

LLM-AS-AN-INTERVIEWER: Beyond Static Testing Through Dynamic LLM Evaluation

Anonymous ACL submission

Abstract

We introduce **LLM-AS-AN-INTERVIEWER**, a novel paradigm for evaluating large language models (LLMs). This approach leverages multi-turn interactions where the LLM interviewer actively provides feedback on responses and poses follow-up questions to the evaluated LLM. At the start of the interview, the LLM interviewer dynamically modifies datasets to generate initial questions, mitigating data contamination. We apply the **LLM-as-an-Interviewer** framework to evaluate six models on the reasoning, factuality and instruction-following tasks. Our results show that the framework effectively provides insights into LLM performance, including the quality of initial responses, adaptability to feedback, and ability to address follow-up queries like clarification or additional knowledge requests. The framework also addresses key limitations of conventional methods like LLM-as-a-Judge, including verbosity bias and inconsistency across runs. Finally, we propose the **Interview Report**, which aggregates insights from the interview process, providing examples and a comprehensive analysis of the LLM’s strengths and weaknesses. This report offers a detailed snapshot of the model’s real-world applicability¹.

1 Introduction

With large language models (LLMs) becoming increasingly proficient in generating fluent free-form responses, it has become crucial to properly assess their capabilities and limitations (Liang et al., 2022; Chang et al., 2024). Recently, LLM-as-a-judge has emerged as a promising framework for automatic free-form response evaluation. Compared to traditional lexical matching-based metrics (e.g., ROUGE (Lin, 2004), BLEU (Papineni et al., 2002)) or embedding-based metrics (e.g.,

BERTScore (Zhang et al., 2019)), previous works on LLM-as-a-Judge have reported higher correlations with human judgments (Chiang and Lee, 2023; Zheng et al., 2023a; Dubois et al., 2024b).

Despite its potential, the LLM-as-a-Judge framework faces several practical limitations that hinder its widespread adoption, primarily due to its *static nature* (Li et al., 2024a; Gu et al., 2024). First, using a *fixed set of test inputs* raises concerns about data contamination (Sainz et al., 2023; Zhou et al., 2023a; Oren et al., 2023b), where the evaluated models may achieve high scores on instances encountered during training. Second, *single-turn interactions* fail to thoroughly probe a model’s true comprehension (Li et al., 2019; Wang et al., 2023; Kwan et al., 2024). For instance, the judge model may assess the confined performance of LLMs distant from the use case, be influenced by superficial factors (e.g., favoring longer responses), and exhibit high variance across runs.

In this study, we propose **LLM-AS-AN-INTERVIEWER**, a new paradigm for evaluating LLMs. Inspired by human interviews, this approach starts with general questions but dynamically adapts by posing different types of questions based on the model’s responses. As shown in Figure 1, the LLM interviewer plays three key roles: (1) Question Modification, adapting benchmark datasets to generate diverse and challenging initial interview questions; (2) Providing Feedback, guiding the model to refine its responses; and (3) Generating Follow-up Questions, exploring related concepts through clarification requests or additional explanations. This dynamic evaluation reveals behaviors that static benchmarks cannot capture, such as the model’s ability to improve through feedback and provide more detailed explanations.

We demonstrate the efficacy of our framework through experiments on three tasks: Reasoning, Factuality, and Instruction Following. In §5, we

¹The code for our framework is publicly available at <https://anonymous.4open.science/r/interview-eval-46C6/>

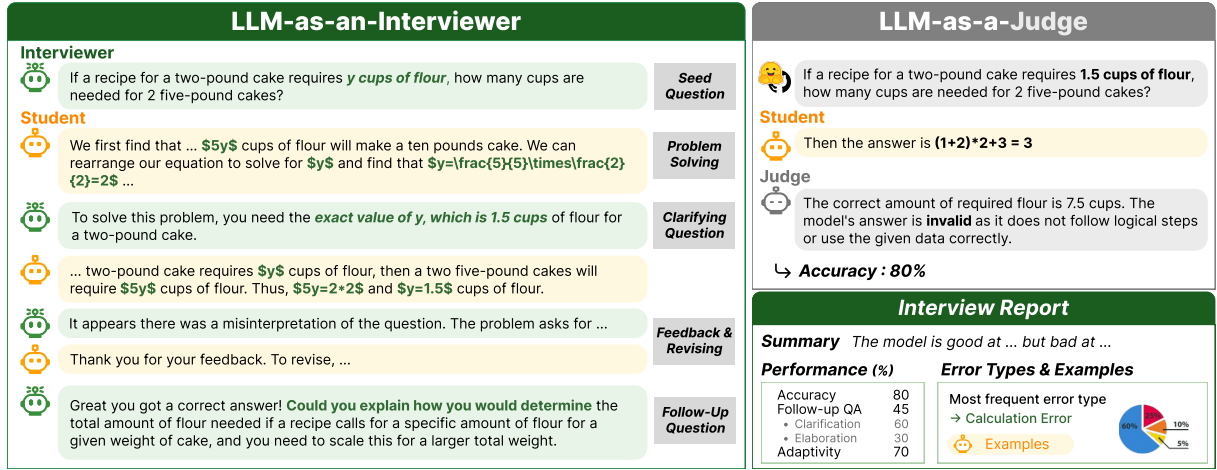


Figure 1: **Overview of LLM-as-an-Interviewer.** *LLM-as-an-Interviewer* aims to assess LLMs by 1) constructing seed questions of interviews based on existing benchmark datasets, and 2) providing feedback or asking additional follow-up questions. Unlike evaluation being limited to a single score with LLM-as-a-Judge, the *Interview Report* provided by our framework offers a snapshot of what the model excels at and where it falls short, along with scores for various abilities and examples.

evaluate six different models using GPT-4o as the interviewer. We demonstrate the impact of providing feedback on the model’s performance and suggest that follow-up questions can offer deeper insights into its behavior. These follow-ups can help uncover failure reasons or assess performance on additional requests. Furthermore, in § 6 and § 7, we show that our interactive and dynamic framework addresses key limitations of static evaluations, such as data contamination, verbosity bias in LLM evaluators, and high variance across multiple runs.

Moreover, we find that the extended interaction between the LLM interviewer and interviewee contains rich information that reveals the LLM interviewee’s strengths and limitations. Based on this, we introduce a new evaluation protocol, **Interview Report**, which summarizes the interaction into a structured format.

Our contributions are as follows:

- Introduce **LLM-as-an-Interviewer**, a novel evaluation paradigm that mimics the dynamic nature of how humans evaluate humans.
- Demonstrate that **LLM-as-an-Interviewer** mitigates several limitations of traditional evaluation approaches, such as verbosity bias, data contamination, and high variance.
- Propose the **Interview Report**, offering a detailed snapshot of an LLM’s capabilities, including examples and its common errors.

2 Related Works

2.1 LLM-based Evaluation

LLM-as-a-Judge is widely used for text evaluation, providing more human-aligned assessments than traditional methods like BLEU and ROUGE, while addressing scalability and consistency issues in human-based approaches (Gu et al., 2024). However, LLM-as-a-Judge faces reliability issues, including self-enhancement bias, sensitivity to response length and order, and low self-consistency. To enhance reliability, Li et al. (2023b) and Dubois et al. (2024b) mitigate biases related to ordering and length by adjusting positions or win rates. Additionally, works such as Wang et al. (2024b); Kim et al. (2024) train open-sourced critique models to tackle the inherent inconsistencies of closed LLMs, ensuring reproducibility.

LLM-based evaluation is also used to simulate multi-turn interactions between evaluators and models in certain tasks. For instance, Wang et al. (2023) evaluate LLMs by simulating user-model interactions in contexts where users provide feedback during tool usage. Similarly, Yu et al. (2024) simulate knowledge-focused dialogues in multiple interactions to mitigate data contamination issues. Li et al. (2023a) share a motivation similar to ours, as they are also inspired by human interviews. However, their approach is tailored to Conversational Question Answering tasks and emphasizes generating new questions during evaluation.

LLM-as-an-Interviewer is a generalized bench-

mark that mimics a multiple interview process to better assess a model’s capabilities. Our framework enables the simulation of diverse interactions, such as giving feedback and asking follow-up questions, which align much more closely with how humans use models in real scenarios.

2.2 Data contamination in LLMs

LLMs are trained on large corpora that can be contaminated with benchmark data (Dodge et al., 2021; Soldaini et al., 2024), undermining benchmark reliability (Zhou et al., 2023b). This has sparked interest in **contamination detection** (Margar and Schwartz, 2022). Shi et al. (2024) and Oren et al. (2023a) propose methods to estimate the likelihood of text in LLM pretraining data using only API access, as is common for frontier models (OpenAI et al., 2024). This has led to further research on contamination (Yax et al., 2024; Deng et al., 2024; Jiang et al., 2024; Dekoninck et al., 2024b; Xu et al., 2024), including methods to evade detection by training on rephrased benchmark data (Dekoninck et al., 2024a).

Overall, efforts to prevent contamination through detection have had limited success, prompting researchers to focus on **contamination mitigation** in benchmarks. This typically involves dynamic evaluation in two ways: (1) The *evaluation data is dynamic*, such as crawling new instances live (White et al., 2024; Jain et al., 2024) or generating new instances on the fly (Wang et al., 2024a; Li et al., 2024b); (2) The *evaluator is dynamic*, as in AlpacaEval (Li et al., 2023b) and MT-Bench (Zheng et al., 2023b), where the same samples are reused but evaluation depends on a dynamic LLM rating completions instead of ground-truth solutions. Some benchmarks, like Chatbot Arena (Chiang et al., 2024), combine dynamic data and evaluators, with users creating and rating responses on the fly.

LLM-as-an-Interviewer is a benchmark that mitigates contamination through both dynamic evaluation data and a dynamic evaluator. The LLM interviewer rephrases questions and proposes novel follow-ups while also evaluating the model’s responses and problem-solving process. This setup makes cheating via contamination difficult. Additionally, it enables fast, low-cost evaluation in minutes, unlike Chatbot Arena, which requires human access and takes days to yield statistically significant results at a higher cost.

3 LLM-AS-AN-INTERVIEWER

We introduce the LLM-AS-AN-INTERVIEWER², which simulates an interview where one LLM dynamically evaluates another. The LLM being evaluated is referred to as the **Interviewee**, while the LLM conducting the evaluation is referred to as the **Interviewer**. In this section, we describe (1) the overall interview framework (§ 3.1) and (2) the Interview Report, which summarizes and presents the results of the interview (§ 3.2).

3.1 Interview Framework

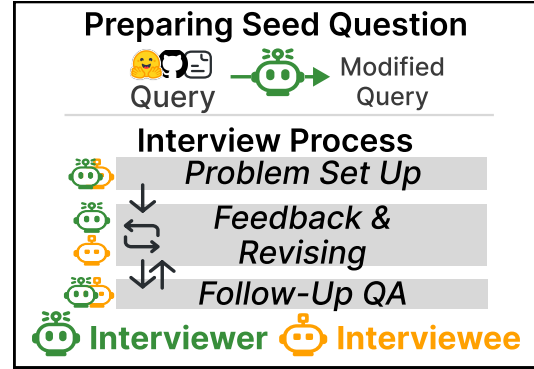


Figure 2: Workflow of LLM-as-an-Interviewer

We design an interview framework for LLMs as illustrated in Figure 2. The **Interviewer** plays three main roles throughout the interview process: (1) **Question Modification**, (2) **Providing Feedback**, and (3) **Generating Follow-up Questions**. This structure ensures a comprehensive evaluation of the **Interviewee**’s capabilities. These roles are **modular and pluggable**, allowing users to enable or disable specific roles based on their evaluation needs. Algorithm 1 provides the implementation of the interview process.

3.1.1 Question Modification

Before conducting the interview, the **Interviewer** prepares seed questions by **modifying queries from existing benchmark datasets**. This ensures diverse scenarios that align with established evaluation standards while preventing data contamination. This approach mirrors human interviewers, who often tweak existing questions or introduce new challenges. In our experiment, we employ two query modification strategies, prompting the model to modify queries. Each strategy is designed to enhance question diversity while main-

²We provide a general Python library applicable to various tasks. We delete the link due to an anonymous policy.

taining ease of real-time adjustments and comparable complexity. We provide a detailed description of these methods in Appendix C.1.

3.1.2 Providing Feedback

During the interview process, the **Interviewer** evaluates the response from the **Interviewee** and provides **constructive feedback**. Providing feedback aims to guide the **Interviewee** in identifying errors, refining the response, and improving solutions. This phase closely resembles real-world LLM interaction scenarios, where LLMs iteratively refine their responses based on user-generated feedback (Wang et al., 2023).

3.1.3 Generating Follow-up Questions

After assessing the response to the initial question, the **Interviewer** poses **follow-up questions** to further evaluate the **Interviewee**’s ability. These questions help gauge additional abilities that may not have been covered in the initial response. Follow-up questioning is crucial in many interview processes (e.g., Google coding interviews) and real-world LLM applications (Bai et al., 2024). Table 1 provides follow-up type examples for Reasoning, Knowledge, and Instruction Following task.

3.2 Interview Report

The Interview Report, generated as a result of the interview process, includes (1) Performance Scores, (2) Error Analysis and Examples, and (3) Comprehensive Summary. The real examples of Interview Report are in Appendix C.1.

(1) Performance Scores We obtain two main performance metrics through the interview process:

- **Problem Solving Ability** – Measures how effectively the model solves a given problem at its n -th interaction with the **Interviewer**. We define $\text{Score}_{\text{seed}}@n$ as the model’s performance at the n -th interaction, where an interaction involves receiving feedback and making revisions. Thus, $\text{Score}_{\text{seed}}@1$ corresponds to the LLM-as-a-Judge setting without feedback, while $\text{Score}_{\text{seed}}@n$ represents performance after $n - 1$ iterations of feedback and revision.
- **Follow-up Question Handling** – Evaluates how well the model responds to follow-up questions. To ensure fairness across models with varying performance levels, we cate-

gorize follow-up question accuracy into three types: accuracy within correctly answered questions, accuracy within incorrectly answered questions, and overall accuracy.

These metrics reflect abilities crucial in real-world interactions with LLMs, where responding effectively to feedback and follow-up questions significantly influences user satisfaction beyond just answering the initial query.

(2) Error Analysis and Examples A detailed breakdown of error types across multiple interactions, along with illustrative examples. These highlight common failure reasons and examples in the models, offering insights into their limitations and areas for improvement.

(3) Comprehensive Summary A summary of the model’s behavior throughout the interview process across all samples. It highlights the model’s strengths and weaknesses in performing the task. (e.g. GPT-4o in DepthQA: “The model struggles with conciseness, often providing overly detailed responses...”)

4 Experimental Setup

In this section, we describe the overall experimental setup used in § 5-6, as well as the design of the **LLM-as-an-Interviewer** framework. In our experiment, we use three tasks: Reasoning, Knowledge, and Instruction Following. For Reasoning, we use MATH (Hendrycks et al., 2021) and the reasoning subset of MINT (Wang et al., 2023). For Knowledge, we use DepthQA (Ko et al., 2024), a dataset of real-world STEM-related questions assessing factuality, rebuilt from the TutorEval dataset. For Instruction Following, we leverage MT-Bench (Chiang et al., 2024), which evaluates an ability to follow multi-turn instructions.

4.1 Models

We list all the models³⁴ used throughout our experiments in Section 5-6.

Interviewer Models We adopt GPT-4o as the default Interviewer, as it demonstrates the best performance (about 90% accuracy as an Interviewer) in our analyses (Table 10 in Appendix C.4).

³We use the Azure API for GPT models, leveraging the GPT-4o-0514 and GPT-3.5-turbo-0125 versions.

⁴We use the **Together API** to access Llama models.

Follow-Up Question	TYPE and Example
Knowledge/Factuality	ADDITIONAL FACTS Can you provide an example of determining the 6th roots of unity and specifying their arguments in radians?
Reasoning	CLARIFICATION What does it mean for a and b to be consecutive integers? RATIONALE How did you determine Jasmine’s rate of water consumption per mile?
Instruction Following	MODIFICATION OF CONDITIONS If we change the previous question and assume that each sister of David has two brothers, how many brothers would David have? ADDITIONAL EXPLANATION Can you provide a specific example of how you would address a potential objection your friend might have about public speaking in your persuasive email? CORRECTION Can you now reformulate your earlier reply, outputting it in JSON format and only including books published after 1980?

Table 1: Examples of Follow-Up Questions posed by the interviewer (GPT-4o) along with the tasks and types.

Interviewee Models We use a total of 8 models for the Interviewee, including 2 proprietary models (GPT-4o (OpenAI, 2024), GPT-3.5) and 6 open-source models (Llama-3.1-{8,70}b (Grattafiori et al., 2024), deepseek-Math-7b (Shao et al., 2024), Qwen2.5-Math-7b (Yang et al., 2024), OLMoE (Muennighoff et al., 2024) and Zephyr-7b (Tunstall et al., 2023)).

4.2 Interview Process Design

LLM-as-an-Interviewer framework adapts to each task type’s characteristics while maintaining rigorous evaluation standards. This section describes the interview process design used in the experiments. Additional details regarding the prompts used are provided in the Appendix C-D. An example of the full interview log is attached in the Appendix F.

Feedback Criteria The Interviewer provides feedback when there is an error in the model’s response or final answer (if the task involves absolute answers). We prompt the Interviewer to guide the model to recognize its mistakes and revise accordingly, but not to provide the correct answer directly. We provide feedback examples in Table 8 of the Appendix, where the Interviewer gives feedback on incorrect responses or missing aspects based on the evaluation criteria.

Follow-Up question Type For each task, we categorize follow-up question types as shown in Table 1. For Knowledge/Factuality tasks, follow-ups probe for additional facts, partially covering the role of a recall metric by identifying missing

information from the model’s previous response—a common limitation in factuality evaluation (Min et al., 2023). For reasoning tasks, follow-ups aim to uncover the underlying rationale or failure reasons. For instruction following, we define three follow-up types based on the type of second-turn questions by users summarized in MT-Bench-101 (Bai et al., 2024).

5 Evaluating LLMs with LLM-AS-AN-INTERVIEWER

We simulate the interview process with various Interviewee models. We present insights from interviews, focusing on the seed question-solving phase in § 5.1, and discuss the follow-up questions phase in § 5.2.

Test set We sample 500 queries⁵ from both MATH and DepthQA, ensuring an equal distribution across the difficulty levels specified in the datasets. We use all 316 samples from the MINT reasoning set and 80 samples from MT-Bench.

Interviewee Models We use 6 models for the Interviewee, GPT-4o, GPT-3.5, Llama-3.1-{8,70}b, DeepSeek-Math-7b, Qwen2.5-Math-7b.

5.1 Seed Question Solving Phase

Table 2 presents the scores of seed questions in MATH and DepthQA for both the LLM-as-a-Judge and LLM-as-an-Interviewer settings. In this experiment, query modification is applied at the

⁵We provide a statistical analysis of the sampling in Appendix C.2.

Model	Reasoning (MATH)					Factuality (DepthQA)				
	Judge	@1	@2	@3	Δ	Judge	@1	@2	@3	Δ
GPT-4o	0.553	0.444	0.566	0.636	0.192	0.989	0.987	0.997	0.999	0.012
GPT-3.5-turbo	0.162	0.082	0.236	0.346	0.264	0.986	0.916	0.975	0.984	0.068
Llama3.1-70b	0.398	0.254	0.436	0.552	0.298	0.990	0.959	0.985	0.989	0.030
Llama3.1-8B	0.222	0.128	0.217	0.300	0.172	0.984	0.906	0.970	0.977	0.071
DeepSeek-math	0.140	0.088	0.222	0.348	0.260	0.977	0.875	0.938	0.958	0.083
Qwen-math	0.454	0.303	0.429	0.483	0.180	0.814	0.878	0.898	0.907	0.029
<i>std.</i>	<i>0.143</i>	<i>0.161</i>	<i>0.123</i>	<i>0.126</i>	-	<i>0.0597</i>	<i>0.0485</i>	<i>0.0482</i>	<i>0.0481</i>	-

Table 2: MATH Problem Solving Accuracy and DepthQA Precision along with Delta (Δ) for six models. @k indicates the performance of the k-th interaction. Judge indicate scores of LLM-as-a-judge.

beginning of the interview, and the same questions are used for the interviewee models. We present the results from MINT in Table 11 of Appendix E.1.

Comparison with LLM-as-a-Judge Score

There is a clear performance drop in the LLM-as-an-Interviewer setting, with most models struggling more when handling modified queries. In MATH, most models achieve performance comparable to the Judge during the second interaction. For DepthQA, all models except GPT-4o experience a significant performance decline, while GPT-4o shows only a minimal decrease of 0.2%. In addition, the evaluation criteria that assess both the reasoning process and the final answers can contribute to this performance drop. For example, GPT-3.5 correctly answers the questions but introduces errors in the solving process in 3% of the MATH test set.

Model’s Capability to Revise Through Feedback

Across all models, performance improves when feedback is provided for incorrect answers (@2, @3), compared to the initial accuracy (@1). Interestingly, as the models receive feedback and iteratively revise their answers, the standard deviation (*std*) of scores across models generally decreases compared to that in the LLM-as-a-Judge setting and @1 score. This is mainly because models that initially perform poorly show substantial improvement through interaction.

5.2 Follow-UP Question Solving Phase

Table 3 presents follow-up question accuracy in Reasoning and Factuality tasks. Overall, models exhibit higher accuracy on follow-up questions for correctly answered seed questions than for failed ones. In § 5.2.1, we provide a discussion based on the MT-Bench evaluation.

5.2.1 Comparing with Static Multi-turn Benchmark

We compare follow-up questions in LLM-as-an-Interviewer with MT-Bench (Chiang et al., 2024), a 2-turn static benchmark with 80 samples across various tasks (e.g., writing, math, coding, role-play).

MT-Bench primarily uses general second-turn questions that are independent of previous responses (e.g., “Explain like I’m five,” “Change format”), limiting diversity, and realism. In contrast, LLM-as-an-Interviewer generates context-aware follow-ups (e.g., “You mentioned Thomas might be in a clinical trial. Can you explain what a clinical trial is and why it requires daily hospital visits?”). Some MT-Bench tasks reference previous responses (e.g., “Explain the N-th example from your previous response.”), which can cause contradictions if the model fails to follow instructions in turn 1, as seen in the following example.

MT-bench Failure Example of Deepseek

Question-Turn-1:

Help me construct a catchy, yet scientifically accurate, headline for an article on the (...). Propose 4 options.

Model Response (Deepseek): (Irrelevant Response)

```
from sympy import symbols, simplify, solve
def solve_equation():
    x = symbols('x')
    equation = x**2 - 16 (...)
```

Question-Turn-2:

Alter your previous response. Make the following adjustments to the 2nd option: 1. Make the tone sound casual (...)

Table 4 shows the number of dependent and independent follow-up questions for each setting. For the Interviewer setting, we compare GPT-4o, the best-performing model at turn-1, and Qwen, the worst-performing model with low instruction-following ability⁶. Aligned with MT-Bench fail-

⁶See Table 13 in Appendix E for the results in MT-bench

Model	Reasoning (MATH)			Reasoning (MINT)			Factuality (DepthQA)		
	Success	Fail	Overall	Success	Fail	Overall	Success	Fail	Overall
GPT-4o	0.99	0.17	0.93	0.96	0.82	0.94	0.93	0.50	0.92
GPT-3.5	0.90	0.56	0.82	0.85	0.70	0.80	0.73	0.64	0.73
Llama3.1-70b	0.91	0.22	0.84	0.91	0.74	0.87	0.85	0.63	0.83
Llama3.1-8B	0.89	0.07	0.76	0.75	0.65	0.72	0.80	1.00	0.80
DeepSeek-math	0.82	0.38	0.69	0.42	0.41	0.42	0.68	0.50	0.65
Qwen-2.5-math	0.90	0.08	0.78	0.55	0.34	0.50	0.61	0.32	0.52

Table 3: Follow-up accuracy of six models on MATH, MINT, and Depth-QA datasets.

	MT-bench	Interviewer	
		GPT-4o	Qwen
Dependent	27	54	44
Independent	53	26	36

Table 4: Dependent and Independent question Counts in 2-turn questions. GPT-4o and Qwen denote the Interviewee.

ure examples, Qwen receives more independent follow-ups (e.g., topic shifts), while GPT-4o receives more dependent ones. As such, considering previous responses dynamically during multi-turn evaluations is crucial for assessing the model in a more diverse and natural context.

6 Can LLM-AS-AN-INTERVIEWER Mitigate Data Contamination Issue?

This section demonstrates the efficiency of the modification strategy in generating questions based on established benchmarks.

6.1 Experimental Setting

We establish various settings depending on the model’s training dataset. Specifically, we use two open-source models, OLMoE and Zephyr-7b, whose training datasets are disclosed, and fine-tune them on different training sets. These settings allow us to evaluate both **Uncontaminated Models**—which do not access the test set during training, and **Contaminated Models**—which are trained with exposure to the test set.

Uncontaminated Model Training As shown in the setting column of Table 5, we train uncontaminated models using three distinct configurations:

- 1 Models trained exclusively on the train set of the target dataset (In-Distribution, $Train_{ID}$).
- 2 Models trained on both the $Train_{ID}$ and additional datasets from the same domain (Out-Of-Distribution, $Train_{OOD}$).
- 3 Models trained solely on $Train_{OOD}$.

Contaminated Model Training For contaminated models, we employ four different configurations:

- 4 Models trained exclusively on the test set of the target dataset (In-Distribution, $Test_{ID}$).
- 5 Models trained on both the $Test_{ID}$ and $Train_{ID}$ sets.
- 6 Models trained on the $Test_{ID}$ combined with an instruction-tuning ($Instruct$) dataset.
- 7 Models trained on both $Test_{ID}$ and $Train_{OOD}$.

Dataset We use the in-distribution test set ($Test_{ID}$) when testing the model, and MATH and DepthQA are used as in-distribution datasets. Since DepthQA is not pre-divided into train and test sets, we manually split it into 851 instances for the $DepthQA_{train}$ set and 848 instances for the $DepthQA_{test}$ set. For other datasets, we randomly select 2000 samples each. For the out-of-distribution dataset of MATH, we use GSM8K. For MATH, we experiment with all seven settings, while for DepthQA, we include only the four settings that do not involve out-of-distribution (OOD) data.

Details about model training, including training configurations, parameters, and datasets used, are provided in Appendix G.1.

6.2 Result

Table 5 presents the effectiveness of the modifying strategy in mitigating data contamination, which is applied while the Interviewer prepares the seed queries. The tables compare accuracy for both the original queries just as LLM-as-an-Judge setting ($Judge$) and modified queries ($Interview$). Settings [1, 2, 3] represent uncontaminated models, serving as a proxy for the model’s true ability on the task. The goal of the strategy is to bring the performance of contaminated models closer to these uncontaminated models.

In MATH, contaminated models often outper-

		OLMoE		Zephyr	
<i>Model.</i>		Judge.	Interv.	Judge.	Interv.
Uncontam.					
Train set _{ID}	[1]	0.10	0.05	0.10	0.01
+Train set _{OOD}	[2]	0.07	0.05	0.08	0.04
Train set _{OOD}	[3]	0.07	0.06	0.06	0.02
Avg.		0.08	0.05	0.08	0.02
Contam.					
Test set _{ID}	[4]	0.84	0.08	0.61	0.10
+Train set _{ID}	[5]	0.77	0.15	0.64	0.17
+Instruct	[6]	0.56	0.12	0.19	0.09
+Train set _{OOD}	[7]	0.75	0.10	0.42	0.08
Avg.		0.69	0.12	0.42	0.11

Table 5: **Performance of uncontaminated and contaminated models on MATH dataset.** "Judge" refers to the LLM-as-a-Judge setting, and "Interv" refers to the LLM-as-an-Interviewer setting.

form uncontaminated ones under the *Judge* setting. These contaminated models achieve significantly higher accuracy, averaging 0.7 for OLMoE and 0.4 for Zephyr. However, applying the modification strategy leads to a substantial performance drop, making their performance comparable to the uncontaminated model. A similar trend is observed in the DepthQA results (Table 14 in Appendix G). The results show that our framework mitigates the contamination issue by having the model solve similar yet modified questions.

7 Discussion

We assess whether LLM-as-an-Interviewer exhibits verbosity bias and maintains robustness across multiple runs, a common consideration in model-based evaluation (Zheng et al., 2023b).⁷

Verbosity Bias To assess whether the framework favors longer answers, we examine the correlation between answer length and score using the Depth-QA task. GPT-4o acts as the *Interviewer*, with six models as *Interviewee*.⁸

Figures 3 in the Appendix H.1 reveal that the linear correlation (r) between length and score weakens as interactions increase. In the LLM-as-a-Judge setting, a significant positive correlation is found ($r = .371$, $p < .05$); however, as interactions progress, the p -value rises and r -value declines. This indicates that our interview pro-

⁷We also provide an analysis of self-enhancement bias in the appendix H.2.

⁸*Interviewee*: GPT-4o, GPT-3.5-turbo, Llama-3.1-70b, Llama-3.1-8b, DeepSeek-Math, Qwen2-Math.

<i>Interact.</i>	@1	@2	@3	@4	@5
<i>DepthQA</i>	0.0244	0.0151	0.0116	0.0090	0.0082
<i>MATH</i>	0.0358	0.0480	0.0583	0.0642	0.0704
<i>easy</i>	0.1115	0.1127	0.0839	0.0797	0.0890
<i>hard</i>	0.0531	0.1032	0.1281	0.1372	0.1332

Table 6: Standard Deviation of Multiple Runs Across Interactions.

cess could reduce verbosity bias and length dependency, alongside the length control mechanism for adjusting win rates or scores (Dubois et al., 2024a).

Does LLM-AS-AN-INTERVIEWER Function Robust?

We assess the robustness of our framework across multiple interactions by repeating the experiment five times for each setting. Specifically, We calculate the standard deviation of these five runs for each combination of different *Interviewer* temperatures (0 or 1) and different *Interviewee* models. The *Interviewee*'s temperature remains fixed at 1. We use GPT-4o as the *Interviewer* and GPT-3.5 and Llama-3.1-{8/70}b as the *Interviewee*.

As shown in Table 6, the standard deviation for DepthQA decreases with more interactions. In contrast, for MATH, it increases over time. Breaking MATH down by difficulty, we find that for easier problems (levels ≤ 2 , MATH-easy), the standard deviation remains stable or decreases, whereas for harder problems (levels ≥ 3 , MATH-hard), it rises due to initially low accuracy and occasional correct answers from feedback. Excluding MATH-hard, the standard deviation remains stable, demonstrating the framework's robustness in most cases.

8 Conclusion

In this paper, we introduce LLM-as-an-Interviewer, a novel approach for evaluating LLMs through interviews. Our multi-turn evaluation framework with Interview Report reveals the abilities of LLMs, such as revising answers based on feedback or handling follow-up questions, while providing comprehensive insights into the task. Our analysis demonstrates that our framework mitigates common issues in static evaluation methods, including verbosity bias and data contamination. We hope our framework is widely adopted and brings a paradigm shift in LLM evaluation.

9 Limitation and Future Work

While our framework addresses several biases and limitations of the LLM-as-a-judge approach, it still inherits some of the inherent limitations of large language models, such as inconsistency and the potential to introduce additional biases, similar to the issues raised in the LLM-as-a-judge approach.

As part of future work, we plan to develop a dataset specifically designed to evaluate the capabilities of interviewer models. Additionally, we aim to assess our framework not only on benchmark datasets but also in user-defined scenarios to better capture real-world applications.

References

- Ge Bai, Jie Liu, Xingyuan Bu, Yancheng He, Jiaheng Liu, Zhanhui Zhou, Zhuoran Lin, Wenbo Su, Tiezheng Ge, Bo Zheng, and Wanli Ouyang. 2024. [MT-bench-101: A fine-grained benchmark for evaluating large language models in multi-turn dialogues](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 7421–7454, Bangkok, Thailand. Association for Computational Linguistics.
- Yupeng Chang, Xu Wang, Jindong Wang, Yuan Wu, Linyi Yang, Kaijie Zhu, Hao Chen, Xiaoyuan Yi, Cunxiang Wang, Yidong Wang, et al. 2024. A survey on evaluation of large language models. *ACM Transactions on Intelligent Systems and Technology*, 15(3):1–45.
- Cheng-Han Chiang and Hung-Yi Lee. 2023. Can large language models be an alternative to human evaluations? In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 15607–15631.
- Wei-Lin Chiang, Lianmin Zheng, Ying Sheng, Anastasios Nikolas Angelopoulos, Tianle Li, Dacheng Li, Hao Zhang, Banghua Zhu, Michael Jordan, Joseph E. Gonzalez, and Ion Stoica. 2024. [Chatbot arena: An open platform for evaluating llms by human preference](#).
- Jasper Dekoninck, Mark Niklas Müller, Maximilian Baader, Marc Fischer, and Martin Vechev. 2024a. [Evading data contamination detection for language models is \(too\) easy](#).
- Jasper Dekoninck, Mark Niklas Müller, and Martin Vechev. 2024b. [Constat: Performance-based contamination detection in large language models](#).
- Chunyuan Deng, Yilun Zhao, Xiangru Tang, Mark Gestein, and Arman Cohan. 2024. [Investigating data contamination in modern benchmarks for large language models](#).
- Jesse Dodge, Maarten Sap, Ana Marasović, William Agnew, Gabriel Ilharco, Dirk Groeneveld, Margaret Mitchell, and Matt Gardner. 2021. [Documenting large webtext corpora: A case study on the colossal clean crawled corpus](#).
- Yann Dubois, Balázs Galambosi, Percy Liang, and Tatsunori B. Hashimoto. 2024a. [Length-controlled alpaca-eval: A simple way to debias automatic evaluators](#).
- Yann Dubois, Percy Liang, and Tatsunori Hashimoto. 2024b. Length-controlled alpaca-eval: A simple debiasing of automatic evaluators. In *First Conference on Language Modeling*.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Allonsius, Daniel Song, Danielle Pintz, Danny Livshits, Danny Wyatt, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Francisco Guzmán, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Govind Thattai, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Ishan Misra, Ivan Evtimov, Jack Zhang, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Karthik Prasad, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer, Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal Lakhotia, Lauren Rantala-Yeary, Laurens van der Maaten, Lawrence Chen, Liang Tan, Liz Jenkins, Louis Martin, Lovish Madaan, Lubo Malo, Lukas Blecher, Lukas Landzaat, Luke de Oliveira, Madeline Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew Oldham, Mathieu Rita, Maya Pavlova, Melanie Kam-badur, Mike Lewis, Min Si, Mitesh Kumar Singh, Mona Hassan, Naman Goyal, Narjes Torabi, Nikolay Bashlykov, Nikolay Bogoychev, Niladri Chatterji, Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick Alrassy, Pengchuan Zhang, Pengwei Li,

695	Petar Vasic, Peter Weng, Prajjwal Bhargava, Pratik	Habeeb, Harrison Rudolph, Helen Suk, Henry As-	758
696	Dubal, Praveen Krishnan, Punit Singh Koura, Puxin	pegren, Hunter Goldman, Hongyuan Zhan, Ibrahim	759
697	Xu, Qing He, Qingxiao Dong, Ragavan Srini-	Damlaj, Igor Molybog, Igor Tufanov, Ilias Leon-	760
698	vasan, Raj Ganapathy, Ramon Calderer, Ricardo Sil-	tiadis, Irina-Elena Veliche, Itai Gat, Jake Weissman,	761
699	veira Cabral, Robert Stojnic, Roberta Raileanu, Ro-	James Geboski, James Kohli, Janice Lam, Japhet	762
700	han Maheswari, Rohit Girdhar, Rohit Patel, Ro-	Asher, Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang,	763
701	main Sauvestre, Ronnie Polidoro, Roshan Sum-	Jennifer Chan, Jenny Zhen, Jeremy Reizenstein,	764
702	baly, Ross Taylor, Ruan Silva, Rui Hou, Rui Wang,	Jeremy Teboul, Jessica Zhong, Jian Jin, Jingyi Yang,	765
703	Saghar Hosseini, Sahana Chennabasappa, Sanjay	Joe Cummings, Jon Carvill, Jon Shepard, Jonathan	766
704	Singh, Sean Bell, Seohyun Sonia Kim, Sergey	McPhie, Jonathan Torres, Josh Ginsburg, Junjie	767
705	Edunov, Shaoliang Nie, Sharan Narang, Sharath Ra-	Wang, Kai Wu, Kam Hou U, Karan Saxena, Kar-	768
706	parthy, Sheng Shen, Shengye Wan, Shruti Bhos-	tikay Khandelwal, Katayoun Zand, Kathy Matosich,	769
707	ale, Shun Zhang, Simon Vandenhende, Soumya Ba-	Kaushik Veeraraghavan, Kelly Michelena, Keqian	770
708	tra, Spencer Whitman, Sten Sootla, Stephane Col-	Li, Kiran Jagadeesh, Kun Huang, Kunal Chawla,	771
709	lot, Suchin Gururangan, Sydney Borodinsky, Tamar	Kyle Huang, Lailin Chen, Lakshya Garg, Laven-	772
710	Herman, Tara Fowler, Tarek Sheasha, Thomas	der A, Leandro Silva, Lee Bell, Lei Zhang, Liang-	773
711	Georgiou, Thomas Scialom, Tobias Speckbacher,	peng Guo, Licheng Yu, Liron Moshkovich, Luca	774
712	Todor Mihaylov, Tong Xiao, Ujjwal Karn, Vedanuj	Wehrstedt, Madian Khabsa, Manav Avalani, Manish	775
713	Goswami, Vibhor Gupta, Vignesh Ramanathan,	Bhatt, Martynas Mankus, Matan Hasson, Matthew	776
714	Viktor Kerkez, Vincent Gonguet, Virginie Do,	Lennie, Matthias Reso, Maxim Groshev, Maxim	777
715	Vish Vogeti, Vitor Albiero, Vladan Petrovic, Wei-	Naumov, Maya Lathi, Meghan Keneally, Miao	778
716	wei Chu, Wenhan Xiong, Wenyin Fu, Whitney	Liu, Michael L. Seltzer, Michal Valko, Michelle	779
717	Meers, Xavier Martinet, Xiaodong Wang, Xiao-	Restrepo, Mihir Patel, Mik Vyatskov, Mikayel	780
718	fang Wang, Xiaoqing Ellen Tan, Xide Xia, Xin-	Samvelyan, Mike Clark, Mike Macey, Mike Wang,	781
719	feng Xie, Xuchao Jia, Xuewei Wang, Yaelle Gold-	Miquel Jubert Hermoso, Mo Metanat, Mohammad	782
720	schlag, Yashesh Gaur, Yasmine Babaei, Yi Wen, Yi-	Rastegari, Munish Bansal, Nandhini Santhanam,	783
721	wen Song, Yuchen Zhang, Yue Li, Yuning Mao,	Natascha Parks, Natasha White, Navyata Bawa,	784
722	Zacharie Delpierre Coudert, Zheng Yan, Zhengx-	Nayan Singhal, Nick Egebo, Nicolas Usunier, Nikhil	785
723	ing Chen, Zoe Papakipos, Aaditya Singh, Aayushi	Mehta, Nikolay Pavlovich Laptev, Ning Dong, Nor-	786
724	Srivastava, Abha Jain, Adam Kelsey, Adam Sha-	man Cheng, Oleg Chernoguz, Olivia Hart, Omkar	787
725	infeld, Adithya Gangidi, Adolfo Victoria, Ahuva	Salpekar, Ozlem Kalinli, Parkin Kent, Parth Parekh,	788
726	Goldstand, Ajay Menon, Ajay Sharma, Alex Boe-	Paul Saab, Pavan Balaji, Pedro Rittner, Philip Bon-	789
727	senberg, Alexei Baevski, Allie Feinstein, Amanda	trager, Pierre Roux, Piotr Dollar, Polina Zvyag-	790
728	Kallet, Amit Sangani, Amos Teo, Anam Yunus, An-	ina, Prashant Ratanchandani, Pritish Yuvraj, Qian	791
729	drei Lupu, Andres Alvarado, Andrew Caples, An-	Liang, Rachad Alao, Rachel Rodriguez, Rafi Ayub,	792
730	drew Gu, Andrew Ho, Andrew Poulton, Andrew	Raghotham Murthy, Raghu Nayani, Rahul Mi-	793
731	Ryan, Ankit Ramchandani, Annie Dong, Annie	tra, Rangaprabhu Parthasarathy, Raymond Li, Re-	794
732	Franco, Anuj Goyal, Aparajita Saraf, Arkabandhu	bekkah Hogan, Robin Battey, Rocky Wang, Russ	795
733	Chowdhury, Ashley Gabriel, Ashwin Bharambe,	Howes, Ruty Rinott, Sachin Mehta, Sachin Siby,	796
734	Assaf Eisenman, Azadeh Yazdan, Beau James, Ben	Sai Jayesh Bondu, Samyak Datta, Sara Chugh, Sara	797
735	Maurer, Benjamin Leonhardi, Bernie Huang, Beth	Hunt, Sargun Dhillon, Sasha Sidorov, Satadru Pan,	798
736	Loyd, Beto De Paola, Bhargavi Paranjape, Bing	Saurabh Mahajan, Saurabh Verma, Seiji Yamamoto,	799
737	Liu, Bo Wu, Boyu Ni, Braden Hancock, Bram	Sharadh Ramaswamy, Shaun Lindsay, Shaun Lind-	800
738	Wasti, Brandon Spence, Brani Stojkovic, Brian	say, Sheng Feng, Shenghao Lin, Shengxin Cindy	801
739	Gamido, Britt Montalvo, Carl Parker, Carly Burton,	Zha, Shishir Patil, Shiva Shankar, Shuqiang Zhang,	802
740	Catalina Mejia, Ce Liu, Changhan Wang, Changkyu	Shuqiang Zhang, Sinong Wang, Sneha Agarwal,	803
741	Kim, Chao Zhou, Chester Hu, Ching-Hsiang Chu,	Soji Sajuyigbe, Soumith Chintala, Stephanie Max,	804
742	Chris Cai, Chris Tindal, Christoph Feichtenhofer,	Stephen Chen, Steve Kehoe, Steve Satterfield,	805
743	Cynthia Gao, Damon Civin, Dana Beaty, Daniel	Sudarshan Govindaprasad, Sumit Gupta, Summer	806
744	Kreymer, Daniel Li, David Adkins, David Xu, Da-	Deng, Sungmin Cho, Sunny Virk, Suraj Subrama-	807
745	vide Testuggine, Delia David, Devi Parikh, Di-	nian, Sy Choudhury, Sydney Goldman, Tal Remez,	808
746	ana Liskovich, Didem Foss, Dingkan Wang, Duc	Tamar Glaser, Tamara Best, Thilo Koehler, Thomas	809
747	Le, Dustin Holland, Edward Dowling, Eissa Jamil,	Robinson, Tianhe Li, Tianjun Zhang, Tim Matthews,	810
748	Elaine Montgomery, Eleonora Presani, Emily Hahn,	Timothy Chou, Tzook Shaked, Varun Vontimitta,	811
749	Emily Wood, Eric-Tuan Le, Erik Brinkman, Esteban	Victoria Ajayi, Victoria Montanez, Vijai Mohan,	812
750	Arcaute, Evan Dunbar, Evan Smothers, Fei Sun, Fe-	Vinay Satish Kumar, Vishal Mangla, Vlad Ionescu,	813
751	lix Kreuk, Feng Tian, Filippas Kokkinos, Firat Oz-	Vlad Poenaru, Vlad Tiberiu Mihailescu, Vladimir	814
752	genel, Francesco Caggioni, Frank Kanayet, Frank	Ivanov, Wei Li, Wenchen Wang, Wenwen Jiang,	815
753	Seide, Gabriela Medina Florez, Gabriella Schwarz,	Wes Bouaziz, Will Constable, Xiaocheng Tang, Xi-	816
754	Gada Badeer, Georgia Swee, Gil Halpern, Grant	aojian Wu, Xiaolan Wang, Xilun Wu, Xinbo Gao,	817
755	Herman, Grigory Sizov, Guanyi, Zhang, Guna Lak-	Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia,	818
756	shminarayanan, Hakan Inan, Hamid Shojanazeri,	Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi	819
757	Han Zou, Hannah Wang, Hanwen Zha, Haroun	Adi, Youngjin Nam, Yu, Wang, Yu Zhao, Yuchen	820

821	Hao, Yundi Qian, Yunlu Li, Yuzi He, Zach Rait,	question answering models . In <i>Proceedings of the</i>	876
822	Zachary DeVito, Zef Rosnbrick, Zhaoduo Wen,	<i>2023 Conference on Empirical Methods in Natural</i>	877
823	Zhenyu Yang, Zhiwei Zhao, and Zhiyu Ma. 2024.	<i>Language Processing</i> , pages 3435–3446, Singapore.	878
824	The llama 3 herd of models .	Association for Computational Linguistics.	879
825	Jiawei Gu, Xuhui Jiang, Zhichao Shi, Hexiang Tan,	Xuechen Li, Tianyi Zhang, Yann Dubois, Rohan Taori,	880
826	Xuehao Zhai, Chengjin Xu, Wei Li, Yinghan Shen,	Ishaan Gulrajani, Carlos Guestrin, Percy Liang, and	881
827	Shengjie Ma, Honghao Liu, et al. 2024. A survey on	Tatsunori B. Hashimoto. 2023b. AlpacaEval: An au-	882
828	llm-as-a-judge. <i>arXiv preprint arXiv:2411.15594</i> .	tomatic evaluator of instruction-following models.	883
829	Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul	https://github.com/tatsu-lab/alpaca_eval .	884
830	Arora, Steven Basart, Eric Tang, Dawn Song, and	Percy Liang, Rishi Bommasani, Tony Lee, Dimitris	885
831	Jacob Steinhardt. 2021. Measuring mathematical	Tsipras, Dilara Soylu, Michihiro Yasunaga, Yian	886
832	problem solving with the math dataset .	Zhang, Deepak Narayanan, Yuhuai Wu, Ananya Ku-	887
833	Naman Jain, King Han, Alex Gu, Wen-Ding Li, Fan-	mar, et al. 2022. Holistic evaluation of language	888
834	jia Yan, Tianjun Zhang, Sida Wang, Armando Solar-	models. <i>arXiv preprint arXiv:2211.09110</i> .	889
835	Lezama, Koushik Sen, and Ion Stoica. 2024. Live-	Chin-Yew Lin. 2004. Rouge: A package for automatic	890
836	codebench: Holistic and contamination free evalua-	evaluation of summaries. In <i>Text summarization</i>	891
837	tion of large language models for code .	<i>branches out</i> , pages 74–81.	892
838	Minhao Jiang, Ken Ziyu Liu, Ming Zhong, Rylan Scha-	Inbal Magar and Roy Schwartz. 2022. Data contami-	893
839	ffer, Siru Ouyang, Jiawei Han, and Sanmi Koyejo.	nation: From memorization to exploitation .	894
840	2024. Investigating data contamination for pre-	Sewon Min, Kalpesh Krishna, Xinxu Lyu, Mike Lewis,	895
841	training language models .	Wen-tau Yih, Pang Koh, Mohit Iyyer, Luke Zettle-	896
842	Seungone Kim, Jamin Shin, Yejin Cho, Joel Jang,	moyer, and Hannaneh Hajishirzi. 2023. FActScore:	897
843	Shayne Longpre, Hwaran Lee, Sangdoo Yun,	Fine-grained atomic evaluation of factual precision	898
844	Seongjin Shin, Sungdong Kim, James Thorne, and	in long form text generation . In <i>Proceedings of the</i>	899
845	Minjoon Seo. 2024. Prometheus: Inducing fine-	<i>2023 Conference on Empirical Methods in Natural</i>	900
846	grained evaluation capability in language models .	<i>Language Processing</i> , pages 12076–12100, Singa-	901
847	Miyoung Ko, Sue Hyun Park, Joonsuk Park, and Min-	pore. Association for Computational Linguistics.	902
848	joon Seo. 2024. Hierarchical deconstruction of	Iman Mirzadeh, Keivan Alizadeh, Hooman Shahrokhi,	903
849	LLM reasoning: A graph-based framework for an-	Oncel Tuzel, Samy Bengio, and Mehrdad Farajtabar.	904
850	alyzing knowledge utilization . In <i>Proceedings of</i>	2024. Gsm-symbolic: Understanding the limitations	905
851	<i>the 2024 Conference on Empirical Methods in Natu-</i>	of mathematical reasoning in large language models .	906
852	<i>ral Language Processing</i> , pages 4995–5027, Miami,	Niklas Muennighoff, Luca Soldaini, Dirk Groeneveld,	907
853	Florida, USA. Association for Computational Lin-	Kyle Lo, Jacob Morrison, Sewon Min, Weijia	908
854	guistics.	Shi, Pete Walsh, Oyvind Tafjord, Nathan Lam-	909
855	Wai-Chung Kwan, Xingshan Zeng, Yuxin Jiang, Yufei	bert, Yuling Gu, Shane Arora, Akshita Bhagia,	910
856	Wang, Liangyou Li, Lifeng Shang, Xin Jiang, Qun	Dustin Schwenk, David Wadden, Alexander Wet-	911
857	Liu, and Kam-Fai Wong. 2024. Mt-eval: A multi-	tig, Binyuan Hui, Tim Dettmers, Douwe Kiela, Ali	912
858	turn capabilities evaluation benchmark for large lan-	Farhadi, Noah A. Smith, Pang Wei Koh, Aman-	913
859	guage models. <i>arXiv preprint arXiv:2401.16745</i> .	preet Singh, and Hannaneh Hajishirzi. 2024. OL-	914
860	Dawei Li, Bohan Jiang, Liangjie Huang, Alimoham-	moe: Open mixture-of-experts language models .	915
861	mad Beigi, Chengshuai Zhao, Zhen Tan, Amrita	OpenAI. 2024. Gpt-4o system card .	916
862	Bhattacharjee, Yuxuan Jiang, Canyu Chen, Tianhao	OpenAI, Josh Achiam, Steven Adler, Sandhini Agar-	917
863	Wu, et al. 2024a. From generation to judgment: Op-	wal, Lama Ahmad, Ilge Akkaya, Florencia Leoni	918
864	portunities and challenges of llm-as-a-judge. <i>arXiv</i>	Aleman, Diogo Almeida, Janko Altmenschmidt,	919
865	<i>preprint arXiv:2411.16594</i> .	Sam Altman, Shyamal Anadkat, Red Avila, Igor	920
866	Margaret Li, Jason Weston, and Stephen Roller. 2019.	Babuschkin, Suchir Balaji, Valerie Balcom, Paul	921
867	Acute-eval: Improved dialogue evaluation with opti-	Baltescu, Haiming Bao, Mohammad Bavarian,	922
868	mized questions and multi-turn comparisons. <i>arXiv</i>	Jeff Belgum, Irwan Bello, Jake Berdine, Gabriel	923
869	<i>preprint arXiv:1909.03087</i> .	Bernadett-Shapiro, Christopher Berner, Lenny Bog-	924
870	Xiang Lisa Li, Evan Zheran Liu, Percy Liang, and Tat-	donoff, Oleg Boiko, Madelaine Boyd, Anna-Luisa	925
871	sunori Hashimoto. 2024b. Autobench: Creating	Brakman, Greg Brockman, Tim Brooks, Miles	926
872	salient, novel, difficult datasets for language models .	Brundage, Kevin Button, Trevor Cai, Rosie Camp-	927
873	Xibo Li, Bowei Zou, Yifan Fan, Yanling Li, Ai Ti Aw,	bell, Andrew Cann, Brittany Carey, Chelsea Carl-	928
874	and Yu Hong. 2023a. Interview evaluation: A novel	son, Rory Carmichael, Brooke Chan, Che Chang,	929
875	approach for automatic evaluation of conversational	Fotis Chantzis, Derek Chen, Sully Chen, Ruby	930
		Chen, Jason Chen, Mark Chen, Ben Chess, Chester	931

932	Cho, Casey Chu, Hyung Won Chung, Dave Cummings, Jeremiah Currier, Yunxing Dai, Cory Decareaux, Thomas Degry, Noah Deutsch, Damien Deville, Arka Dhar, David Dohan, Steve Dowling, Sheila Dunning, Adrien Ecoffet, Atty Eleti, Tyna Eloundou, David Farhi, Liam Fedus, Niko Felix, Simón Posada Fishman, Juston Forte, Isabella Fulford, Leo Gao, Elie Georges, Christian Gibson, Vik Goel, Tarun Gogineni, Gabriel Goh, Rapha Gontijo-Lopes, Jonathan Gordon, Morgan Grafstein, Scott Gray, Ryan Greene, Joshua Gross, Shixiang Shane Gu, Yufei Guo, Chris Hallacy, Jesse Han, Jeff Harris, Yuchen He, Mike Heaton, Johannes Heidecke, Chris Hesse, Alan Hickey, Wade Hickey, Peter Hoeschele, Brandon Houghton, Kenny Hsu, Shengli Hu, Xin Hu, Joost Huizinga, Shantanu Jain, Shawn Jain, Joanne Jang, Angela Jiang, Roger Jiang, Haozhun Jin, Denny Jin, Shino Jomoto, Billie Jonn, Heewoo Jun, Tomer Kaftan, Łukasz Kaiser, Ali Kamali, Ingmar Kanitscheider, Nitish Shirish Keskar, Tabarak Khan, Logan Kilpatrick, Jong Wook Kim, Christina Kim, Yongjik Kim, Jan Hendrik Kirchner, Jamie Kiros, Matt Knight, Daniel Kokotajlo, Łukasz Kondraciuk, Andrew Kondrich, Aris Konstantinidis, Kyle Kopic, Gretchen Krueger, Vishal Kuo, Michael Lampe, Ikai Lan, Teddy Lee, Jan Leike, Jade Leung, Daniel Levy, Chak Ming Li, Rachel Lim, Molly Lin, Stephanie Lin, Mateusz Litwin, Theresa Lopez, Ryan Lowe, Patricia Lue, Anna Makanju, Kim Malfacini, Sam Manning, Todor Markov, Yaniv Markovski, Bianca Martin, Katie Mayer, Andrew Mayne, Bob McGrew, Scott Mayer McKinney, Christine McLeavey, Paul McMillan, Jake McNeil, David Medina, Aalok Mehta, Jacob Menick, Luke Metz, Andrey Mishchenko, Pamela Mishkin, Vinnie Monaco, Evan Morikawa, Daniel Mossing, Tong Mu, Mira Murati, Oleg Murk, David Mély, Ashvin Nair, Reiichiro Nakano, Rameev Nayak, Arvind Neelakantan, Richard Ngo, Hyeonwoo Noh, Long Ouyang, Cullen O’Keefe, Jakub Pachocki, Alex Paino, Joe Palermo, Ashley Pantuliano, Giambattista Parascandolo, Joel Parish, Emy Parparita, Alex Passos, Mikhail Pavlov, Andrew Peng, Adam Perelman, Filipe de Avila Belbute Peres, Michael Petrov, Henrique Ponde de Oliveira Pinto, Michael, Pokorny, Michelle Pokrass, Vitchyr H. Pong, Tolly Powell, Alethea Power, Boris Power, Elizabeth Proehl, Raul Puri, Alec Radford, Jack Rae, Aditya Ramesh, Cameron Raymond, Francis Real, Kendra Rimbach, Carl Ross, Bob Rotsted, Henri Roussez, Nick Ryder, Mario Saltarelli, Ted Sanders, Shibani Santurkar, Girish Sastry, Heather Schmidt, David Schnurr, John Schulman, Daniel Selsam, Kyla Shepard, Toki Sherbakov, Jessica Shieh, Sarah Shoker, Pranav Shyam, Szymon Sidor, Eric Sigler, Maddie Simens, Jordan Sitkin, Katarina Slama, Ian Sohl, Benjamin Sokolowsky, Yang Song, Natalie Staudacher, Felipe Petroski Such, Natalie Summers, Ilya Sutskever, Jie Tang, Nikolas Tezak, Madeleine B. Thompson, Phil Tillet, Amin Tootoonchian, Elizabeth Tseng, Preston Tuggle, Nick Turley, Jerry Tworek, Juan Felipe Cerón Uribe, Andrea Vallone, Arun Vijayvergiya, Chelsea Voss, Carroll Wainwright, Justin Jay Wang, Alvin Wang, Ben Wang, Jonathan Ward, Jason Wei, CJ Weinmann, Akila Welihinda, Peter Welinder, Jiayi Weng, Lilian Weng, Matt Wiethoff, Dave Willner, Clemens Winter, Samuel Wolrich, Hannah Wong, Lauren Workman, Sherwin Wu, Jeff Wu, Michael Wu, Kai Xiao, Tao Xu, Sarah Yoo, Kevin Yu, Qiming Yuan, Wojciech Zaremba, Rowan Zellers, Chong Zhang, Marvin Zhang, Shengjia Zhao, Tianhao Zheng, Juntang Zhuang, William Zhuk, and Barret Zoph. 2024. Gpt-4 technical report .	995 996 997 998 999 1000 1001 1002 1003 1004 1005
943	Yonatan Oren, Nicole Meister, Niladri Chatterji, Faisal Ladhak, and Tatsunori B. Hashimoto. 2023a. Proving test set contamination in black box language models .	1006 1007 1008 1009
944	Yonatan Oren, Nicole Meister, Niladri S Chatterji, Faisal Ladhak, and Tatsunori Hashimoto. 2023b. Proving test set contamination in black-box language models. In <i>The Twelfth International Conference on Learning Representations</i> .	1010 1011 1012 1013 1014
945	Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. Bleu: a method for automatic evaluation of machine translation. In <i>Proceedings of the 40th annual meeting of the Association for Computational Linguistics</i> , pages 311–318.	1015 1016 1017 1018 1019
946	Oscar Sainz, Jon Campos, Iker García-Ferrero, Julen Etxaniz, Oier Lopez de Lacalle, and Eneko Agirre. 2023. Nlp evaluation in trouble: On the need to measure llm data contamination for each benchmark. In <i>Findings of the Association for Computational Linguistics: EMNLP 2023</i> , pages 10776–10787.	1020 1021 1022 1023 1024 1025
947	Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. 2024. Deepseekmath: Pushing the limits of mathematical reasoning in open language models .	1026 1027 1028 1029 1030
948	Weijia Shi, Anirudh Ajith, Mengzhou Xia, Yangsibo Huang, Daogao Liu, Terra Blevins, Danqi Chen, and Luke Zettlemoyer. 2024. Detecting pretraining data from large language models .	1031 1032 1033 1034
949	Luca Soldaini, Rodney Kinney, Akshita Bhagia, Dustin Schwenk, David Atkinson, Russell Authur, Ben Bogin, Khyathi Chandu, Jennifer Dumas, Yanai Elazar, Valentin Hofmann, Ananya Harsh Jha, Sachin Kumar, Li Lucy, Xinxu Lyu, Nathan Lambert, Ian Magnusson, Jacob Morrison, Niklas Muennighoff, Aakanksha Naik, Crystal Nam, Matthew E. Peters, Abhilasha Ravichander, Kyle Richardson, Zejiang Shen, Emma Strubell, Nishant Subramani, Oyvind Tafjord, Pete Walsh, Luke Zettlemoyer, Noah A. Smith, Hannaneh Hajishirzi, Iz Beltagy, Dirk Groeneveld, Jesse Dodge, and Kyle Lo. 2024. Dolma: an open corpus of three trillion tokens for language model pretraining research .	1035 1036 1037 1038 1039 1040 1041 1042 1043 1044 1045 1046 1047 1048
950	Lewis Tunstall, Edward Beeching, Nathan Lambert, Nazneen Rajani, Kashif Rasul, Younes Belkada, Shengyi Huang, Leandro von Werra, Clémentine	1049 1050 1051

1052	Fourrier, Nathan Habib, Nathan Sarrazin, Omar Sanseviero, Alexander M. Rush, and Thomas Wolf. 2023. Zephyr: Direct distillation of lm alignment .	1109
1053		1110
1054		1111
1055	Siyuan Wang, Zhuohan Long, Zhihao Fan, Zhongyu Wei, and Xuanjing Huang. 2024a. Benchmark self-evolving: A multi-agent framework for dynamic llm evaluation .	1112
1056		1113
1057		
1058		
1059	Xingyao Wang, Zihan Wang, Jiateng Liu, Yangyi Chen, Lifan Yuan, Hao Peng, and Heng Ji. 2023. Mint: Evaluating llms in multi-turn interaction with tools and language feedback. In <i>The Twelfth International Conference on Learning Representations</i> .	1114
1060		1115
1061		1116
1062		1117
1063		1118
1064	Yidong Wang, Zhuohao Yu, Zhengran Zeng, Linyi Yang, Cunxiang Wang, Hao Chen, Chaoya Jiang, Rui Xie, Jindong Wang, Xing Xie, Wei Ye, Shikun Zhang, and Yue Zhang. 2024b. Pandalm: An automatic evaluation benchmark for llm instruction tuning optimization .	1119
1065		1120
1066		1121
1067		1122
1068		
1069		
1070	Colin White, Samuel Dooley, Manley Roberts, Arka Pal, Ben Feuer, Siddhartha Jain, Ravid Shwartz-Ziv, Neel Jain, Khalid Saifullah, Siddhartha Naidu, Chinmay Hegde, Yann LeCun, Tom Goldstein, Willie Neiswanger, and Micah Goldblum. 2024. Livebench: A challenging, contamination-free llm benchmark .	
1071		
1072		
1073		
1074		
1075		
1076		
1077	Ruijie Xu, Zengzhi Wang, Run-Ze Fan, and Pengfei Liu. 2024. Benchmarking benchmark leakage in large language models .	
1078		
1079		
1080	An Yang, Beichen Zhang, Binyuan Hui, Bofei Gao, Bowen Yu, Chengpeng Li, Dayiheng Liu, Jianhong Tu, Jingren Zhou, Junyang Lin, Keming Lu, Mingfeng Xue, Runji Lin, Tianyu Liu, Xingzhang Ren, and Zhenru Zhang. 2024. Qwen2.5-math technical report: Toward mathematical expert model via self-improvement .	
1081		
1082		
1083		
1084		
1085		
1086		
1087	Nicolas Yax, Pierre-Yves Oudeyer, and Stefano Palminteri. 2024. Assessing contamination in large language models: Introducing the logprober method .	
1088		
1089		
1090	Zhuohao Yu, Chang Gao, Wenjin Yao, Yidong Wang, Wei Ye, Jindong Wang, Xing Xie, Yue Zhang, and Shikun Zhang. 2024. KIEval: A knowledge-grounded interactive evaluation framework for large language models . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 5967–5985, Bangkok, Thailand. Association for Computational Linguistics.	
1091		
1092		
1093		
1094		
1095		
1096		
1097		
1098		
1099	Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Q Weinberger, and Yoav Artzi. 2019. Bertscore: Evaluating text generation with bert. <i>arXiv preprint arXiv:1904.09675</i> .	
1100		
1101		
1102		
1103	Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. 2023a. Judging llm-as-a-judge with mt-bench and chatbot arena. <i>Advances in Neural Information Processing Systems</i> , 36:46595–46623.	
1104		
1105		
1106		
1107		
1108		

1123	A AI Writing Assistance		
1124	We used ChatGPT for grammar and language re-		
1125	finement of the manuscript.		
1126	B LLM-as-an-Interviewer		
1127	B.1 Implementation Details		
1128	Building upon the interview framework described		
1129	in § 3.1, this section explains our interview sys-		
1130	tem implementation. The system operates as a		
1131	two-agent dialogue with additional complexity in		
1132	controlling the interviewer to conduct evaluations		
1133	accurately and efficiently. Controlling interviewer		
1134	is done in two stages: (1) keeping the interview		
1135	state in track, and (2) using pre-written prompts		
1136	for each state. While we implemented the inter-		
1137	view process specifically for MATH and STEM		
1138	domains, the framework can be easily adapted to		
1139	other domains with minor modifications.		
1140	B.2 Adapting to New Tasks		
1141	To adapt a new task, one can modify the Main		
1142	Interview process, specifically the response eval-		
1143	uation criteria and follow-up question generation		
1144	strategy. This can be done by rewriting the pre-		
1145	written prompts. Here are examples of adaptations		
1146	for different domains:		
1147	Programming Interviews:		
1148	• Response Evaluation: Code correctness, ef-		
1149	ficiency, and style		
1150	• Follow-up Generation: Edge cases, opti-		
1151	mization opportunities, or alternative imple-		
1152	mentations		
1153	Language Proficiency:		
1154	• Response Evaluation: Grammar accuracy,		
1155	vocabulary usage, and fluency		
1156	• Follow-up Generation: Complex linguistic		
1157	scenarios or cultural context questions		
1158	Customer Service Evaluation		
1159	• Response Evaluation: Empathy in customer		
1160	interactions, response tone and professional-		
1161	ism, policy adherence		
1162	Follow-up Generation:		
1163	• For acceptable responses: Escalate scenario		
1164	complexity (e.g., adding multiple complaints		
1165	or time pressure), Rationale of responses		
		• For problematic responses: Probe under-	1166
		standing of customer service principles or	1167
		store policies	1168
	B.3 Algorithmic Implementation		1169
	Algorithm 1 presents the algorithmic implementa-		1170
	tion of the interview process.		1171

Algorithm 1 Interview Process

Require: Q_{init} : Initial question

Require: $\text{Modify}(\cdot)$: Function to modify the initial question

Require: I : Interviewer

Require: i : Interviewee

Require: G : Grader

Require: $\text{AskFollowup}(\cdot)$: Function to ask a follow-up question

Require: MAX_QUESTIONS : Maximum number of questions

Require: MAX_RETRIES : Maximum number of retries for each question

Constants:

1: $\text{MAX_QUESTIONS} \leftarrow N$

▷ Maximum number of questions

2: $\text{MAX_RETRIES} \leftarrow M$

▷ Maximum retries per question

Variables:

3: $Q_{\text{modified}} \leftarrow \text{Modify}(Q_{\text{init}})$

▷ Modified list of seed questions

4: $\text{feedbacks} \leftarrow []$

5: $\text{questions_count} \leftarrow 0$

6: $\text{termination} \leftarrow \text{False}$

Execution:

7: **for** each question $q_{\text{seed}} \in Q_{\text{modified}}$ **do**

8: $\text{questions_count} \leftarrow \text{questions_count} + 1$

9: $q \leftarrow q_{\text{seed}}$

▷ Start with the seed question

10: **while not** termination **do**

11: $\text{response} \leftarrow i(q)$

▷ Interviewee answers the question

12: $\text{feedback}, \text{is_correct} \leftarrow G(q, \text{response})$

▷ Grader evaluates response

13: $\text{feedbacks.append}(\text{feedback})$

14: **if** is_correct **then**

15: **break**

▷ Move to the next question if correct

16: **end if**

Retries for Incorrect Answers:

17: $\text{retries} \leftarrow 0$

18: **while not** is_correct **and** $\text{retries} < \text{MAX_RETRIES}$ **do**

19: $\text{hint} \leftarrow \text{GetHint}(q, \text{response}, \text{feedback})$

▷ Generate hint

20: $\text{response} \leftarrow i(\text{hint})$

▷ Interviewee retries the question

21: $\text{feedback}, \text{is_correct} \leftarrow G(q, \text{response})$

22: $\text{feedbacks.append}(\text{feedback})$

23: $\text{retries} \leftarrow \text{retries} + 1$

24: **end while**

Follow-up Question:

25: **if** $\text{questions_count} > \text{MAX_QUESTIONS}$ **then**

26: $\text{termination} \leftarrow \text{True}$

▷ Terminate interview after max questions

27: **end if**

28: $q \leftarrow \text{AskFollowup}(q)$

▷ Switch to a follow-up question

29: **end while**

30: **end for**

C Experimental Setup

C.1 Experimental Details

Preparing Seed Question: Interviewer model generates modified versions of seed questions through unclarification or paraphrasing

For Arithmetic Reasoning tasks (MATH), LLM Interviewer replaces numerical values in a query with unknown variables (e.g., x) or omits specific details. For example, "If there were 5 chocolates and you ate 1, how many are left?" becomes "If there were 5 chocolates and you ate x , how many are left?". This maintains the same reasoning level while preventing the direct application of memorized solutions, as seen in methods like those of Mirzadeh et al. (2024). For factuality-based long-form tasks (depth-QA), new problems are created by referencing ground truth solutions and questions from the original dataset to ensure relevance and diversity. Examples of these strategies are in Table 9.

Interview Process: Consists of two primary components:

1. **Response Evaluation:** Interviewer assesses interviewee responses using domain-specific criteria (mathematical correctness for MATH; factual accuracy and completeness for STEM) and identifies error sources.
2. **Follow-up Question Generation:** Based on evaluation results, interviewer either probes deeper concepts (for correct responses) or provides targeted feedback to reveal misconceptions (for incorrect responses).

Evaluation Metric: For MATH, our evaluation framework implements binary correctness assessment for both final answers and solution processes, along with error categorization that distinguishes between conceptual understanding, misinterpretation, and calculation errors. The framework also incorporates step-by-step verification of mathematical reasoning and generates targeted follow-up questions when errors are detected.

For DepthQA (Ko et al., 2024), our evaluation framework uses FActScore (Min et al., 2023) as a metric to calculate factual precision in the model’s long-form generation. These follow-up questions partially address a limitation of previous fact-based evaluations by enabling a somewhat broader assessment that includes aspects

of recall in addition to precision. employs a more nuanced approach, decomposing responses into atomic facts for systematic assessment. The framework also evaluates quality across multiple dimensions, including completeness, redundancy, readability, and depth, while generating follow-up questions to probe understanding of missing concepts.

Interview Report: The interview report, based on the detailed interaction logs from the Main Interview stage, is shown to the user at the end of the interview. This content for the report is explained in § 3.2. Followings are the example of Interview Report.

C.2 Statistical analysis of the Sampling

We compare the sampling variability of our framework using bootstrapping (sampling with replacement). From the 500 samples, we randomly sample n samples, repeated the process 10,000 times, and calculated the p-value. The results showed no significant statistical difference between $n \geq 100$ samples and 500 samples, with results comparable to those of LLM-as-a-Judge.

Number of Samples	100	300	500
Interviewer	p = 0.809	p = 0.832	p = 0.766
Judge	p = 0.464	p = 0.740	p = 0.714

Table 7: Statistical Comparison of Sampling Variability

MATH, Interviewee: GPT-4o
(Interviewer: GPT-4o)

1. Performance Scores(%)

- **Accuracy(@ # of interaction):**
72(@1), 82(@2), 84(@3)
- **Follow-up Accuracy:**
 - Total 93
 - Rationale Questions 99
 - Clarification 17

2. Error Types & Examples

* Frequent Rate for Each Type
* Examples are omitted

- Misinterpret.: 0.17
- Calculation: 0.26
- Conceptual: 0.16

3. Summary

The model demonstrates strong mathematical problem-solving abilities, particularly in algebra, calculus, and logical reasoning. It excels in providing detailed, step-by-step explanations and is effective in handling follow-up questions and user feedback. However, **it struggles with interpreting geometric properties**, maintaining accuracy in complex calculations, and integrating specific **user-provided details without explicit prompts**. The model often requires user intervention to correct initial assumptions and may provide overly verbose explanations. Improvements in handling ambiguous or incomplete information and better error-checking mechanisms could enhance its effectiveness.

MATH, Interviewee: Llama-3.1-8b
(Interviewer: GPT-4o)

1. Performance Scores(%)

- **Accuracy(@ # of interaction):**
58(@1), 68(@2), 75(@3)
- **Follow-up Accuracy:**
 - Total 76
 - Rationale Questions 89
 - Clarification 7

2. Error Types & Examples

* Frequent Rate for Each Type
* Examples are omitted

- Misinterpret.: 0.45
- Calculation: 0.26
- Conceptual: 0.27

3. Summary

The model demonstrates strong mathematical problem-solving skills, particularly in algebra, quadratic equations, and logical reasoning. It excels in providing clear, step-by-step explanations, making it useful for educational purposes. The model is responsive to follow-up questions and user feedback, showing adaptability and a willingness to improve its responses. However, it struggles with **maintaining accuracy in complex calculations, avoiding unnecessary complexity, and handling more abstract or less structured problems**. It may also require user **intervention to correct its approach and simplify its methods**. Overall, the model is effective in structured mathematical tasks but needs improvement in dealing with more complex, ambiguous, or creative problems.

DepthQA, Interviewee: GPT-4o
(Interviewer: GPT-4o)

1. Performance Scores(%)

- **Precision(@ # of interaction):**
96.8(@1), 98.8(@2), 99.2(@3)
- **Follow-up Accuracy (Additional Facts):**
 - Total 92

2. Response Quality & Examples

- * Score for Each Type
- * Examples are omitted

- completeness.: 83.5
- redundancy: 64.9
- readability: 96.9
- depth: 64.9

3. Summary

The model excels in providing detailed, structured, and accurate explanations, particularly in technical, scientific, and mathematical domains. It effectively breaks down complex concepts into understandable parts and maintains coherence and continuity in responses, making it suitable for educational purposes. The model handles follow-up questions well and positively acknowledges user feedback, although it may not dynamically adapt based on feedback within a single session. **However, the model struggles with conciseness, often providing overly detailed responses that can overwhelm users seeking brief answers.** It may also have difficulty with highly specialized or nuanced queries, maintaining context over multiple interactions, and handling ambiguous or poorly defined questions. Additionally, the model may not always incorporate the latest information or trends beyond its training data. Overall, the model demonstrates strong capabilities in delivering comprehensive and accurate information but could improve in providing concise answers, handling more abstract or context-specific queries, and better incorporating user feedback.

DepthQA, Interviewee: Llama-3.1-8b
(Interviewer: GPT-4o)

1. Performance Scores(%)

- **Precision(@ # of interaction):**
89.8(@1), 98.1(@2), 99.6(@3)
- **Follow-up Accuracy (Additional Facts):**
 - Total 80

2. Response Quality & Examples

- * Score for Each Type
- * Examples are omitted

- completeness.: 86.3
- redundancy: 59.1
- readability: 90.9
- depth: 81.8

3. Summary

The model excels in providing detailed, structured, and comprehensive explanations, particularly in scientific, technical, and mathematical fields. It effectively breaks down complex topics into understandable segments and uses examples to enhance clarity. The model handles follow-up questions well, maintaining coherence and continuity, and incorporates user feedback constructively to refine its responses. **However, the model struggles with highly specialized or niche topics, ambiguous queries, and may initially provide overly detailed or redundant information. It sometimes requires user prompts to provide deeper insights,** maintain conciseness, and ensure technical accuracy. Additionally, it may face challenges with real-time data, subjective questions, and practical complexities without further user guidance.

1246

C.3 Feedback Example

1247

See Table 8.

1248

C.4 Human Evaluation of LLM-as-an-Interviewer

This section details the evaluation process and criteria for validating the three roles of **Interviewer** models. Four authors conducted the human evaluation to systematically assess the effectiveness of each model in three key dimensions: (1) *Query Modification*, (2) *Feedback Generation*, and (3) *Follow-up Question Generation*.

In this study, we evaluate GPT-4o, Llama-3.1-70B, and Llama-3.1-8B as interviewers, while using GPT-4o, GPT-3.5, Llama-3.1-70B and Llama-3.1-8B as interviewees to simulate the interview with models of varying capabilities. We validate at least 100 samples per dimension, resulting in over 300 samples per model. This sample size is comparable to or larger than the human-annotated datasets used in prior studies (Wang et al., 2023; Yu et al., 2024).

Table 10 presents the evaluation results, showing the accuracy of each model’s responses along with the number of evaluated samples in parentheses.

C.4.1 Evaluation Criteria

Each model’s responses were evaluated to determine whether they correctly followed the intended prompts and adhered to the reference solution. Annotators referenced both the problem statement and the reference solution throughout the evaluation process. The following criteria outline the assessment methodology:

Query Modification For MATH tasks, we assessed whether the query retained only the key information necessary for solving the problem, avoiding unnecessary modifications that might alter the problem’s complexity. For DepthQA, we evaluated whether the query was reformulated into a well-structured and solvable question aligned with the reference solution.

Feedback Generation We evaluated whether the feedback was contextually appropriate to the interviewee’s prior response. The feedback had to avoid contradictions with the reference solution while also refraining from directly revealing the correct answer. Annotators additionally reviewed the model’s intermediate reasoning (chain-of-thought) to ensure logical coherence.

Follow-up Question Generation Follow-up questions were assessed based on their alignment

with the problem type and intent. They needed to be relevant to the previous discussion, logically structured, and free from redundancy. For STEM-related questions, we ensured that follow-up queries were grounded in the reference solution and contributed meaningfully to the progression of the interview.

By adopting this structured evaluation framework, we ensure a comprehensive and transparent assessment of **Interviewer** models, allowing for a clear comparison of their capabilities across different interview phases.

Feedback	Example
Knowledge/Factuality	Your response effectively outlines the purpose and communication methods of the AQI and is clearly articulated. To improve further, ensure that the AQI scale range examples are accurate and align with commonly accepted values. Additionally, providing a more detailed explanation of the AQI calculation and data collection process would enhance the depth and completeness of your answer.
Reasoning	Your initial steps show some understanding of the order of operations, but you incorrectly simplified the expression inside the brackets. Remember to first evaluate the power, then apply the subtraction , and follow through with the multiplication and division before finally adding the constant outside the brackets.

Table 8: Examples of Feedback posed by the interviewer (GPT-4o).

	Original	Modified
MATH	If a recipe for a two-pound cake requires 1.5 cups of flour, how many cups are needed for 2 five-pound cakes?	If a recipe for a two-pound cake requires y cups of flour, how many cups are needed for 2 five-pound cakes?
DepthQA	Question What are the properties of straight lines in geometry? Reference Solution 1. A straight line is the shortest distance (...) 8. Parallel lines are always the same distance apart and never meet.(...) 10. Two lines on a plane that never meet are called parallel lines. 11. Two lines that intersect at a right angle (90 degrees) are called perpendicular lines.	Question What are parallel and perpendicular lines in geometry? Reference Solution 1. A straight line is the shortest distance (...) 8. Parallel lines are always the same distance apart and never meet.(...) 10. Two lines on a plane that never meet are called parallel lines. 11. Two lines that intersect at a right angle (90 degrees) are called perpendicular lines.

Table 9: Example of Query Modification. Changed parts are colored blue in the original and purple in the modified question.

Interview Phase	Interviewer Model		
	GPT-4o	Llama-3.1-70B	Llama-3.1-8B
Query Modification	89.5% (105)	84% (100)	73% (100)
Feedback Generation	85.8% (162)	73.5% (102)	58.1% (105)
Follow-up Question Generation	93.8% (145)	85% (100)	62% (100)

Table 10: Each Interviewer’s Performance Across Different Phases (%)

D Prompt⁹

D.1 Preparing Seed Questions

MATH

You are given a mathematical expression where key information must be removed, but the remaining structure and operations should not change. Only remove critical data that makes the problem unsolvable without that information.

Guidelines:

1. Remove only key information required for the solution (e.g., constants, values) without changing the operations and convert to unknown variables or ambiguous words.
2. Ensure the revised question retains the same structure and other information.
3. Avoid removing trivial or subtle details.
4. Do not modify the mathematical operations or alter the structure of the equation.
5. In the explanation, detail exactly how and where deleted information fits into the equation or text. For example, explain which unknown variable or ambiguous words denote what value.

Response format must be in JSON as shown below:

Example 1: {Example}

Example 2: {Example}

Input Question: {Question}

Output:

DepthQA

You are provided with a question and its solution. Your task is to create a unique question based on these.

Instructions:

1. Create one new question that is distinct from the original question:
 - The new question should be directly answerable using the provided solution.
 - It should avoid focusing on overly minor or trivial details.
2. Ensure the new question aligns well with the given solution and maintains relevance to the original context.

Output Format: {Format Example}

Input:

Question: {Question}

Solution: {Solution}

1314

D.2 Interview Process

1315

D.2.1 Feedback Generation

1316

The feedback generation follows two steps approach. First, the response is graded (using the Grader Prompt below), and then feedback for the interviewee is generated based on that grading (using the Feedback Generator Prompt below).

1317

1318

1319

1320

1321

MATH Grader

Task: You are provided with a model's output, a reference solution, and a correct answer. Your goal is to determine the correctness of the model's output by comparing it with the reference solution and the correct answer. Note that the reference solution is just a guide—there could be other valid methods to solve the problem.

Possible Error Types:

1. **Concept:** This error type indicates the model lacks the concept used to solve the problem. In this case, the feedback should contain a question that can check the understanding of the concept. (e.g., What is the Pythagorean theorem?)
2. **Misinterpret:** The model misunderstood or misinterpreted the question. In this case, the feedback should include a follow-up question or clarification that helps the model reassess and correctly interpret the original question. This ensures the model understands the context and requirements before proceeding.
3. **Calculation:** The model made a mistake in calculation.
4. **N/A:** None of the above.

Response Format: {Format Example}

Input:

Initial question: {Question}

Correct answer: {answer}

Reference solution: {Solution}

Model's Output: history

1322

⁹Not all prompts are included here. Refer to our GitHub repository for additional prompts.

MATH Feedback Generator

Refer to the pre-generated evaluation on the given question below and generate feedback for the model.

Your response must include: 1. Feedback: Provide the model with concise, constructive guidance based on the evaluation and reference solution.

2. Feedback Type: Choose one of the following:

- Conceptual Guidance: Focus on understanding key concepts.
- Error Identification and Correction: Address specific mistakes and guide corrections.
- Process and Strategy Guidance: Improve the model's approach or strategy.
- Precision and Accuracy Emphasis: Stress precision in calculations or answer format.
- Encouragement and Affirmation: Motivate and reinforce correct actions.

Guidelines: 1. Do not reveal the solution. Guide the assistant toward understanding the correct approach.

2. Ensure feedback is unique, avoiding repetition.

3. Provide progressively more specific hints without focusing on trivial issues.

4. If the model does not seem to understand the question, explain what the question is.

Response format must be in JSON as shown below: {example}

Input:

Question: {Question}

Reference Solution (DO NOT disclose): {Solution}

Correct Answer (DO NOT disclose): {answer}

Pre-generated Evaluation: evaluation

Previous Feedback: model history

Model Output (to evaluate): model output

Output:

model has revised the incorrect parts of its original output, and these revisions are provided in the model's correction statement.

Your task is to update the model's original output by replacing the corresponding revised facts with those found in the model's correction statement. Ensure that updated facts are from the model's correction statement, not the reference solution. For facts that were not revised, maintain the same content and correctness as before, so that your output contains the same number of atomic facts as the model's original output.

Output Format: {Format Example}

Input:

Question: {question}

Solution: {solution}

Previous Feedback: {feedback}

Model's original output with correctness: {history}

Model's correction statement: {correction}

Output:

STEM Feedback Generator

You are an expert tasked with evaluating and providing feedback on an assistant's performance. Your response should contain the following elements:

1. **"status"**: Indicate whether the answer is "complete" (fully correct), "partially_correct" (some elements are correct but not all), or "incorrect" (entirely wrong).

2. **"feedback"**: Provide specific, constructive feedback based on the differences between the correct solution and the model's output. Your feedback should:

- Highlight what aspects of the answer are correct (if any).
- Identify specific areas where the model's output differs from the correct solution, without revealing the entire solution.
- Offer guidance or hints that can help the model improve its answer, focusing on the concepts or steps that need correction.
- If the process is incorrect but the final answer is right, explain that the reasoning needs improvement.
- Encourage the model to think step-by-

STEM Grader

Task:

You are given a question, a reference solution, and the model's output, which is broken down into atomic fact units, along with their correctness and justification. The

1325

1326

step and retry if necessary.

Remember: - Do not give away the complete solution or tell exactly which step is incorrect.

- If it's not the first attempt, provide more detailed hints, such as mentioning a relevant equation or concept, without repeating previous feedback.

- Tailor your feedback to the specific errors or misconceptions evident in the model's output.

- If there is a lack in at least one aspect among completeness, redundancy, readability, or depth, consider it incomplete and provide feedback on the identified shortcomings.

Examples of the expected format and style of feedback: {Examples}

Input:

Question: question

Reference Answer (DO NOT disclose this to the assistant): answer

Correct Solution (DO NOT disclose this to the assistant): solution

Previous Feedbacks: model_history

Model Output (This is the part to be evaluated): model_output

Ensure your feedback is unique and does not repeat previous feedback.

Expert Feedback:

D.2.2 Follow-Up Question Generation

MATH Fail:

If the seed question is answered incorrectly, generate a clarification question based on the model's error type. The question should focus on the concept or interpretation of the seed question to help identify the source of the misunderstanding.

MATH Success:

If the seed question is answered correctly, generate a follow-up rationale question to assess the model's understanding of the reasoning behind its solution.

DepthQA: Regardless of whether the seed question is answered correctly or not, generate a follow-up question based on the reference material. The question should inquire about additional facts or details that the model did not address in its response.

MATH Fail

Task:

You are a question generator for assessing the model. Below are the list of the answer history of an Evaluatee model that keeps getting the question wrong, its error type, the corresponding problem, and the solution. Your role is to generate a proper follow-up question based on the model's error type.

Instructions: 1. Model's Error Type is {error_type}. Do not disclose the error type to the model.

2. Based on the error type, generate appropriate feedback as described below.

Possible Error Types:

1. **Concept:** This error type indicates the model lacks the concept used to solve the problem. The feedback should contain a question to assess the understanding of the concept. (e.g., What is the Pythagorean theorem?)

2. **Misinterpret:** The model misunderstood or misinterpreted the question. The feedback should include a follow-up question or clarification to help the model reassess and correctly interpret the original question. This ensures the model understands the context and requirements before proceeding.

3. **Calculation:** The model made a mistake in calculation. Feedback should focus on identifying the part that requires recalculation.

4. **N/A:** None of the above. Provide a custom description.

Response Format (JSON): {Format Example}

Input Data:

Question: {question}

Correct Answer (DO NOT disclose): {answer}

Reference Solution (DO NOT disclose): {solution}

Previous Feedback: {Dialogue_History}

Model's Error Type: {error_type}

Output:

MATH Success

Task:

You are an evaluator assessing the Evaluatee model. The Evaluatee has successfully answered the problem, and the following is the Evaluatee's solution. Your role is to determine whether the Evaluatee truly understands their solution or if they arrived at the answer through memorization without proper understanding. To do this, you should ask questions that can assess the model's understanding of its own solution.

Guidelines: 1. If there are errors or missing steps in the Evaluatee's solution, ask questions to clarify or correct these issues.
2. If there are no errors in the solution, ask questions to confirm the Evaluatee's understanding of the solution steps.
3. Generate one question and provide the answer in JSON format. Be sure not to ask questions that are already present in the previous history below.
4. Create the question solely based on the model's solution.

Response Format (JSON): {Format Example}

Input Data:

Initial Question: {initial_question}

Correct Answer (DO NOT disclose): {answer}

Reference Solution (DO NOT disclose): {solution}

Model's Solution: {model_solution}

Output:

Reference Answer that were labeled as unsupported by the model's output.

2. The follow-up questions should be designed to test whether the model understands the unsupported facts.

3. Never include any facts from the reference solution to the question. Instead, ask questions to verify knowledge.

4. The questions should indirectly reference the concepts in a way that allows us to assess the model's understanding.

5. Your follow-up questions should not simply ask the fact directly but should guide the model to demonstrate whether it knows the fact.

Output Format (JSON): {Format Example}

Input Data:

Question: question

Reference Answer: solution

Output:

1349

D.3 Summarizing Interview

1350

Session Summarize Prompt

Task:

The following is a conversation between user and system. Summarize the conversation.

Instructions:

1. Instead of focusing on the specific details of the questions and answers, provide a summary that highlights:

- The overall flow of the conversation.
- system's problem-solving abilities.

Input:

Session History: {session_history}

Output:

Summary:

1351

DepthQA

Task:

You are given three inputs: 1. **Question:** The original question that was asked.

2. **Reference Answer:** A list of atomic facts with labels indicating whether each fact is supported or unsupported by a model's output.

3. **Model's Output:** The model's output and the corresponding correctness for the given question.

Instruction:

1. Your task is to generate follow-up questions based on the atomic facts from the

Summarize Prompt for Interview Report

Summarize the following problem summaries of the system’s problem-solving abilities. In your summary, try to provide **general insights** into the capabilities of the model, such as:

Strength and weakness of the model

Does the model respond well to the user’s follow-up questions?

Does the model handle the user’s feedback effectively?

What types of information or tasks does the model handle well?

What types of problems or details does the model struggle with?

Answer the following questions one by one. Do **not** focus on specific examples, but rather offer a general overview of the model’s strengths and weaknesses.

Here are the system’s problem-solving abilities to summarize. Don’t generate too long: {chunk_dict}

E Evaluating LLMs with LLM-as-an-Interviewer

E.1 Evaluation Result of MINT

As highlighted in the Related Work section, prior studies of multi-turn dynamic evaluation are designed with distinct objectives and methodologies, making direct comparisons difficult. Among the available benchmarks, we select **MINT** (Wang et al., 2023), which evaluates *Multi-turn Interaction with Tool usage*, as it aligns closely with our focus. Specifically, we choose the **Reasoning** subset of MINT, which excludes external tools or code execution and consists of **316 samples** sourced from *GSM8K*, *MATH*, *HotpotQA*, *TheoremQA*, and *MMLU*.

For this evaluation, we leverage our **LLM-as-an-Interviewer framework**, implemented as a **PyPI module**, showcasing its applicability beyond datasets like *MATH* and *DepthQA*.

E.1.1 Evaluation

We assess six models and compare three models—Llama-{8B/70B} and GPT-3.5-Turbo—that have prior results reported in MINT. To maintain consistency in measuring accuracy **without feedback or interaction**—comparing

Acc_seed(1) from our framework with k=1 from MINT—we do not apply any query modifications.

Table 11 presents the performance assessed by the LLM-as-an-Interviewer framework, while Table 12 shows the results from the MINT framework. Discrepancies between **Acc_seed(1) and k=1** may arise due to differences in *tool usage and prompt structures*. MINT employs *longer initial prompts* that explicitly guide tool usage and define output formats, potentially influencing accuracy at $k=1$.

E.2 Evaluation Result of MT-bench

We compare performance on MT-Bench (Chiang et al., 2024), a 2-turn static benchmark with 80 samples across various tasks (e.g., writing, reasoning, math, coding, roleplay). For the baseline, model responses are evaluated sequentially using turn = 1 and turn = 2 queries from the dataset. In the **Interviewer** setting, turn = 1 queries serve as seed questions without modification, followed by follow-up questions posed by the **Interviewer** (turn = 2). The scores range from 1 to 10 (refer to the method in (Chiang et al., 2024)).

E.3 Follow-Up Question Example of DepthQA

In the following example, the follow-up question targets the n-th roots mentioned in the reference solution but missing from the model’s previous response. This type of follow-up question is central to DepthQA. It partially address a key limitation of traditional factuality scoring, which only evaluates the facts the model provides (i.e., precision). By probing further with additional questions, we can assess the model’s broader factuality, capturing missing or incomplete information.

Model	Judge	Acc_seed(1)	Acc_seed(2)	Acc_seed(3)	Δ
GPT-4o	0.753	0.753	0.832	0.858	0.105
GPT-3.5-turbo	0.467	0.467	0.644	0.717	0.250
Llama-3.1-70b	0.648	0.648	0.759	0.794	0.146
Llama-3.1-8b	0.537	0.537	0.673	0.733	0.196
DeepSeek-math	0.155	0.155	0.389	0.509	0.354
Qwen-math	0.215	0.215	0.304	0.342	0.127

Table 11: **Performance of different models using LLM-as-an-Interviewer Framework in MINT reasoning dataset.** Δ represents the difference between Acc_seed(3) and Acc_seed(1).

Model	k = 1	k = 2	k = 3	k = 5	Δ
GPT-3.5-turbo	0.0316	0.117	0.323	0.535	0.2914
Llama-3.1-70b	0.370	0.411	0.484	0.658	0.114
Llama-3.1-8b	0.282	0.358	0.364	0.377	0.082

Table 12: **Performance of different models using MINT Framework in MINT reasoning dataset.** K denotes the number of interactions. Δ represents the difference between k = 3 and k = 1.

Model	Baseline		Interviewer	
	1	2	1	2
GPT4o	9.08	8.30	9.03	8.52
GPT-3.5-turbo	8.40	7.54	8.38	8.18
Llama-3.1-70B	8.70	8.08	8.70	8.36
Llama-3.1-8B	8.52	7.58	8.48	7.75
Deepseek-Math-7B	5.68	5.56	5.51	5.71
Qwen-Math-7B	4.91	2.73	4.6	4.9

Table 13: Performance comparison of models on MT-Bench and the Interviewer setting.

DepthQA Follow-up Question Example

User: What is the significance of the roots of unity in mathematics and physics?

System: The roots of unity are important in mathematics and physics because they have connections to many different areas of study. In mathematics, they are used in algebra, number theory, and complex analysis. In physics, they appear in wave functions, quantum mechanics, and signal processing. The roots of unity also have applications in computer science, cryptography, and engineering. Overall, they play a key role in understanding the behavior of periodic functions and symmetries in various mathematical and physical systems.

Follow-up Question

User: Can you provide an example of determining the **6th roots of unity** and specifying their arguments in radians?

Reference Solution

The roots of unity refer to the complex solutions of the equation $z^n = 1$, where n is a positive integer. These roots are regularly spaced around the circumference of the unit circle in the complex plane and hence, are of significant importance in many areas of mathematics and physics.

The term "unity" signifies the number 1, and "roots" refer to the solutions of an equation. So, the roots of unity are essentially the solutions to the equation that equal 1 when raised to the power n .

Each n -th root of unity is given by the formula: $z = \cos\left(\frac{2\pi k}{n}\right) + i \sin\left(\frac{2\pi k}{n}\right)$

for $k = 0, 1, \dots, n - 1$. Here, $\cos$ and $\sin$ are the standard trigonometric functions, i refers to the imaginary unit, and π is the mathematical constant Pi (approximately equal to 3.14).

Let's say if you wanted to determine the 4th roots of unity, (omitted)

F Full Interview Log

MATH Full Log

Interviewee : Llama-3.1-70b, Interviewer : GPT-4o

User: Question: The number $\sqrt{x\sqrt{6}} + y\sqrt{10} + z\sqrt{15} + w$ can be written as $a\sqrt{2} + b\sqrt{3} + c\sqrt{5}$, where a, b , and c are positive integers. Find $a \cdot b \cdot c$.

The given expression can be written as $a\sqrt{2} + b\sqrt{3} + c\sqrt{5}$,

where a, b , and c are positive integers.

We are given the expression:

$$\sqrt{x\sqrt{6}} + y\sqrt{10} + z\sqrt{15} + w$$

Squaring the expression to get rid of the outer square root:

$$(x\sqrt{6} + y\sqrt{10} + z\sqrt{15} + w)^2 = (a\sqrt{2} + b\sqrt{3} + c\sqrt{5})^2$$

Expanding the right side:

$$(x\sqrt{6} + y\sqrt{10} + z\sqrt{15} + w)^2 = 2a^2 + 3b^2 + 5c^2 + 2ab\sqrt{6} + 2ac\sqrt{10} + 2bc\sqrt{15} + 2abc\sqrt{30}$$

Comparing the coefficients of $\sqrt{6}$, $\sqrt{10}$, and $\sqrt{15}$:

$$x = 2ab, y = 2ac, z = 2bc$$

Also, the constant term w is:

$$w = 2a^2 + 3b^2 + 5c^2$$

We know that a, b , and c are positive integers.

From the coefficients of $\sqrt{6}$, $\sqrt{10}$, and $\sqrt{15}$:

$$x = 2ab, y = 2ac, z = 2bc$$

We can see that x, y , and z can be expressed as multiples of 2.

Also, the constant term w is:

$$w = 2a^2 + 3b^2 + 5c^2$$

Since x, y , and z are multiples of 2, w must be a multiple of 2 as well.

The least possible values for a, b , and c that satisfy the conditions are:

$$a = 1, b = 1, c = 1$$

Substituting these values in:

$$x = 2ab = 2(1)(1) = 2$$

$$y = 2ac = 2(1)(1) = 2$$

$$z = 2bc = 2(1)(1) = 2$$

$$w = 2a^2 + 3b^2 + 5c^2 = 2(1)^2 + 3(1)^2 + 5(1)^2 = 10$$

However, $w = 10$ is not a multiple of 2 that satisfies the conditions.

User: You got a correct answer! Good Job.

I'm glad I was able to solve the problem correctly in the end.

It was a bit of a journey, but using the exact values of 104, 468, 144, and 2006 for x, y, z , and w respectively, helped to simplify the problem and arrive at the correct solution.

If you have any more problems you'd like help with, feel free to ask!

User: How do you determine the values of 'a', 'b', and 'c'?

from the equations $ab = 52$, $ac = 234$, and $bc = 72$?

To determine the values of 'a', 'b', and 'c' from the equations

$ab = 52$, $ac = 234$, and $bc = 72$, we can use the following steps:

1. Divide the equation $ac = 234$ by the equation $ab = 52$:

$$\frac{ac}{ab} = \frac{234}{52}$$

2. Multiply both sides by b to get:

$$c = \frac{9}{4}b$$

3. Now, substitute this expression for c into the equation $bc = 72$:

$$b \cdot \left(\frac{9}{4}b\right) = 72$$

$$\frac{9}{4}b^2 = 72$$

4. Multiply both sides by $\frac{4}{9}$ to get:

$$b^2 = 72 \cdot \frac{4}{9}$$

$$b^2 = 32$$

5. Take the square root of both sides:

$$b = \sqrt{32}$$

$$b = 4 \text{ (since } b \text{ is a positive integer)}$$

6. Now that we have the value of b , we can find the value of a using the equation $ab = 52$:

$$a = \frac{52}{b}$$

$$a = \frac{52}{4}$$

$$a = 13$$

7. Next, find the value of c using the equation $ac = 234$:

$$c = \frac{234}{a}$$

$$c = \frac{234}{13}$$

$$c = 18$$

These values satisfy the original equations $ab = 52$, $ac = 234$, and $bc = 72$.

User: Your explanation is generally correct, but there are some inconsistencies in your steps.

Specifically, your calculation of b and the adjustment for c requires clarification.

You initially found $b = 4$, which is correct, but then multiplied b by 2 without clear reasoning.

Finally, you identified the correct values, but the intermediate explanations need to be streamlined for consistency. Could you revisit and clarify these steps to arrive at a consistent conclusion?

You are correct, my initial explanation had some inconsistencies.

Here's a revised explanation for determining the values of 'a', 'b', and 'c' from the equations $ab = 52$, $ac = 234$, and $bc = 72$:

1. Divide the equation $ac = 234$ by the equation $ab = 52$:

$$\frac{ac}{ab} = \frac{234}{52}$$

$$\frac{c}{b} = \frac{9}{4}$$

2. Multiply both sides by b to get:

$$c = \frac{9}{4}b$$

3. Now, substitute this expression for c into the equation $bc = 72$:

$$b \cdot \left(\frac{9}{4}b\right) = 72$$

$$\frac{9}{4}b^2 = 72$$

4. Multiply both sides by $\frac{4}{9}$ to get:

$$b^2 = 72 \cdot \frac{4}{9}$$

$$b^2 = 32$$

5. Take the square root of both sides:

$$b = \sqrt{32}$$

$b = 4$ (since b is a positive integer)

6. Now that we have the value of b , we can find the value of a using the equation $ab = 52$:

$$a = \frac{52}{b}$$

$$a = \frac{52}{4}$$

$$a = 13$$

7. Next, find the value of c using the equation $ac = 234$:

$$c = \frac{234}{a}$$

$$c = \frac{234}{13}$$

$$c = 18$$

This way, we arrive at the correct values for a , b , and c in a consistent and straightforward manner:

$$a = 13, b = 4, c = 18.$$

Thank you for pointing out the inconsistencies in my initial explanation.

I hope this revised explanation helps clarify the steps.

Interviewee : Llama-3.1-70b, Interviewer : GPT-4o

User: Question: What are some scenarios where the forward or backward difference approximations are preferred over the central difference approximation?

System: There are several scenarios where the forward or backward difference approximations are preferred over the central difference approximation:

1. ****Boundary conditions**:** When approximating the derivative at a boundary point, the central difference approximation is not feasible because it requires points on both sides of the boundary. In such cases, the forward or backward difference approximation is used, depending on whether the boundary point is at the beginning or end of the interval.
2. ****Non-uniform grids**:** When the grid points are not uniformly spaced, the central difference approximation may not be the most accurate choice. In such cases, the forward or backward difference approximation can be used, taking into account the non-uniform spacing of the grid points.
3. ****Noise or irregular data**:** When the data is noisy or irregular, the central difference approximation may amplify the noise or irregularities. In such cases, the forward or backward difference approximation can be used to reduce the effect of noise or irregularities.
4. ****Stability and convergence**:** In some numerical methods, such as finite difference methods for solving partial differential equations, the forward or backward difference approximation may be preferred over the central difference approximation for stability and convergence reasons.
5. ****Computational efficiency**:** In some cases, the forward or backward difference approximation may be computationally more efficient than the central difference approximation, especially when the number of grid points is large.

6. ****Advection-dominated problems**:** In problems where advection dominates diffusion, the backward difference approximation may be preferred over the central difference approximation to reduce numerical diffusion.

7. ****Time-dependent problems**:** In time-dependent problems, the backward difference approximation may be preferred over the central difference approximation to ensure stability and convergence.

In summary, the choice of difference approximation depends on the specific problem, the nature of the data, and the requirements of the numerical method.

User: Your explanation covers several important scenarios where forward or backward difference approximations might be preferred, but some areas need more detailed context and justification. For instance, could you explain when grid points are not uniformly spaced and why central difference may not be accurate there? Also, clarify how the central difference method might amplify noise or irregularities compared to other methods. Additionally, elaborating on why the backward difference might be used in advection-dominated problems and the suitability of backward difference for time-dependent problems would strengthen your response. Overall, consider diving deeper into these points to improve the completeness and depth of your answer.

System: I'd be happy to provide more detailed explanations for the scenarios where the forward or backward difference approximations are preferred over the central difference approximation.

****Non-uniform grids****

When the grid points are not uniformly spaced, the central difference approximation may not be the most accurate choice. This is because the central difference formula assumes that the grid points are evenly spaced, and it uses the average of the function values at the neighboring points to approximate the derivative. However, when the grid points are not uni-

formly spaced, this assumption is no longer valid.

For example, consider a scenario where the grid points are spaced more closely near the boundary of the domain than in the interior. In this case, the central difference formula would use the same weighting for the function values at the neighboring points, even though the points near the boundary are more closely spaced. This can lead to inaccurate results, especially if the function is rapidly changing near the boundary.

In contrast, the forward or backward difference approximations can be adapted to handle non-uniform grids by taking into account the spacing between the grid points. For example, the forward difference formula can be modified to use the distance between the current point and the next point, rather than assuming a uniform spacing.

****Noise or irregular data:****

When the data is noisy or irregular, the central difference approximation may amplify the noise or irregularities. This is because the central difference formula uses the difference between the function values at the neighboring points, which can accentuate any noise or irregularities in the data.

For example, consider a scenario where the data has a high-frequency noise component. The central difference formula would use the difference between the function values at the neighboring points, which would amplify the noise component. In contrast, the forward or backward difference approximations would use the difference between the current point and the next or previous point, respectively, which can help to reduce the effect of the noise.

Additionally, the central difference formula can also be sensitive to outliers or irregularities in the data. For example, if there is a single outlier in the data, the central difference formula would use this outlier in the calculation of the derivative, which can lead to inaccurate results. In contrast, the forward or backward difference approximations can be more robust to outliers, as they only use the current point and the next

or previous point, respectively.

****Advection-dominated problems:****

In problems where advection dominates diffusion, the backward difference approximation may be preferred over the central difference approximation to reduce numerical diffusion. Advection-dominated problems are characterized by a strong convective flow, which can lead to numerical diffusion when using the central difference approximation.

Numerical diffusion is a phenomenon where the numerical solution of the problem becomes smeared or diffused, even though the underlying physics of the problem does not involve diffusion. This can happen when the central difference approximation is used, as it can introduce artificial diffusion into the solution.

The backward difference approximation, on the other hand, can help to reduce numerical diffusion by using the upwind direction to calculate the derivative. The upwind direction is the direction in which the flow is moving, and using this direction can help to reduce the artificial diffusion introduced by the numerical method.

For example, consider a problem where a pollutant is being transported by a fluid flow. The central difference approximation would use the average of the pollutant concentrations at the neighboring points to calculate the derivative, which can lead to numerical diffusion. In contrast, the backward difference approximation would use the pollutant concentration at the upwind point to calculate the derivative, which can help to reduce numerical diffusion.

****Time-dependent problems:****

In time-dependent problems, the backward difference approximation may be preferred over the central difference approximation to ensure stability and convergence. Time-dependent problems involve solving a partial differential equation that changes over time, and the numerical method must be able to capture the time-dependent behavior of the solution.

The backward difference approximation is often preferred in time-dependent prob-

lems because it is unconditionally stable, meaning that it will converge to the correct solution regardless of the time step size. The central difference approximation, on the other hand, can be conditionally stable, meaning that it will only converge to the correct solution if the time step size is small enough.

For example, consider a problem where a wave is propagating through a medium over time. The central difference approximation would use the average of the wave amplitude at the neighboring points to calculate the derivative, which can lead to instability if the time step size is too large. In contrast, the backward difference approximation would use the wave amplitude at the previous time step to calculate the derivative, which can help to ensure stability and convergence.

In summary, the choice of difference approximation depends on the specific problem, the nature of the data, and the requirements of the numerical method. The forward or backward difference approximations can be preferred over the central difference approximation in scenarios such as non-uniform grids, noisy or irregular data, advection-dominated problems, and time-dependent problems."

User: Your output provides a comprehensive and detailed explanation of various scenarios where forward or backward difference approximations may be preferred over central difference. However, consider if the preference for specific difference approximations might depend on factors beyond stability and convergence, and reflect on whether computational efficiency is a significant deciding factor. Additionally, some redundancy in the discussion of non-uniform grids and noisy data could be consolidated to improve conciseness. Overall, the depth and readability of your explanation are strong, but ensure all points are necessary to the question.

System: I'd be happy to refine my explanation to address your points and provide a

more concise response.

****Non-uniform grids and noisy data:****

In scenarios where the grid points are not uniformly spaced or the data is noisy or irregular, the forward or backward difference approximations may be preferred over the central difference approximation. The central difference formula assumes uniform spacing and can be sensitive to noise or irregularities, which can lead to inaccurate results. In contrast, the forward or backward difference approximations can be adapted to handle non-uniform grids and reduce the effect of noise or irregularities. For example, in problems with high-frequency noise or outliers, the forward or backward difference approximations can provide a more robust estimate of the derivative by using the difference between the current point and the next or previous point, respectively. Similarly, in non-uniform grids, the forward or backward difference approximations can take into account the varying spacing between grid points to provide a more accurate estimate of the derivative.

****Advection-dominated problems:****

In problems where advection dominates diffusion, the backward difference approximation may be preferred over the central difference approximation to reduce numerical diffusion. The central difference approximation can introduce artificial diffusion into the solution, which can be mitigated by using the backward difference approximation and the upwind direction to calculate the derivative.

****Time-dependent problems:****

In time-dependent problems, the backward difference approximation may be preferred over the central difference approximation to ensure stability and convergence. The backward difference approximation is unconditionally stable, whereas the central difference approximation can be conditionally stable, requiring a small enough time step size to converge.

****Computational efficiency:****

While computational efficiency can be a consideration in choosing a difference ap-

proximation, it is often secondary to stability and convergence. In some cases, the forward or backward difference approximations may be computationally more efficient than the central difference approximation, especially when the number of grid points is large. However, this advantage may be outweighed by the need for stability and convergence in many applications.

****Additional considerations:****

Beyond stability and convergence, the choice of difference approximation may depend on factors such as:

* ****Boundary conditions:**** The forward or backward difference approximations may be preferred at boundary points, where the central difference approximation is not feasible. * ****Problem physics:**** The choice of difference approximation may depend on the specific physics of the problem, such as the dominance of advection or diffusion. * ****Numerical method:**** The choice of difference approximation may be influenced by the overall numerical method being used, such as finite difference or finite element methods.

In summary, the choice of difference approximation depends on a range of factors, including stability, convergence, computational efficiency, and problem-specific considerations. While the central difference approximation is often a good choice, the forward or backward difference approximations may be preferred in scenarios with non-uniform grids, noisy data, advection-dominated problems, or time-dependent problems.

User: You got a correct explanation! Good Job.

User: In the context of boundary conditions, how would you choose between forward and backward difference approximations at the start and end of the domain?

System: When choosing between forward and backward difference approximations at the start and end of the domain, it ultimately depends on the specific boundary

conditions and the direction of information flow. Here's a general guideline:

****Start of the domain:****

* If the boundary condition at the start of the domain is a ****Dirichlet boundary condition**** (i.e., the value of the function is specified at the boundary), a ****forward difference approximation**** is often used. This is because the forward difference approximation looks ahead to the next point, which is inside the domain, to estimate the derivative. This approach is consistent with the direction of information flow, where the boundary condition is imposed from the outside and influences the behavior of the solution inside the domain.

* If the boundary condition at the start of the domain is a ****Neumann boundary condition**** (i.e., the derivative of the function is specified at the boundary), a ****backward difference approximation**** is not suitable, as it would look backwards outside the domain. In this case, a forward difference approximation or a ghost point method (where a fictional point is introduced outside the domain to apply the Neumann boundary condition) can be used.

****End of the domain:****

* If the boundary condition at the end of the domain is a ****Dirichlet boundary condition****, a ****backward difference approximation**** is often used. This is because the backward difference approximation looks backwards to the previous point, which is inside the domain, to estimate the derivative. This approach is consistent with the direction of information flow, where the boundary condition is imposed from the outside and influences the behavior of the solution inside the domain. * If the boundary condition at the end of the domain is a ****Neumann boundary condition****, a ****forward difference approximation**** is not suitable, as it would look forward outside the domain. In this case, a backward difference approximation or a ghost point method can be used.

****Additional considerations:****

* When using forward or backward difference approximations at the boundaries, it's

essential to ensure that the chosen method is consistent with the overall numerical method being used (e.g., finite difference or finite element methods). * In some cases, a **ghost point method** or a **one-sided finite difference** can be used to apply boundary conditions. These methods introduce fictional points outside the domain to apply the boundary conditions, which can be useful when dealing with complex boundary conditions or non-uniform grids. In summary, the choice between forward and backward difference approximations at the start and end of the domain depends on the specific boundary conditions and the direction of information flow. A forward difference approximation is often used at the start of the domain with Dirichlet boundary conditions, while a backward difference approximation is often used at the end of the domain with Dirichlet boundary conditions."

G Mitigating Contamination Issue with LLM-as-an-Interviewer

G.1 Training Details and Configurations

For training, we utilized four 80GB NVIDIA A100 GPUs. We trained Zephyr models using the implementation from the Alignment Handbook repository.¹⁰ We trained OLMoE models using the instructions provided on its GitHub repository.¹¹

Training hyperparameters are detailed in Table 15 and 16.

G.2 Experimental Results in DepthQA

		OLMoE		Zephyr	
<i>Model.</i>		Judge.	Interv.	Judge.	Interv.
Uncontam.					
Train set _{ID}	[1]	0.70	0.71	0.84	0.80
Avg.		0.70	0.71	0.84	0.80
Contam.					
Test set _{ID}	[4]	1.00	0.80	0.86	0.85
+Train set _{ID}	[5]	0.95	0.73	0.87	0.88
+Instruct	[6]	0.91	0.70	0.86	0.83
Avg.		0.93	0.76	0.87	0.85

Table 14: **Performance of uncontaminated and contaminated models on DepthQA dataset.** "Judge" refers to the LLM-as-a-Judge setting, and "Interv" refers to the LLM-as-an-Interviewer setting.

Table 14 shows the contamination experiment results in DepthQA.

Base Model	alignment-handbook/zephyr-7b-sft-full
Torch dtype	bfloat16
Epoch	10
Max Seq Length	4096
Learning Rate	1e-5
Train Batch Size	32
Random Seed	42
Training Method	Supervised Fine-tuning

Table 15: Hyperparameters used to train Zephyr 7B.

Base Model	allenai/OLMoE-1B-7B-0924-Instruct
Torch dtype	bfloat16
Epoch	10
Max Seq Length	4096
Learning Rate	2e-5
Train Batch Size	16
Training Method	Supervised Fine-tuning

Table 16: Hyperparameters used to train OLMoE-1B-7B model.

¹⁰Alignment Handbook Repository

¹¹OLMoE GitHub Repository

H Reliability of LLM-as-an-Interviewer

assigning higher scores to other models when acting as the interviewer.

H.1 Verbosity Bias

Table 3 presents the correlation between scores and interaction lengths across different stages of the interaction.

H.2 Self-Enhancement Bias

Existing papers primarily analyze the self-enhancement bias of LLM-as-a-Judge through relative rating—comparing the responses of two models and selecting the preferred one (Chiang et al., 2024). However, since LLM-as-an-Interviewer uses absolute rating—evaluating only a single model’s response—it is difficult to directly analyze self-enhancement bias. Therefore, we only check whether an apparently observable self-enhancement bias exists.

We analyze self-enhancement bias using GPT-4o, Llama-3.1-70B, and Llama-3.1-8B as both Interviewers and Interviewee. Figures 4-5 indicate that self-enhancement bias, defined as assigning higher scores to oneself, is observed only in Llama-3.1-70B among the models. In the MATH dataset, GPT-4o and Llama-3.1-8B tend to assign themselves slightly lower scores compared to other Interviewers, whereas Llama-3.1-70B assigns itself marginally higher scores. Despite this, the behavior has negligible impact on the overall rankings. As interactions progress, this bias diminishes further, and the score gap between GPT-4o and Llama-3.1-70B increases significantly. A similar pattern of diminishing bias and widening score differences is observed in the DepthQA dataset.

Figure 4 and 5 show the trend of self-enhancement bias in MATH and DepthQA. Similar to MATH (Figure 4), no significant self-enhancement bias is observed except for Llama-3.1-70b. Figure 6 illustrates the trends in MATH and DepthQA for follow-up questions. In MATH, as in previous results, Llama-3.1-70b consistently gives itself higher scores when grading its own responses compared to other models. However, Llama-3.1-70b also tends to give higher scores when grading other models, making it difficult to determine whether this is due to self-enhancement bias or because it generates easier follow-up questions compared to other models. In DepthQA, GPT-4 shows a similar pattern, giving itself higher scores when grading its own responses, but also

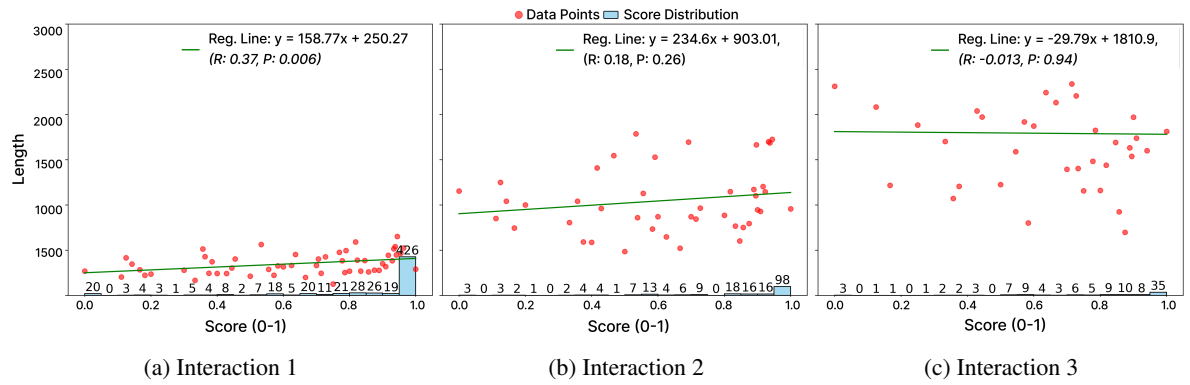


Figure 3: **Distribution of Response Lengths (measured in tokens) Across Scores in DepthQA.** The regression correlation coefficient (r) decreases with the number of Interaction, suggesting a reduction in **Verbosity Bias** as Interactions progress.

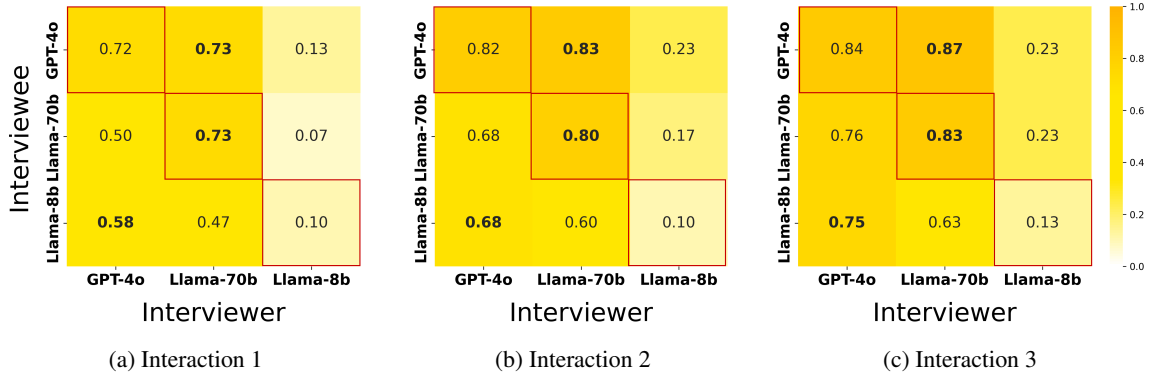


Figure 4: **Accuracy of Models as Interviewees by Number of Interactions in MATH.** The red-highlighted boxes indicate self-evaluated cases. Llama-70B consistently gives itself higher scores compared to other interviewers, but this does not result in a significant **Self-Enhancement Bias** that alters its ranking. No such bias is observed in the other models.

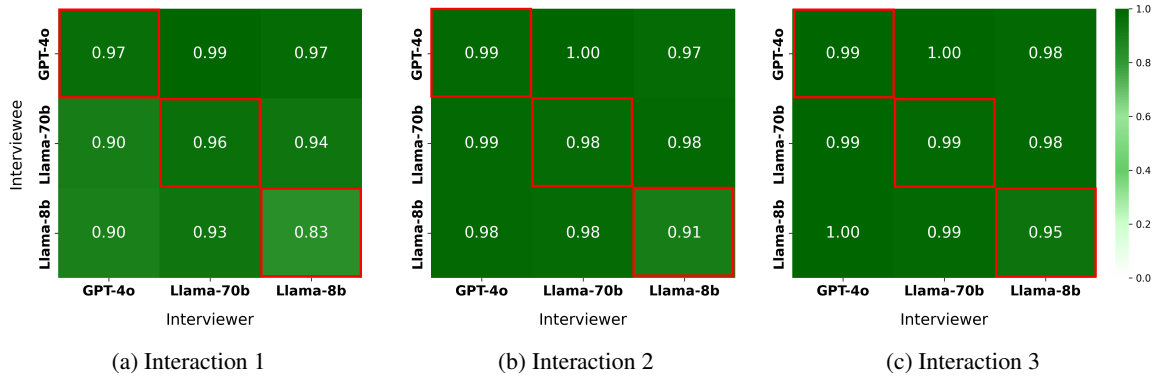


Figure 5: **Performance Score(Precision) by Number of Interactions in DepthQA.** The red-highlighted boxes indicate self-evaluated cases.

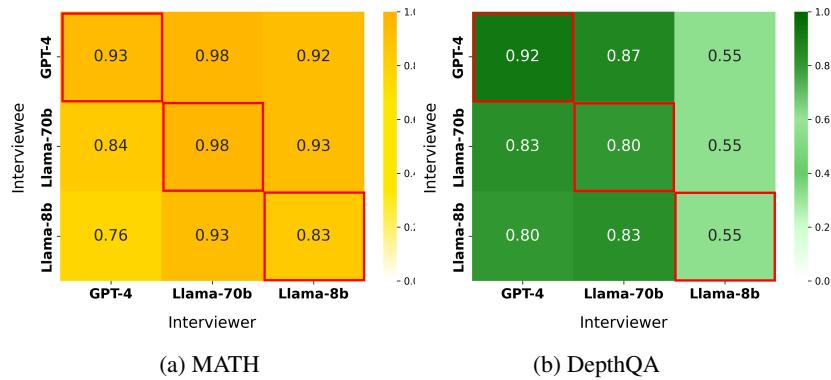


Figure 6: **Follow-Up Accuracy in MATH and DepthQA.** The red-highlighted boxes indicate self-evaluated cases.