$ 0.57	Human	Mostly Wikipedia
BanditQA	2070	35.56	261.09 $\pm$ 145.19	8.37 $\pm$ 2.43	1.70 $\pm$ 0.72	Human	Wikipedia
BBC	278	21.22	477.75 $\pm$ 289.41	16.37 $\pm$ 3.49	3.31 $\pm$ 1.04	GPT-4	News
Reddit	313	36.10	477.43 $\pm$ 307.18	14.27 $\pm$ 2.75	2.93 $\pm$ 0.78	GPT-4	Social Media
Yelp	165	30.91	387.70 $\pm$ 147.83	15.27 $\pm$ 3.04	3.04 $\pm$ 0.76	GPT-4	Review

Table 2: An overview of COULDASK. Unans% is the percentage of unanswerable questions.

As this task is challenging to formally define, we begin with a qualitative study over a set of example reformulations by a human user, shown in Table 1. Different strategies are applied for different presupposition errors in the reformulation process. For handling presuppositions that are contradictory to the documents, the human user corrects the presuppositions. When it comes to presuppositions that are unverifiable given the documents, we observe three strategies. The first strategy takes a step back by asking about a less specific event than the event in the original question. The second strategy seeks a “nearest match” question that the document can answer due to a flaw in the original. The third strategy refines the original question by asking about something more specific that can be verified by the document. While these strategies are not exhaustive, they demonstrate the challenging nature of the problem and the necessity for establishing sources of ground truth in the document.

4 The COULDASK Benchmark

Motivated by the need to reformulate questions both to rely on verified information and to avoid contradictions, we develop an evaluation bench-

mark. We consider two important challenges in constructing benchmarks for this task. (1) The benchmark should cover a wide range of domains to cover different types of presuppositions. (Existing QA datasets that study unanswerable questions mostly rely on Wikipedia articles (Rajpurkar et al., 2018; Gao et al., 2023; Kim et al., 2023).) (2) The evaluation method should be capable of fairly assessing equally good reformulated questions, considering the subjective nature of question reformulation.

4.1 Datasets

Following the desiderata, we select three existing datasets – SQuADv2, BanditQA (Gao et al., 2023), and QA² (Kim et al., 2023) – and to cover a broader range of domains, we create three new datasets in the domains of news, review, and Reddit, where the questions are generated by models and verified by crowdworkers. We summarize statistics for all datasets in Table 2.

BBC, Reddit, and Yelp. We synthetically generate questions with a question generation model.

Example	Revised question	Ans	Rel
Question: When did Chick-fil-A open their first restaurant in Pennsylvania? Document: [...] he registered the name Chick-fil-A, Inc. From 1964 to 1967, the sandwich was licensed to over fifty eateries, including Waffle House and the concession stands of the new Houston Astrodome. The Chick-Fil-A sandwich was withdrawn from sale at other restaurants when the first standalone location opened in 1967, in the food court of the Greenbriar Mall, in a suburb of Atlanta. Since 1997, the Atlanta-based company has been the title sponsor of the Peach Bowl, an annual college football bowl game played in Atlanta on New Year’s Eve.	• Did the article mention the first Chick-fil-A restaurant in Pennsylvania?	✓	✓
	• Which Chick-fil-A mentioned in the article is the closest to Pennsylvania?	✓	✓
	• When did Chick-fil-A open their first free-standing location?	✓	✗
	• How many eateries of Chick-fil-A are licensed from 1964 to 1967?	✓	✗
	• Where is the first Chick-fil-A restaurant in Pennsylvania?	✗	✓
	• When did Chick-fil-A open its first restaurant in a northern state, specifically in Pennsylvania?	✗	✓
	• When were chicken sandwiches invented at Chick-fil-A?	✗	✗

Table 3: Different aspects of question reformulation. Ans indicates whether the reformulated question can be answered by the document, and Rel indicates whether the reformulated question is relevant to the original question.

We use the documents from BBC news articles⁴, Reddit pages⁵, and Yelp reviews⁶, respectively. To not artificially craft unanswerable questions, we instruct the question generation model to produce questions normally and later identify the unanswerable ones. To produce a dataset that is challenging for LLMs, we additionally leverage a question checking model⁷. We search for questions that confuse the question checking model. Specifically, we sample the question checking model five times to produce an answer for each of the questions. We gather questions where the question checking model flags the questions are unanswerable any of those five times. We then ask three crowdworkers from Amazon Mechanical Turk (MTurk) to independently verify whether the question is answerable or not⁸. The question generation model produces 9500, 9964, and 10000 QA pairs for the BBC, Reddit, and Yelp datasets, respectively. We keep 278, 313, and 165 examples that have confused the question checking model. Finally, the crowdworkers identify 21.22%, 36.10%, and 30.91% of the questions are truly unanswerable. We thus construct examples that are both challenging to LLMs and high-quality.

⁴<https://huggingface.co/datasets/SetFit/bbc-news>

⁵https://huggingface.co/datasets/reddit_tifu

⁶https://huggingface.co/datasets/yelp_review_full

⁷GPT-4 is used for both question generation model and question checking model.

⁸In particular, to annotate whether a question is answerable, we instruct the crowdworkers to select a span from the document to answer the question. If they cannot identify a span, the question is deemed unanswerable. To reduce noise, we take the majority vote from three crowdworkers. To further ensure a low noise level in annotations, we remove the examples where answer spans annotated by different crowdworkers are disjoint from each other.

SQuADv2, QA², and BanditQA. We adapt these datasets to be used for document-grounded QA. Questions in SQuADv2 are mechanistically formulated with the intention of being unanswerable, whereas the other two datasets feature naturally occurring unanswerable questions. QA² is composed of natural Google search queries, and BanditQA includes questions formulated by users during interactions with LLMs. More information on the modifications we make to these datasets can be found in Appendix A.

4.2 Evaluating Reformulation

To improve the utility of QA systems in responding to unanswerable questions, the reformulated questions must be (1) answerable by the documents and (2) relevant to the original questions posed by the users. As illustrated by the examples in Table 3, a reformulation could be answerable but not relevant, or relevant but not answerable. A successful reformulation must satisfy both conditions. There are multiple equally good reformulations for each question. For instance, when given the original question, it is hard to determine which of the first two reformulations in Table 3 would be closer to the user’s intent. As a result, we opt for a reference-free evaluation approach, where the evaluation does not rely on gold reformulations.

Measuring the relatedness of two questions involves determining how closely their topics, intents, and meanings are aligned. With these goals in mind, we propose three reference-free relevance metrics: edit distance, entity overlap ratio⁹, and cosine similarity between the original question and

⁹We justify our choice of entity overlap ratio as a relevance metric with human evaluation in Section 7.

		SQuADv2	QA ²	BanditQA	BBC	Reddit	Yelp	Average
GPT-3.5	ZS	31.86	37.61	50.58	9.09	24.46	19.35	28.83
	ZS CoT	26.49	39.52	44.74	14.08	19.40	13.33	26.26
	FS	35.75	43.26	67.34	0.00	19.12	16.67	30.35
	FS CoT	<i>51.08</i>	37.90	<i>70.33</i>	<i>16.22</i>	<i>37.65</i>	<i>33.77</i>	<i>41.16</i>
GPT-4	ZS	85.47	67.40	71.24	51.93	61.82	58.74	66.10
	ZS CoT	84.97	68.55	67.33	52.02	55.51	61.65	65.01
	FS	77.34	64.17	80.48	52.07	65.95	60.74	66.79
	FS CoT	76.11	64.09	80.85	53.93	71.48	60.43	67.82
Llama2	ZS	21.29	35.26	44.51	24.44	15.49	24.32	27.55
	ZS CoT	27.76	37.39	48.22	19.61	19.23	27.16	29.89
	FS	55.36	54.24	66.77	27.91	42.11	45.83	48.70
	FS CoT	35.44	35.65	58.53	25.45	35.87	34.09	37.51
Mistral	ZS	38.30	30.43	46.61	30.23	26.14	17.39	31.52
	ZS CoT	31.87	26.5	47.39	29.27	21.48	21.92	29.74
	FS	41.98	34.13	47.45	38.83	42.74	40.38	40.92
	FS CoT	46.65	44.32	58.23	38.10	48.28	32.97	44.76
Zephyr	ZS	63.73	55.37	67.61	45.31	38.04	32.50	50.43
	ZS CoT	65.42	50.36	68.27	39.72	34.07	39.53	49.56
	FS	67.29	63.06	67.67	42.86	50.00	30.23	53.52
	FS CoT	64.25	71.10	68.61	42.03	49.61	43.40	56.50

Table 4: Comparing F1 scores for unanswerable question detection with different prompting methods using both proprietary and open-source models on COULDASK. The best method for each base model is italicized, and the best method across all base models is bolded.

the reformulation to indicate the level of relevance. We consider the reformulation to be unanswerable and have zero relatedness for the unanswerable questions that are not successfully detected by the system. To automatically evaluate (1), we train a Llama2-7B model on COULDASK to classify whether the reformulation is answerable or not. The classifier achieved 95% accuracy on a held-out validation set. We release the classifier on Hugging Face . For (2), we calculate the Levenshtein edit distance, use GPT-4 to tag entities for computing entity overlap ratios, and apply OpenAI embedding models to produce question embeddings for computing cosine similarities. Finally, a reformulation that is answerable but irrelevant, or vice versa, is not yet helpful. To consolidate the evaluation into a single unified score, using entity overlap ratios as an example, we assign a score of 1 to a reformulation if it is both answerable and has an entity overlap ratio of more than 50%; otherwise, we assign a score of 0. We refer to this binary score as *success rate*.

5 Experimental Setup

Models. We test a range of instruction-finetuned LLMs. For proprietary models, we consider GPT-3.5 (gpt-3.5-turbo-0125) and GPT-4 (gpt-4-0613). For open-source models, we consider a list of

7-billion-parameter models: Llama2¹⁰ (Touvron et al., 2023), Mistral¹¹ (Jiang et al., 2023), and Zephyr¹² (Tunstall et al., 2023).

Comparisons. We consider several prompting approaches: zero-shot (ZS) and few-shot (FS) prompting and ZS and FS chain-of-thought (CoT) prompting. We first prompt LLMs to determine whether the input question cannot be answered by the provided document. For ZS and FS prompting, we use the prompt provided by SurgeAI¹³, which explicitly instructs the model to not produce an answer if the answer cannot be determined from the document alone. For ZS and FS CoT prompting, we expand the aforementioned prompt by asking the model to think step by step to come up with a reason to explain and support its decision. Only questions determined to be unanswerable proceed to the question reformulation stage. In this stage, all methods are provided with their previous turns where the models determine the questions are unanswerable. ZS and FS prompting instruct the model to make minimum edits to the original question to make it answerable. ZS and FS CoT prompting

¹⁰huggingface.co/meta-llama/Llama-2-7b-chat-hf

¹¹huggingface.co/mistralai/Mistral-7B-Instruct-v0.2

¹²huggingface.co/HuggingFaceH4/zephyr-7b-beta

¹³SurgeAI prompt: “Answer the following from the above passage alone, and if you can’t determine the answer based on the passage, say that you don’t know the answer.”

		SQuADv2	QA ²	BanditQA	BBC	Reddit	Yelp	Average
GPT-3.5	ZS	3.55	6.88	5.71	1.69	4.42	9.80	5.34
	ZS CoT	3.55	3.64	4.35	1.69	2.65	3.92	3.30
	FS	3.35	8.91	6.25	0.00	0.88	5.88	4.21
	FS CoT	5.52	8.50	7.20	10.17	3.54	7.84	7.13
GPT-4	ZS	17.16	23.89	11.68	52.54	8.85	43.14	26.21
	ZS CoT	15.78	18.22	8.29	35.59	14.16	35.29	21.22
	FS	10.45	17.00	10.19	42.37	8.85	31.37	20.04
	FS CoT	9.86	17.81	10.60	44.07	7.96	17.65	17.99
Llama2	ZS	2.56	6.88	5.30	13.56	4.42	11.76	7.42
	ZS CoT	3.75	3.24	2.99	10.17	1.77	9.80	5.29
	FS	7.10	14.17	8.15	16.95	8.85	17.65	12.14
	FS CoT	3.16	4.45	5.84	13.56	3.54	5.88	6.07
Mistral	ZS	1.58	2.02	4.35	11.86	1.77	1.96	3.92
	ZS CoT	2.96	1.21	4.08	6.78	0.88	7.84	3.96
	FS	3.35	4.05	3.67	22.03	6.19	11.76	8.51
	FS CoT	5.92	4.05	6.39	16.95	5.31	13.73	8.72
Zephyr	ZS	7.89	12.15	7.61	30.51	4.42	9.80	12.06
	ZS CoT	8.09	7.69	7.34	20.34	6.19	7.84	9.58
	FS	5.33	15.38	7.07	13.56	4.42	7.84	8.93
	FS CoT	8.88	13.77	10.19	35.59	7.96	13.73	15.02

Table 5: Question Reformulation. Success rates using different prompting methods with both proprietary and open-source models on COULDAK. The best method for each base model is italicized, and the best method across all base models is bolded.

instruct the model to think step by step to reason about the minimum edits they can make.

6 Results

Detecting Unanswerable Questions Being able to detect unanswerable questions is a necessary precondition for successful question reformulation. Table 4 presents the F1 scores for unanswerable question detection using different prompting approaches with each base model. Performance is often better for existing datasets than new ones, which indicates our approach’s effectiveness in generating more challenging questions. Among all models, GPT-4 performs best at identifying unanswerable questions. Surprisingly, both Mistral and Zephyr are more accurate at detecting unanswerable questions than GPT-3.5. Among all prompting techniques, FS CoT consistently improves upon ZS, with a larger degree of improvement observed in smaller models compared to larger ones.

Reformulating Questions Table 5 presents success rates for question reformulation vary dramatically from domain to domain. News (BBC) appears to be the least challenging domain, with success rates ranging from 10.17 to 52.54 depending on base models. Reddit is a challenging domain, with success rates ranging from 4.42 to 14.16. The results on Wikipedia are mixed. While SQuADv2, QA², and BanditQA are in the Wikipedia domain,

LLMs achieve the lowest success rates on BanditQA. We hypothesize that user queries written during interaction require deeper revision.

Among all base models, GPT-4 achieves the highest average success rate (26.21), while GPT-3.5 has the lowest average success rate (7.13). Among all open-source models, Zephyr has the best performance. When it comes to prompting methods, there is not a clear winner. Different LLMs can be improved with different prompting approaches. GPT-3.5, Mistral, and Zephyr benefit the most from FS CoT prompting, GPT-4 from ZS, Llama2 from FS. We include individual metrics for question reformulation in Appendix B.

7 Analysis

Qualitative analysis We randomly sample and analyze 20 reformulated questions from each base model (a total of 100 questions) that cannot be answered by the corresponding documents. We summarize the results in Table 6. We identify three major types of errors. The most frequent type is that the models simply rephrase or generate the same questions. Most of the errors in this type are contributed by open-source models such as Llama2.

Another type of error that occurs 14% of the time is that the models generate a question by copying a document span that looks similar to the original question. For example, given the original question

Error Category	%	Original question	Revised question
Simply rephrasing / producing the same questions	62	When was he born?	When was Jay Chou born and raised in Taipei, Taiwan?
		Which team won the most NFL Europe titles?	Who was the team that won the most NFL Europe titles?
Producing questions by copying documents	14	When did the Cubs play their first home game at Wrigley Field?	When did the Cubs move to Wrigley Field?
		what are the characteristics of the Classical Roman architecture?	What are some of the characteristics of Classical Roman architecture that were adopted by Renaissance architects?
Producing irrelevant questions	17	What is Snow Patrol’s most famous song?	Who are some famous musicians that have expressed admiration for Snow Patrol?
		which campus is the biggest in University of California?	Which college in the University of California system is not controlled by the Regents or President?
Other	11	How many arms does Krishna have?	I apologize for the confusion. The passage does not mention how many arms Krishna has.

Table 6: Error analysis on 100 revised questions generated by different models.

about the characteristics of the Classical Roman architecture, the model identifies a span that is similar to the original question and generates a question based on this span. However, the model does not plan thoroughly about whether the new question can be answered or not.

For 17% of the time, the models propose questions that are irrelevant to the original questions. For example, in the Snow Patrol question, the subject of the reformulation has entirely changed from asking about the most famous song to famous musicians. Finally, we group the remaining eleven errors in the Other type, where there usually is not even a question produced.

Limited benefit from FS CoT on Question Reformulation. We explore why FS CoT improves the detection of unanswerable questions but not question reformulation by conducting a qualitative analysis of GPT-4 outputs on a BanditQA example, as illustrated in Figure 2. It is relatively straightforward to determine why a question is unanswerable — either the document does not provide the necessary information or there is a presupposition that conflicts with the document. However, question reformulation demands compositional reasoning. The model needs to first decide on a reformulation strategy and then plan the specific steps to achieve the reformulation. The reformulated question generated by GPT-4 via FS CoT prompting is closer to being answerable. However, the model misses a subtle detail — the document only mentions the exercise habits of 30.8% of residents, not the general population. Therefore, FS CoT alone, without further methodological innovation, does not fully

Document:
In December 2007, Sport England published a survey which revealed that residents of St Albans were the 10th most active in England in sports and other fitness activities. 30.8% of the population participate at least 3 times a week for 90 minutes.

Question:
What is St Albans Sport?

ZS CoT for unanswerable question detection:
The article did not provide information on what St Albans Sport is, so I don’t know.

ZS CoT for question reformulation:
The article discusses the activity level of the residents of St Albans in sports and fitness but does not detail a specific sport. Instead of asking "What is St Albans Sport?" we could ask, "How often do residents of St Albans participate in sports and fitness activities?" to get an answer from the given passage.

ZS for question reformulation:
What is the level of participation in sports and other fitness activities in St Albans?

Figure 2: Example outputs produced by GPT-4. Via prompting, the model detects unanswerable questions then reformulates the questions with a second prompt.

address the challenge of question reformulation.

Compositional Modifications vs. Answerability. We hypothesize that it is more difficult to reformulate an unanswerable question to be answerable when it requires compositional modifications, which means making global edits to the question instead of making local edits.

To quantify edits required, we follow Lee et al. (2020) and annotate the minimum span of a question to explain why the question cannot be answered by the document. We use GPT-4 to annotate unanswerable questions in BanditQA, QA², and Yelp. We divide the questions into two categories.

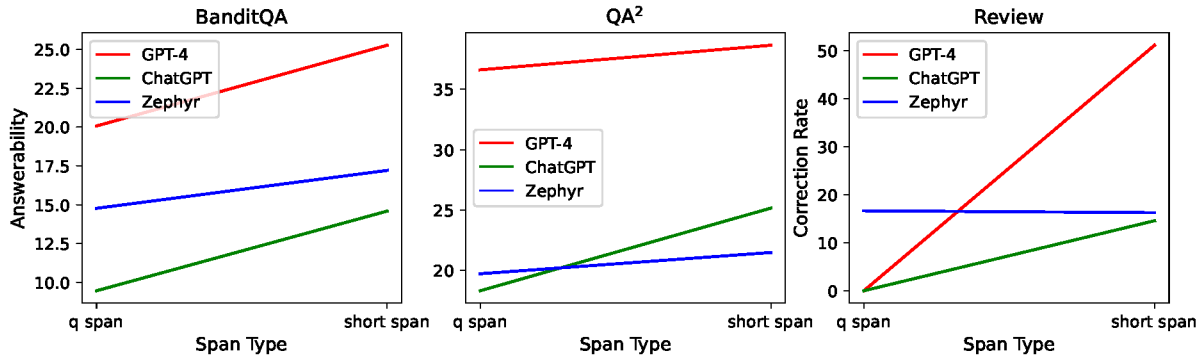


Figure 3: The relation between the percentage of reformulations answerable by the documents and the type of minimum spans in the original questions. q span are the examples where the minimum span is the full question, while short span includes all other instances.

q span	
Original Q	What medicine is made from Coca?
Reformulated Q	What substance is derived from the Coca leaf?
short span	
Original Q	When was USB-C developed?
Reformulated Q	When was USB developed?

Table 7: Examples of a question with the minimum span being the entire question (q span) and a question where the minimum span is a noun (short span).

The first category comprises questions whose minimum span is the entire question. We call these *q span*. The second category covers questions that have a shorter minimum span, typically a noun or a noun phrase. We call these *short span*. We calculate how many reformulations are answerable, broken down by the type of minimum span. We consider the reformulations generated by zero-shot prompting with GPT-4, GPT-3.5, and Zephyr.

The results of this analysis are in Figure 3. Our findings consistently show that fewer reformulations are answerable when the minimum span constitutes the entire question compared to when the minimum span is shorter, with only one exception in zero-shot prompting with Zephyr on Yelp.

We present examples for each span type in Table 7. When the minimum span is the entire question, we cannot attribute the presupposition error to a smaller segment of the question. As a result, the modification has to be applied to the full question. For example, for the q-span example in Table 7, modifying individual parts of the question cannot correct the presupposition error, and therefore the model needs to make global changes. On the contrary, when the minimum span is short, it only re-

	Edit Distance	Entity Overlap
Fleiss' κ	40.96	93.99

Table 8: The correlation between the tested metrics and human-judged relevance between two questions, evaluated using Fleiss' κ .

quires local edits. For the short-span question in Table 7, replacing USB-C with USB is enough to correct the question. Future efforts should be more devoted to the more challenging cases to make progress on question reformulation.

Sufficiency of entity overlap ratios. Our metric for relevance is based on a minimum entity overlaps. To judge the sufficiency of this metric, we compare it to human evaluation. As a baseline we consider Levenshtein edit distance, a method to measure the similarity between two questions, as a baseline metric. We have a human annotator evaluate 200 pairs of reformulated questions produced by zero-shot prompting with GPT-4 and GPT-3.5 on BanditQA. The question pairs are randomly shuffled to ensure the annotator remains unaware of the model source of each reformulation. The annotator then selects the more relevant question from each pair, based on alignment with the original question's intent. Subsequently, we identify the question with a higher entity overlap ratio and the question with a lower edit distance as the more relevant ones, respectively.

We calculate the Fleiss' κ score between human evaluations and each of the metrics, with the results summarized in Table 8. Compared to edit distance, the entity overlap ratio more accurately represents the relevance between the original and reformulated questions (93.99 vs. 40.96). A Fleiss' κ score

of 93.99 also suggests a near-perfect agreement between the entity overlap ratio and human judgments. We hypothesize that the specific semantic properties of the questions are more central than the surface form represented by edit distance.

8 Conclusion

Users often ask unanswerable questions when they seek information from unfamiliar documents. Existing LLMs identify these questions but do not aid users in reformulating new questions, resulting in ineffective user-LLM interactions. We introduce COULDASK, a compiled set of grounded-document QA datasets designed to study both unanswerable question detection and question reformulation.

Limitations

While our benchmark offers advantages over existing sources, we acknowledge the following limitations. Questions in BBC, Reddit, and Yelp are generated by GPT-4, and they may not accurately represent questions posed by humans. Despite best efforts to ensure high-quality annotations, occasional human errors are possible. Additionally, our benchmark only collects English questions and thus lacks language diversity. Finally, regarding evaluation, the way we currently measure success rates only focuses on mistakes made on unanswerable questions. If an answerable question is detected to be unanswerable, we do not evaluate question reformulation in such cases.

Acknowledgments

This work was supported by NSF CAREER #2037519 and NSF #1901030. We also thank Lillian Lee and the anonymous reviewers for their helpful feedback.

References

- Monica Agrawal, Stefan Hegselmann, Hunter Lang, Yoon Kim, and David Sontag. 2022. [Large language models are few-shot clinical information extractors](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 1998–2022, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Pedro Faustini, Zhiyu Chen, Besnik Fetahu, Oleg Rokhlenko, and Shervin Malmasi. 2023. [Answering unanswered questions through semantic reformulations in spoken QA](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 5: Industry Track)*, pages 729–743, Toronto, Canada. Association for Computational Linguistics.
- Ge Gao, Hung-Ting Chen, Yoav Artzi, and Eunsol Choi. 2023. [Continually improving extractive QA via human feedback](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 406–423, Singapore. Association for Computational Linguistics.
- Neel Guha, Julian Nyarko, Daniel E. Ho, Christopher Re, Adam Chilton, Aditya Narayana, Alex Chohlas-Wood, Austin Peters, Brandon Waldon, Daniel Rockmore, Diego Zambrano, Dmitry Talisman, Enam Hoque, Faiz Surani, Frank Fagan, Galit Sarfaty, Gregory M. Dickinson, Haggai Porat, Jason Hegland, Jessica Wu, Joe Nudell, Joel Niklaus, John J Nay, Jonathan H. Choi, Kevin Tobia, Margaret Hagan, Megan Ma, Michael Livermore, Nikon Rasumov-Rahe, Nils Holzenberger, Noam Kolt, Peter Henderson, Sean Rehaag, Sharad Goel, Shang Gao, Spencer Williams, Sunny Gandhi, Tom Zur, Varun Iyer, and Zehua Li. 2023. [Legalbench: A collaboratively built benchmark for measuring legal reasoning in large language models](#). In *Thirty-seventh Conference on Neural Information Processing Systems Datasets and Benchmarks Track*.
- Shengding Hu, Yifan Luo, Huadong Wang, Xingyi Cheng, Zhiyuan Liu, and Maosong Sun. 2023. [Won’t get fooled again: Answering questions with false premises](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5626–5643, Toronto, Canada. Association for Computational Linguistics.
- Albert Q Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, et al. 2023. Mistral 7b. *arXiv preprint arXiv:2310.06825*.
- Najoung Kim, Phu Mon Htut, Samuel R. Bowman, and Jackson Petty. 2023. [\(QA\)²: Question answering with questionable assumptions](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 8466–8487, Toronto, Canada. Association for Computational Linguistics.
- Najoung Kim, Ellie Pavlick, Burcu Karagol Ayan, and Deepak Ramachandran. 2021. [Which linguist invented the lightbulb? presupposition verification for question-answering](#). 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 3932–3945, Online. Association for Computational Linguistics.
- Gyeongbok Lee, Seung-won Hwang, and Hyunsouk Cho. 2020. [SQuAD2-CR: Semi-supervised annotation for cause and rationales for unanswerability](#).

- in SQuAD 2.0. In *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 5425–5432, Marseille, France. European Language Resources Association.
- Bodhisattwa Prasad Majumder, Sudha Rao, Michel Galley, and Julian McAuley. 2021. [Ask what’s missing and what’s useful: Improving clarification question generation using global knowledge](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 4300–4312, Online. Association for Computational Linguistics.
- Sewon Min, Julian Michael, Hannaneh Hajishirzi, and Luke Zettlemoyer. 2020. [AmbigQA: Answering ambiguous open-domain questions](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 5783–5797, Online. Association for Computational Linguistics.
- Valentina Pyatkin, Jena D. Hwang, Vivek Srikumar, Ximing Lu, Liwei Jiang, Yejin Choi, and Chandra Bhagavatula. 2023. [ClarifyDelphi: Reinforced clarification questions with defeasibility rewards for social and moral situations](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 11253–11271, Toronto, Canada. Association for Computational Linguistics.
- Pranav Rajpurkar, Robin Jia, and Percy Liang. 2018. Know what you don’t know: Unanswerable questions for squad. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 784–789.
- Sudha Rao and Hal Daumé III. 2018. [Learning to ask good questions: Ranking clarification questions using neural expected value of perfect information](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2737–2746, Melbourne, Australia. Association for Computational Linguistics.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.
- Son Quoc Tran, Gia-Huy Do, Phong Nguyen-Thuan Do, Matt Kretchmar, and Xinya Du. 2023. Agent: A novel pipeline for automatically creating unanswerable questions. *arXiv preprint arXiv:2309.05103*.
- Lewis Tunstall, Edward Beeching, Nathan Lambert, Nazneen Rajani, Kashif Rasul, Younes Belkada, Shengyi Huang, Leandro von Werra, Cl  mentine Fourrier, Nathan Habib, et al. 2023. Zephyr: Direct distillation of lm alignment. *arXiv preprint arXiv:2310.16944*.
- Julia White, Gabriel Poesia, Robert Hawkins, Dorsa Sadigh, and Noah Goodman. 2021. [Open-domain clarification question generation without question examples](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 563–570, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Xinyan Yu, Sewon Min, Luke Zettlemoyer, and Hannaneh Hajishirzi. 2023. [CREPE: Open-domain question answering with false presuppositions](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 10457–10480, Toronto, Canada. Association for Computational Linguistics.

A Dataset Details

We provide more details about the existing datasets we include in COULDASK.

SQuADv2 is a collection of QA pairs from Wikipedia articles, including questions with presupposition errors. Such questions are written by crowdworkers who are instructed to craft questions that are plausible but have presuppositions that cannot be verified by the associated articles. These annotators either modify valid questions or create new questions based on entities or topics related to the text, ensuring the questions appear relevant and deceptively valid.

QA² contains QA pairs from general web articles. QA² sources questions from Google’s autocomplete API. It calls the API on prefix strings – when, where, which, how, what, why, who, whose – to complete the queries. Crowdworkers are recruited to (1) search for relevant documents and (2) check whether these questions contain presupposition errors. QA² features naturally occurring presupposition errors and diverse document sources. The original QA² dataset is open-domain QA and only provides annotations for relevant URLs. To obtain the context, we scrape the text from their annotated URL and remove the examples where the contexts are not clear.

BanditQA investigate whether continual human feedback can improve extractive QA systems; we repurpose the dataset to study naturally occurring presupposition errors. During data collection, users are presented with Wikipedia passages, and they are required to write questions about things they are curious to know. In a later verification step, over 30% of these questions are identified to have presupposition errors. BanditQA is closest to the

	SQuADv2	QA ²	BanditQA	BBC	Reddit	Yelp	Average
Prompting	18.74	21.86	10.60	47.46	7.96	33.33	23.33
Explicit Prompting	11.43	8.91	4.07	18.64	7.07	17.64	11.00
Rule-based	4.50	2.43	1.22	11.86	0.88	3.90	4.13

Table 9: Comparing the standard zero-shot prompting method to various baselines to hack the proposed metric, success rates.

		SQuADv2	QA ²	BanditQA	BBC	Reddit	Yelp	Average
Answerability	Prompting	28.01	36.84	36.84	71.19	13.27	47.06	36.58
	Rule-based	28.01	35.22	24.86	62.71	10.61	52.94	35.73
Entity Overlap Ratio	Prompting	38.64	24.87	29.13	45.00	44.73	48.69	38.51
	Rule-based	21.36	13.60	13.15	23.50	15.55	21.41	18.10

Table 10: Analysis of individual metrics used to compute success rates.

setting of this study – when users are unfamiliar with the documents, they come up with presuppositions that cannot be grounded in the documents.

B More Analysis

Hacking success rates. We explore the potential for our proposed metrics to be manipulated in order to artificially achieve perfect scores. We consider two approaches. The first approach directly prompts LLMs to produce an answerable question while leaving all the entities unchanged. In the second approach, we test a rule-based heuristic that has the following steps: (1) tag entities in the original question via LLM prompting, (2) select the sentence in the document that has the highest entity overlap ratio with the original question, and (3) replace something in the select sentence to create a wh-question, again using LLM prompting. The results, summarized in Table 9, indicate that our success rates remain robust, even when faced with methods explicitly designed to exploit our metric.

We seek to understand where performance degrades in a rule-based heuristic by inspecting individual metrics used to compute success rates. We present the results in Table 10. While using the rule-based heuristic leads to similar answerability performance, the entity overlap ratio drops significantly. We show further qualitative analysis of the error types in Figure 4.

Individual Metrics in Question Reformulation

We present individual metrics for question reformulation. Table 13 presents the answerability (top) and entity overlap ratios (bottom) achieved by each method using each base model. We additionally report cosine similarities in Table 14 and edit distance in Table 15.

	Precision	Recall
GPT-4	99	94

Table 11: Performance of applying GPT-4 to tag entities.

	BBC	Reddit	Yelp
Fleiss’ κ	65.07	59.94	58.76

Table 12: Inter-annotator agreement rates among three workers for annotating unanswerable questions.

Assessing the Accuracy of Entity Tagging Models through Human Evaluation

To determine the entity overlap ratio, we first utilize GPT-4 to identify entities within both the original and revised questions. To ensure the credibility of the tagged entities, we conduct a human evaluation on a set of 100 questions. The findings of this evaluation are shown in Table 11. This analysis reveals that GPT-4 exhibits both high precision and recall in the identification of entities, affirming its effectiveness for this task.

C More Annotation Details

Annotation Guidelines We present annotation guideline for annotating unanswerable questions in Figure 5. The crowdworkers are those who were identified to contribute high-quality annotations from our previous annotation tasks. For every completed HIT, we pay the crowdworker USD 0.5.

Annotation Agreement We report inter-annotator agreement rates between crowdworkers in Table 12. Specifically, we compute how often all three workers agree with each other. The agreement rates on the three datasets are close to

		SQuADv2	QA ²	BanditQA	BBC	Reddit	Yelp	Average
Answerability								
GPT-3.5	ZS	4.73	12.55	8.42	5.08	4.42	9.80	7.50
	ZS CoT	5.92	10.93	8.29	5.08	2.65	5.88	6.46
	FS	7.30	18.22	14.54	0.00	0.88	5.88	7.80
	FS CoT	9.07	16.19	15.35	10.17	3.54	11.76	11.02
GPT-4	ZS	26.82	36.84	22.96	72.88	8.85	58.82	37.86
	ZS CoT	27.22	34.82	19.16	66.10	14.16	52.94	35.73
	FS	21.10	34.01	25.41	66.10	8.85	47.06	33.76
	FS CoT	21.30	34.82	24.18	71.19	7.96	52.94	35.40
Llama2	ZS	4.73	12.55	8.15	18.64	4.42	11.76	10.05
	ZS CoT	5.13	5.67	5.30	13.56	1.77	9.80	6.87
	FS	13.41	25.51	12.36	28.81	10.62	23.53	19.04
	FS CoT	5.13	8.50	8.42	15.25	5.31	7.84	8.41
Mistral	ZS	3.55	6.48	6.39	15.25	2.65	3.92	6.37
	ZS CoT	5.13	2.02	7.74	11.86	2.65	7.84	6.21
	FS	6.90	8.50	7.47	28.81	6.19	19.61	12.92
	FS CoT	10.85	10.12	9.92	23.73	5.31	19.61	13.26
Zephyr	ZS	15.78	23.89	12.77	42.37	4.42	11.76	18.50
	ZS CoT	14.20	12.96	12.23	33.90	6.19	11.76	15.21
	FS	20.32	34.82	16.98	38.98	4.42	15.69	21.87
	FS CoT	15.38	25.91	15.49	38.98	7.96	21.57	20.88
Entity Overlap Ratios								
GPT-3.5	ZS	10.44	10.56	19.82	2.37	11.28	8.50	10.50
	ZS CoT	7.27	5.94	11.39	1.81	7.82	3.27	6.25
	FS	6.75	10.32	23.38	0.00	5.68	9.15	9.21
	FS CoT	14.97	8.03	28.17	7.91	16.74	15.20	15.17
GPT-4	ZS	38.73	24.93	31.26	49.49	45.13	48.86	39.73
	ZS CoT	29.99	20.28	23.73	34.97	32.98	38.30	30.04
	FS	26.13	18.42	31.66	43.98	45.84	47.32	35.56
	FS CoT	24.85	18.76	33.93	41.27	48.27	39.51	34.43
Llama2	ZS	5.78	8.40	18.07	12.43	6.98	12.25	10.65
	ZS CoT	6.53	5.16	13.45	7.37	5.84	10.62	8.16
	FS	21.09	17.88	35.15	14.58	30.22	31.86	25.13
	FS CoT	8.16	7.42	17.53	12.74	10.34	12.42	11.43
Mistral	ZS	6.85	4.93	10.73	10.88	7.67	4.41	7.58
	ZS CoT	6.27	2.63	11.38	6.10	4.57	9.31	6.71
	FS	8.88	5.74	12.99	19.07	26.55	26.63	16.64
	FS CoT	10.46	6.07	16.66	15.37	22.05	20.26	15.15
Zephyr	ZS	21.25	14.24	32.94	26.50	17.92	16.01	21.48
	ZS CoT	17.13	11.57	23.01	19.58	14.90	17.32	17.25
	FS	12.19	14.88	28.60	12.43	18.07	12.81	16.50
	FS CoT	21.37	16.94	35.29	30.42	33.26	25.98	27.21

Table 13: Question reformulation quality broken down by answerability (top) and entity overlap ratios (bottom) on COULDASK. The best method for each base model is highlighted with an underscore, and the best method across all base models is bolded.

		SQuADv2	QA ²	BanditQA	BBC	Reddit	Yelp	Average
GPT-3.5	ZS	14.33	18.98	26.24	4.94	13.55	10.74	14.80
	ZS CoT	7.27	11.18	14.30	3.60	7.52	3.97	7.97
	FS	6.75	20.46	33.78	0.00	8.66	9.47	13.19
	FS CoT	14.97	16.76	38.43	8.86	21.90	20.22	20.19
GPT-4	ZS	57.92	42.45	44.52	64.72	55.18	59.91	54.12
	ZS CoT	44.51	33.66	31.64	51.80	32.54	42.49	39.44
	FS	43.98	35.28	47.66	59.4	54.18	53.71	49.03
	FS CoT	42.31	34.43	48.84	60.27	59.93	52.49	49.71
Llama2	ZS	8.26	16.42	23.30	17.22	8.66	15.94	14.97
	ZS CoT	7.90	8.62	14.60	9.31	6.41	10.61	9.57
	FS	33.02	35.71	51.55	28.46	38.01	38.38	37.52
	FS CoT	9.78	11.18	19.66	13.01	11.10	14.16	13.15
Mistral	ZS	10.62	10.12	15.42	16.88	10.71	8.19	11.99
	ZS CoT	7.55	5.20	12.60	8.36	5.63	8.64	8.00
	FS	14.34	12.71	19.56	28.50	36.13	32.55	23.97
	FS CoT	12.64	12.93	19.64	19.72	22.86	20.50	18.05
Zephyr	ZS	31.48	27.49	42.97	38.23	22.54	20.85	30.59
	ZS CoT	21.60	16.43	26.66	26.55	16.45	19.63	21.22
	FS	27.89	32.76	39.19	25.91	26.50	16.05	28.05
	FS CoT	30.72	28.60	42.91	35.22	42.45	33.67	35.59

Table 14: Cosine similarity using different prompting methods with both proprietary and open-source models on COULDAK. The best method for each base model is highlighted with an underscore, and the best method across all base models is bolded.

each other.

Failure case 1: Not all entities present in the same sentence
Original unanswerable question: What are balance in an open system of particles?
Entities: balance, open system, particles
Most-overlapped sentence in the document: This means that in a closed system of particles, there are no internal forces that are unbalanced.
Reformulated Question: What does this mean about internal forces in a closed system of particles?
Note: The three entities are mentioned in different sentences. Therefore, the most overlapped sentence only has one overlapped entity. Although this question can be answered, the reformulation is not successful due to the low entity overlap ratio. Note that most of our questions have an average number of two to three entities.

Failure case 2: Reformulated questions are still unanswerable
Original unanswerable question: What job requires no qualifications?
Entities: job, qualifications
Most-overlapped sentence in the document: As in the House of Commons, a number of qualifications apply to being an MSP.
Reformulated Question: What qualifications apply to being an MSP as in the House of Commons?
Note: Even though a part of the sentence was replaced with a what question, this process does not necessarily make the reformulated question become answerable. The article does not mention what specific qualifications, and successful reformulations require a deeper understanding of the document as a whole.

Figure 4: Qualitative analysis on the rule-based heuristic approach to hack our proposed metric, success rates.

		SQuADv2	QA ²	BanditQA	BBC	Reddit	Yelp	Average
GPT-3.5	ZS	1.23	2.04	2.37	0.24	0.95	0.51	1.22
	ZS CoT	4.88	11.70	11.58	3.07	3.35	6.25	6.81
	FS	1.79	2.90	4.24	-	0.94	0.37	1.71
	FS CoT	3.12	2.77	5.57	1.22	2.66	2.20	2.92
GPT-4	ZS	4.09	4.80	5.63	6.63	5.69	6.47	5.55
	ZS CoT	13.64	16.77	16.72	15.92	18.76	22.27	17.35
	FS	4.18	4.53	6.31	6.46	7.21	7.78	6.08
	FS CoT	5.05	5.86	7.47	8.41	9.98	8.29	7.51
Llama2	ZS	2.24	4.87	5.40	1.47	0.82	1.02	2.64
	ZS CoT	6.61	13.05	15.10	7.08	4.70	9.33	9.31
	FS	2.53	3.19	2.66	1.54	2.13	2.37	2.40
	FS CoT	9.02	10.27	18.04	8.29	12.09	9.55	11.21
Mistral	ZS	5.93	5.20	7.85	3.29	2.19	2.18	4.44
	ZS CoT	11.20	9.17	20.88	11.39	9.21	8.88	11.79
	FS	6.05	4.23	6.00	4.29	4.69	5.27	5.09
	FS CoT	16.05	17.85	25.98	20.17	24.45	15.49	20.00
Zephyr	ZS	11.53	12.28	13.04	10.93	6.22	3.41	9.57
	ZS CoT	40.03	29.28	47.48	30.58	15.26	15.53	29.69
	FS	8.32	8.64	9.70	7.00	6.06	3.84	7.26
	FS CoT	14.95	25.63	22.88	16.14	13.13	9.90	17.11

Table 15: Edit distance using different prompting methods with both proprietary and open-source models on COULDASK. The best method for each base model is highlighted with an underscore, and the best method across all base models is bolded.

Instructions (click to expand/collapse)

Thanks for participating in this HIT! Please read the instructions carefully.

You will be given a **question** and an associated **document**.

Your task is to determine whether the **question** is answered with the given **document**:

- When the **question** is answerable, we need you to select the span that you used to answer the **question**.
- Please read our [example google doc](#) to learn how to select an answer span.
- If a word-to-word answer cannot be found within the **document**, then the **question** is considered unanswerable.

More Instructions:

- If the **question** is not sensible (e.g., it is not even asking for any information) or contains multiple sub-questions (when was person X born and when was X married?), please select unanswerable and check "The document or the question is difficult to interpret."
- You are not required to read the entire **document**, it may be faster to just search for the keywords to find potential answers :)
- We will both manually and automatically review your response. Please complete the HIT carefully. Workers who make too many mistakes will be de-qualified. Thanks!

Question:
\${question}

Document:
\${context}

Answerability: Can you answer the **question** by finding a span in the given **document**?

☐ Yes ☐ No

If you answered Yes, please select a span of the document to answer the question.
Your **highlighted span**:

REMEMBER: **Make sure your answer ...**
addresses the question directly
is as concise as possible

☐ The document or the question is difficult to interpret (I don't understand some of the elements).

Figure 5: Annotation guideline for crowdworkers to annotate unanswerable questions.