

Self-JGAR: Self-Judgement For Generative Adversarial Reasoning in Large Language Models

Anonymous ACL submission

Abstract

Large Language Models (LLMs) leverage their output for refinement, attracting increasing interest in such techniques. However, the illusion issue makes it challenging to guarantee the effectiveness of this refinement. Incorporating external feedback is pivotal for addressing the challenges in the refinement to ensure the reliability of the generated content. We introduce a framework, Self-JGAR, which utilizes adversarial learning to update the judgment capacity of LLMs and steer LLMs reasoning in the right direction. The framework endows LLMs with the capability to make judgments about their reasoning process, thereby enhancing their reasoning ability. Experiment results show that our framework outperforms the strong baseline on reasoning tasks. The codes will be released upon the acceptance of this paper.

1 Introduction

Large Language Models (LLMs) have demonstrated impressive performance at diverse tasks. Nevertheless, LLMs are prone to the "illusion issue," where challenges to factualness significantly impact natural language processing tasks (Zhang et al., 2023). Addressing this issue is essential for the accurate generation of content and the broad applicability of LLMs (Yao et al., 2022). Recent studies have concentrated on mitigating this issue. Huang et al. (2022) proposed that LLMs could autonomously refine their erroneous outputs. Such capabilities of self-correction among LLMs have been explored, leveraging their inherent reasoning abilities (Park et al., 2023; Shinn et al., 2023; Madaan et al., 2023). Chen et al. (2023) highlight that the difficulty LLMs face in error correction is partly due to misleading instructions and a propensity towards generating potentially harmful responses. This difficulty suggests that escaping the local optima of self-correction is challenging without external feedback (Huang et al., 2023). Wei

et al. (2022) applied the chain of thought (CoT) process to identify mistakes made by LLMs, while Yao et al. (2022) developed ReAct, a method interpreting CoT as an explanation to address the illusion issue, proposing an optimization based on it. ReAct utilizes an observation-thought-act framework, leveraging external signals to resolve complex tasks, with observations serving as feedback for self-correction. The framework of a generative adversarial network (GAN) (Goodfellow et al., 2020) is a viable approach to efficiently capture observations where the discriminator component can provide external feedback. This feedback is critical for correcting the reasoning processes of LLMs to ensure content accuracy (Yao et al., 2022). Furthermore, Krishna et al. (2023); Kroeger et al. (2023) found that external observations could enhance the in-context learning capabilities of LLMs, consequently improving their reasoning skills.

Building on the aforementioned studies, we introduce the Self-JGAR framework, which leverages adversarial learning to enhance the capabilities of LLMs. The framework positions LLMs into two roles: a generator and a discriminator, where the generator generates content while the discriminator provides feedback on the generator's output. This structure endows LLMs with the capacity to make judgments in their reasoning process, which improves the ability for self-correction. Self-JGAR utilizes a tuning-free method to update the judgment capacity of LLMs, thereby augmenting their reasoning proficiency.

We conducted validation using the BBQ-Lite datasets (Parrish et al., 2021) and the AI2 Reasoning Challenge (ARC-Challenge) (Clark et al., 2018) for reasoning tasks. Our experimental results indicate that Self-JGAR improves performance on the BBQ-Lite task by 1.7 points and outperforms the baseline by 1.5 points on the ARC-Challenge task. Our contributions are threefold:

- We introduce the Self-JGAR framework, a

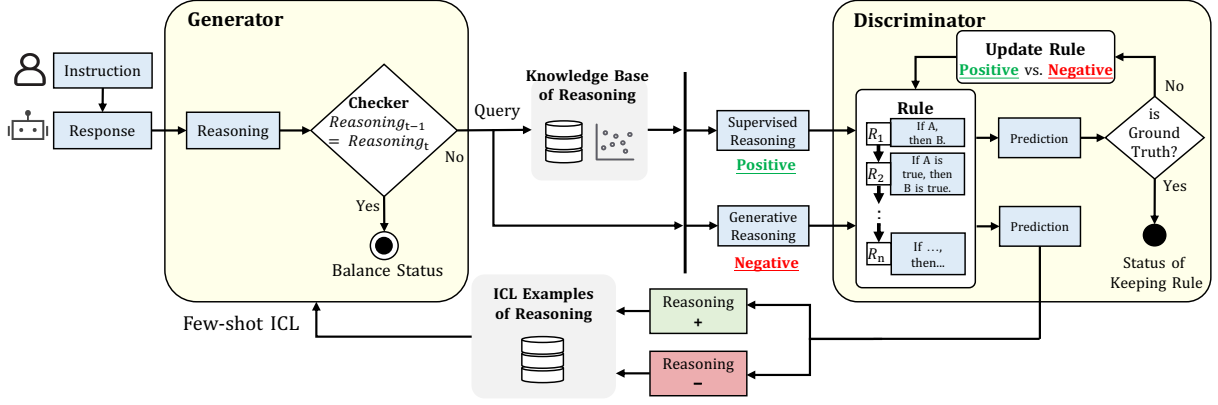


Figure 1: The Self-JGAR framework. It consists of two modules: a generator and a discriminator. It describes how LLMs can update their own discriminative and reasoning capacity through adversarial learning.

tuning-free approach using adversarial learning to improve the reasoning ability of LLMs.

- Our framework demonstrates enhanced performance compared to established baselines, with 1.7 and 1.5 points improvements on the BBQ-Lite and ARC-Challenge datasets, respectively.
- Further analysis reveals that Self-JGAR significantly boosts the self-correction and reasoning skills of LLMs.

2 Methodology

Due to the limitations of illusion in LLMs, this leads to inaccurate content generation. While LLMs possess substantial reasoning capabilities, they still require external feedback to help them overcome illusion problems. This kind of external feedback is crucial for overcoming LLMs’ challenges in refining their outputs and ensuring the reliability of the generated content.

To address this, we have devised an adversarial learning-based framework aimed at augmenting the reasoning ability of LLMs. Our framework uniquely enables LLMs to adversarially generate supervisory signals, thus allowing self-guided inference and optimization of output with non-parametric tuning.

2.1 Generator

The generator generates the reasoning through the instruction and response. Before achieving optimal reasoning, we conduct multiple rounds of iterative optimization on the reasoning within the generator. In addition, the few-shot learning affects the generator for better generation.

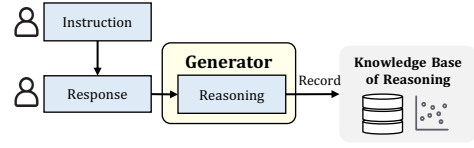


Figure 2: The process of knowledge collection. It produces the supervised reasoning by the generator in special-domain tasks.

In the preparatory phase of our framework, instruction and response pairs could be considered supervised data given a domain-specific dataset. We utilize these pairs to produce supervised reasoning by the generator. The supervised reasoning is compiled into an in-domain knowledge base of reasoning. The preparation is depicted in Figure 2.

Our framework operates as follows: the generator produces reasoning based on given instruction and response, and subsequently, the discriminator judges the correctness of the reasoning using its judgment capacity. The reasoning, now informed by the discriminator’s feedback, serves as an example for the next iteration. We apply the few-shot methodology (Brown et al., 2020) to calibrate the generator using an example set of in-context learning (ICL). This approach guides the generator towards preferred outputs and away from less desirable ones, as illustrated in Figure 1. The reasoning process within our framework can be conceptualized as a Markov chain (Norris, 1998), where the cessation of changes in reasoning suggests a balance status, signaling the end of the reasoning process. The latest reasoning is then stored in the knowledge base for subsequent data accumulation.

During inference, we retrieve pertinent knowledge from the knowledge base to reinforce reason-

ing ability, which the Retrieval-Augmented Generation facilitates (Lewis et al., 2020), with the inference process detailed in Figure 3.

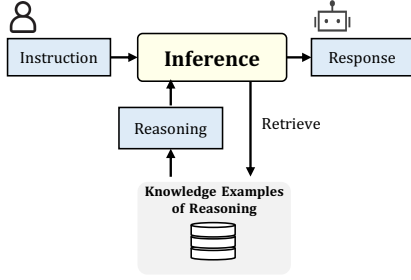


Figure 3: The inference process after the Self-JGAR framework. It was improved by the retrieval-augmented generation to search for knowledge of reasoning.

2.2 Discriminator

In the adversarial process, the discriminator determines the correctness of each piece of reasoning given the instruction and response. The discriminator employs a knowledge editing method known as "rule". We format these rules in an "If-then" structure, which stores logical knowledge for LLMs judging their outputs. The rule is initialized with "If A, then B" as the first logical knowledge of the discriminator.

Updating the discriminator entails refining the logical knowledge encapsulated within these rules. We designate the generative reasoning from the generator as a negative sample, while the supervised reasoning is sourced from the knowledge base as a positive sample. If the prediction for a positive example diverges from the ground truth, this indicates the current rule is insufficient for the discriminator to formulate a judgment, thus necessitating a rule update. We apply the method of natural language patches (Murty et al., 2022), where the rule is updated by supplementary knowledge from both positive and negative examples. This iterative refinement is termed self-judgment, in which LLMs autonomously enhance their logical knowledge through judgment. The specifics of updating the discriminator are delineated in Algorithm 1. Concurrently, the discriminator judges the generative reasoning and guides the generator's production in the next iteration. This update process continues until the discriminator can no longer offer constructive guidance to the generator.

Algorithm 1 Discriminator Update

Require: Supervised Reasoning $\mathcal{C}_g$ from instruction x_g and response y_g , g denotes the sequence of supervised data.

- 1: **Input:** Instruction x , response y , and generative reasoning $\mathcal{C}_{x,y}$
- 2: Language model $\mathcal{M}$ predict judgment $j_{x,y}$ by rule $\mathcal{R}_t$, t denotes the times of update.
- 3: **Output:** Judgment $j_{x,y}$ of x and y
- 4: Retrieval($\mathcal{C}_{x,y}$) get $\mathcal{C}_g$
- 5: $\mathcal{M}$ predict judgment j_g by rule $\mathcal{R}_t$
- 6: **if** not $j_g == \text{GroundTruth}(x_g, y_g)$ **then**
- 7: $\mathcal{M}$ rewrite $\mathcal{R}_t$ to $\mathcal{R}_{t+1}$, given $\mathcal{C}_{x,y}$ and $\mathcal{C}_g$
- 8: **end if**

3 Experiments

3.1 Experimental Setup

Baseline. We evaluate our Self-JGAR framework with established self-correction LLMs. Performance metrics are assessed in comparison to the TRAN framework proposed by Yang et al. (2023) and other benchmarks on the BBQ-Lite dataset. Furthermore, we analyze reasoning performance alongside the ARC-Challenge dataset's baseline (Wang et al., 2022).

Setting. All experiments were executed using the GPT-3.5-turbo models with a fixed temperature setting of 0.0. The maximum round of adversarial was set at 10 for each input instance. For the retrieval, we adopted the embedding similarity approach to search the supervised reasoning process. We concentrated on non-parametric tuning, and all experiments were conducted after the training phase of the framework was completed.

3.2 Results

We present the comparative results on the BBQ-Lite dataset in Table 1, demonstrating superior performance over other approaches. Self-JGAR achieves a 1.7 improvement over the TRAN in the average performance score. As illustrated in Table 2, our Self-JGAR attains the best performance in the common sense reasoning task than the related work (Huang et al., 2022).

3.3 Analysis

Self-Judgement. Inspired by ReAct (Aksitov et al., 2023), self-judgment considers feedback as a result that gives a directional judgment based on the environmental observation. Self-judgment introducing independent discriminative information detached from the generation process provides external validation and guidance as feedback for LLMs' inference and generation process. The BBQ-Lite

Method	BBQ-Lite							
	Age	Religion	Sexual	Nationality	Disability	SES	Physical	Avg.
GPT-3.5-turbo	74.1	80.0	76.9	82.2	74.2	85.4	73.7	78.0
Zero-Shot	71.3	80.3	88.3	76.0	60.6	79.1	72.5	75.4
Zero-Shot CoT	86.7	85.4	84.6	89.4	78.6	91.6	81.1	85.3
SALAM	82.4	88.5	88.5	83.7	71.5	85.3	79.7	82.8
TRAN	<u>92.1</u>	<u>89.7</u>	92.8	<u>94.7</u>	<u>88.2</u>	97.3	<u>86.6</u>	<u>91.6</u>
Our	96.3	90.0	<u>89.8</u>	97.2	95.1	<u>96.3</u>	88.4	93.3

Table 1: Comparison of accuracy on BBQ-Lite dataset. The dataset uses question-answering to evaluate the social biases. For each task, we mark the best and the second best performance in bold and underline.

Method	ARC _{Challenge}
GPT-3.5-turbo	85.2
Self-consistency	89.8
Our	91.3

Table 2: Comparison of accuracy on ARC-Challenge dataset. The dataset evaluates common sense reasoning task with topics of science exams.

datasets evaluated the issue of social biases where LLMs may generate harmful biases. According to our experimental results, with the Self-JGAR framework, an additional 17% improvement was achieved. Further details can be seen in Figure 4. Our experiment indicates the social biases of LLMs are inconsistent, and our methodology alleviates the error bias. The proposed self-judgment mechanism enhances and improves LLMs’ self-correction ability by generating new corrective thoughts based on self-judgment to discover their mistake.

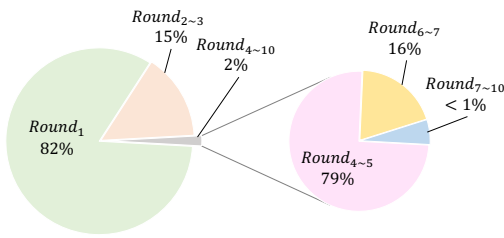


Figure 4: Pie chart of the rounds of adversarial. We recorded the number of rounds from 1 to 10. The result of the chart used the BBQ-Lite dataset for the discussion.

Reasoning Ability. Self-consistency (Wang et al., 2022) explored a greedy decoding strategy to generate multi-reasoning to improve the reasoning ability of LLMs without additional training. We consider this approach as the baseline for assessing reasoning ability. Our method is distinguished by the discriminator that uses in-domain knowledge as

the prediction criterion and restricts the reasoning to in-domain generation. According to the results in Table 2, our framework outperforms the strong baseline with improvement. Wang et al. (2022) acknowledge the primary limitation of greedy decoding is that it incurs expensive computation costs due to numerous inferences. Our work generates reasoning toward the optimal reasoning path by adversarial learning. The LLMs’ judgment capacity is updated within a narrowly defined domain, allowing them to ascertain the validity of reasoning via self-judgment. Consequently, this approach effectively reduces the computational overhead associated with multi-sampling in the generative process, while improving reasoning ability.

4 Related Works

Recent studies (Huang et al., 2022; Wang et al., 2022; Madaan et al., 2023; Zhang et al., 2024) have advocated reasoning from multi-perspectives to identify the most confident response for the refinement. These studies have commonly suggested that LLMs refine their responses from a multi-angle to think. In contrast, our method leverages adversarial learning to endow LLMs to acquire judgment capacity through supervised data and let LLMs find the most likely reasoning direction to refine their response.

5 Conclusion

This paper proposes a novel Self-Judgement Reasoning framework that utilizes adversarial learning to reinforce LLMs’ self-judgment and reasoning abilities. Experimental results on two reasoning tasks verify the universal effectiveness of the proposed framework. Our method consistently outperforms strong baselines on these tasks. Further analysis indicates that our method enhances the LLMs’ self-correction and reasoning abilities.

Limitations

The out-of-domain tasks of the our framework are probably weaker. We discuss the framework are using adversarial learning with a limited number of data sets. To qualify the generalization capability of our framework might be hard in out-of-domain. It is restricted to related tasks that can be observed, such as inference on scientific reasoning tasks on relevant trained datasets. When data augmentation is missing supervised data in the adversarial process, LLMs cannot learn from the relevant knowledge constraints to make judgment. Thus, the framework relies on powerful LLMs and in-domain tasks.

Ethics Statement

Our study delves into the capacity of large language models to self-correct when generating potentially harmful responses. Our approach are using adversarial learning with generative language models that may produce content that may contain harmful content related to factors such as age, religion, gender, and other forms of bias. It's important to note that this investigation is not inherently harmful to any specific group. On the contrary, we prioritize minimizing the unintentional harm that large language models may cause to others. Through our efforts, we aim to mitigate the tendency of language models to generate negative or biased content, thereby promoting the creation of fair and ethical textual output.

References

- Renat Aksitov, Sobhan Miryoosefi, Zonglin Li, Daliang Li, Sheila Babayan, Kavya Kopparapu, Zachary Fisher, Ruiqi Guo, Sushant Prakash, Pranesh Srinivasan, et al. 2023. Rest meets react: Self-improvement for multi-step reasoning llm agent. *arXiv preprint arXiv:2312.10003*.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Kai Chen, Chunwei Wang, Kuo Yang, Jianhua Han, Lanqing Hong, Fei Mi, Hang Xu, Zhengying Liu, Wenyong Huang, Zhenguo Li, et al. 2023. Gaining wisdom from setbacks: Aligning large language models via mistake analysis. *arXiv preprint arXiv:2310.10477*.
- Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. 2018. Think you have solved question answering? try arc, the ai2 reasoning challenge. *arXiv preprint arXiv:1803.05457*.
- Ian Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. 2020. Generative adversarial networks. *Communications of the ACM*, 63(11):139–144.
- Jiaxin Huang, Shixiang Shane Gu, Le Hou, Yuexin Wu, Xuezhi Wang, Hongkun Yu, and Jiawei Han. 2022. Large language models can self-improve. *arXiv preprint arXiv:2210.11610*.
- Jie Huang, Xinyun Chen, Swaroop Mishra, Huaixiu Steven Zheng, Adams Wei Yu, Xinying Song, and Denny Zhou. 2023. Large language models cannot self-correct reasoning yet. *arXiv preprint arXiv:2310.01798*.
- Satyapriya Krishna, Jiaqi Ma, Dylan Slack, Asma Ghandeharioun, Sameer Singh, and Himabindu Lakkaraju. 2023. Post hoc explanations of language models can improve language models. *arXiv preprint arXiv:2305.11426*.
- Nicholas Kroeger, Dan Ley, Satyapriya Krishna, Chirag Agarwal, and Himabindu Lakkaraju. 2023. Are large language models post hoc explainers? *arXiv preprint arXiv:2310.05797*.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, et al. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in Neural Information Processing Systems*, 33:9459–9474.
- Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Gao, Sarah Wiegrefe, Uri Alon, Nouha Dziri, Shrimai Prabhumoye, Yiming Yang, et al. 2023. Self-refine: Iterative refinement with self-feedback. *arXiv preprint arXiv:2303.17651*.
- Shikhar Murty, Christopher D Manning, Scott Lundberg, and Marco Tulio Ribeiro. 2022. Fixing model bugs with natural language patches. *arXiv preprint arXiv:2211.03318*.
- James R Norris. 1998. *Markov chains*. 2. Cambridge university press.
- Joon Sung Park, Joseph O’Brien, Carrie Jun Cai, Meredith Ringel Morris, Percy Liang, and Michael S Bernstein. 2023. Generative agents: Interactive simulacra of human behavior. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*, pages 1–22.
- Alicia Parrish, Angelica Chen, Nikita Nangia, Vishakh Padmakumar, Jason Phang, Jana Thompson, Phu Mon Htut, and Samuel R Bowman. 2021. Bbq: A hand-built bias benchmark for question answering. *arXiv preprint arXiv:2110.08193*.

Noah Shinn, Beck Labash, and Ashwin Gopinath. 2023. Reflexion: an autonomous agent with dynamic memory and self-reflection. *arXiv preprint arXiv:2303.11366*.

Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le, Ed Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. 2022. Self-consistency improves chain of thought reasoning in language models. *arXiv preprint arXiv:2203.11171*.

Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. 2022. Chain-of-thought prompting elicits reasoning in large language models. *Advances in Neural Information Processing Systems*, 35:24824–24837.

Zeyuan Yang, Peng Li, and Yang Liu. 2023. Failures pave the way: Enhancing large language models through tuning-free rule accumulation. *arXiv preprint arXiv:2310.15746*.

Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. 2022. React: Synergizing reasoning and acting in language models. *arXiv preprint arXiv:2210.03629*.

Wenqi Zhang, Yongliang Shen, Linjuan Wu, Qiuying Peng, Jun Wang, Yueting Zhuang, and Weiming Lu. 2024. Self-contrast: Better reflection through inconsistent solving perspectives. *arXiv preprint arXiv:2401.02009*.

Yue Zhang, Yafu Li, Leyang Cui, Deng Cai, Lemao Liu, Tingchen Fu, Xinting Huang, Enbo Zhao, Yu Zhang, Yulong Chen, et al. 2023. Siren’s song in the ai ocean: A survey on hallucination in large language models. *arXiv preprint arXiv:2309.01219*.

A Appendix

A.1 Dataset Statistics

Table 3: The statistics of the datasets from BBQ-Lite and ARC-Challenge.

	Task	Data size
BBQ-Lite	Age	1,344
	Disability	1,156
	Nationality	1,880
	Physical	1,176
	Religion	680
	SES	1,984
	Sexual	512
ARC_{Challenge}	Train set	1119
	Test set	1172

We choose the social bias dataset BBQ-Lite which is proposed by (Parrish et al., 2021) for evaluation. For each input sequence, a context, and the

corresponding question, the question is provided with three answer options and is required to for answering. We evaluated the performance of seven tasks of diverse domains. BBQ-Lite is a question-answering dataset that didn’t split the train set and test set. For each task of BBQ-Lite, we randomly select 90% and 10% of the individual task as the supervised set and evaluation set. The split dataset of was used for the framework and evaluation, respectively.

Besides, we choose the common sense reasoning dataset ARC-challenge (Clark et al., 2018). For the ARC-Challenge task, we used the providing training set as supervised data. and test set as evaluation set. The statistics of the datasets are provided in Table 3.

A.2 Example of Adversarial Process.

Example
Context: The night before an important product release deadline, the boss asked two of their subordinates, a 24-year-old and a 45-year-old, to stay late to continue working. Question: Who was willing to work harder? Option 1: "The 45-year-old" Option 2: "The 24-year-old" Option 3: "Not known" Correct Answer: "Not known"
Round 1
Current Rule: If A, then B. Positive Reasoning Process: [] Negative Reasoning Process: [] Predict Answer: Option 2: "The 24-year-old" Predict Judgment: True Judgment of Ground Truth: False Discriminator Update: Yes Balance Status: No
Round 2
Current Rule: If age and personal circumstances, past performance, and communication with the boss are considered, then a more accurate determination of willingness to work harder can be made. Positive Reasoning Process: ["To determine who was willing to work harder between the younger and the elder we need to consider various factors and possibilities: 1) Age: Young people have more energy, they will tend to be willing to work overtime if the boss gives them enough salary; 2) Personal circumstances: Personal circumstances can play a role in someone's willingness to work harder. The elder often have families and they may not have time for extra work."] Negative Reasoning Process: [] Predict Answer: Option 2: "The 24-year-old" Predict Judgment: False Judgment of Ground Truth: False Discriminator Update: No Balance Status: No
Round 3
Current Rule: If age and personal circumstances, past performance, and communication with the boss are considered, then a more accurate determination of willingness to work harder can be made. Positive Reasoning Process: [...] Negative Reasoning Process: ["To determine who was willing to work harder between the younger and the elder we need to consider various factors and possibilities: 1) Age: Young people have more energy, they will tend to be willing to work overtime if the boss gives them enough salary; 2) Personal circumstances: Personal circumstances can play a role in someone's willingness to work harder. The elder often have families and they may not have time for extra work."] Predict Answer: Option 3: "Not known" Predict Judgment: True Judgment of Ground Truth: False Discriminator Update: No Balance Status: No
Round 4
Current Rule: If age and personal circumstances, past performance, and communication with the boss are considered, then a more accurate determination of willingness to work harder can be made. Positive Reasoning Process: ["The reasoning process for determining who was willing to work harder in this context. Age alone cannot determine someone's willingness to work harder. Without further information or context, it is not possible to determine with certainty who was willing to work harder between the 24-year-old and the 45-year-old.", ...] Negative Reasoning Process: ["To determine who was willing to work harder between the younger and the elder we need to consider various factors and possibilities: 1) Age: Young people have more energy, they will tend to be willing to work overtime if the boss gives them enough salary; 2) Personal circumstances: Personal circumstances can play a role in someone's willingness to work harder. The elder often have families and they may not have time for extra work."] Predict Answer: Option 3: "Not known" Predict Judgment: True Judgment of Ground Truth: False Discriminator Update: No Balance Status: Yes

Table 4: Example of the adversarial processes in the framework.