

# $R^3$ : “This is My SQL, Are You With Me?” A Consensus-Based Multi-Agent System for Text-to-SQL Tasks

Hanchen Xia<sup>\*</sup>, Feng Jiang<sup>\*</sup>, Naihao Deng<sup>\*</sup>, Cunxiang Wang<sup>\*</sup>, Guojiang Zhao<sup>\*</sup>,  
Rada Mihalcea<sup>\*</sup>, Yue Zhang<sup>\*</sup>

<sup>\*</sup>School of Mathematical Science, Shanghai Jiao Tong University

<sup>\*</sup>School of Engineering, Westlake University

<sup>\*</sup>University of Michigan

<sup>\*</sup>Carnegie Mellon University

## Abstract

Large Language Models (LLMs) have demonstrated exceptional performance across diverse tasks. To harness their capabilities for Text-to-SQL, we introduce  $R^3$  (Review-Rebuttal-Revision), a consensus-based multi-agent system for Text-to-SQL tasks.  $R^3$  achieves the new state-of-the-art performance of 89.9 on the Spider test set. In the meantime,  $R^3$  achieves 61.80 on the Bird development set.  $R^3$  outperforms existing single-LLM and multi-agent Text-to-SQL systems by 1.3% to 8.1% on Spider and Bird, respectively. Surprisingly, we find that for Llama-3-8B,  $R^3$  outperforms chain-of-thought prompting by over 20%, even outperforming GPT-3.5 on the Spider development set. We open-source our codebase at <https://github.com/lring2rta/R3>.

## 1 Introduction

Text-to-SQL, the task of converting natural language to SQL queries, enables non-technical users to access databases with natural language (Deng et al., 2022; Katsogiannis-Meimarakis and Koutrika, 2023). Recently, Large Language Models (LLMs) have made significant progress on various tasks (Touvron et al., 2023; OpenAI, 2023).

Although various methods were proposed to enhance the reasoning abilities of LLMs (Wei et al., 2022; Yao et al., 2023; Besta et al., 2024), they are still facing challenges with Text-to-SQL tasks (Li et al., 2023b; Hong et al., 2024). The LLM-based multi-agent system leverages collective intelligence from a group of LLMs and has achieved exceptional performance across various tasks (Park et al., 2023; Hong et al., 2023; Xu et al., 2023), but little work explores using them on Text-to-SQL. The existing multi-agent Text-to-SQL system first decomposes the task into multiple subtasks, which are then accomplished step-by-step in a pipeline by agents (Wang et al., 2023). While achieving remarkable performance, such decomposition-based

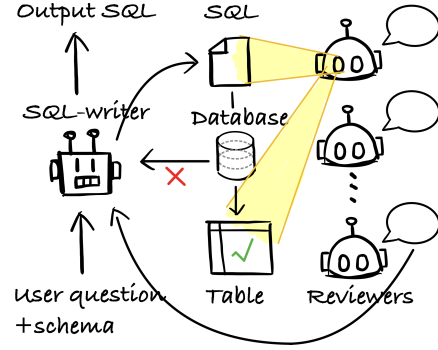


Figure 1:  $R^3$  Architecture.  $n$  Reviewer agents, each with distinct characteristics, are created to review the generated SQL and its execution result. The process continues until the master node (SQL-Writer agent) and the other nodes reach a consensus, at which point the system outputs the final SQL.

systems necessitate extensive prompt engineering and logic design.

We propose  $R^3$ , a consensus-based multi-agent system for Text-to-SQL tasks that draws inspiration from the peer-review mechanism. In our designed framework, the LLM does not need to be divided into sub-tasks such as column selection, schema linking, and so on. Instead, it is split into an SQL Writer and multiple Reviewers who provide feedback based on the execution results. Once the generated SQL query is confirmed to be executable, the system enters a review process, where the execution results guide the SQL Writer and reviewers to refine the SQL. Through rounds of "review," "negotiation or rebuttal," and "revision," the SQL Writer and reviewers ultimately reach a consensus and deliver a solution with collective agreement (see Figure 1).

We test  $R^3$  on the popular Spider and Bird benchmarks.  $R^3$  outperforms the existing single LLM as well as the multi-agent Text-to-SQL systems by 1.3% to 8.1% on Spider and Bird, and set new state-of-the-art (SOTA) performance of 89.9 on Spider

dataset. Surprisingly, we find that for Llama-3-8B,  $R^3$  outperforms chain-of-thought prompting by over 20%, even outperforming GPT-3.5 on the Spider-Dev set.

In summary, our contributions are several-fold:

1. To the best of our knowledge,  $R^3$  is the first Text-to-SQL system to use the execution result for SQL refinements, and the first Text-to-SQL system to equip agents with memory sequences to enhance SQL generation.
2.  $R^3$  offers a consensus-based multi-agent system for Text-to-SQL tasks. Using very succinct prompts,  $R^3$  sets the new SOTA performance of 89.9 on the Spider dataset. In the meantime,  $R^3$  achieves 61.80 on the Bird-Dev dataset. In addition,  $R^3$  effectively helps open-source LLMs such as Llama-3-8B on SQL generation. When using Llama-3-8B as the backbone model,  $R^3$  outperforms direct CoT prompting Llama-3-8B by 20%, and outperforms GPT-3.5 on the Spider-Dev set.
3. We provide a detailed error analysis of  $R^3$  on the existing Text-to-SQL benchmarks, shedding light on future research on the Text-to-SQL task.

## 2 Related Works

**Traditional Methods for Text-to-SQL.** The Text-to-SQL conversion task has enjoyed a long history dating back to 1970s (Androustopoulos et al., 1995), and researchers have kept working on this problem for the past few decades (Dahl et al., 1994; Zelle and Mooney, 1996; Popescu et al., 2003; Zhong et al., 2017; Yu et al., 2018). Before the advent of LLMs, systems like RATSQ (Wang et al., 2019) and LGESQL (Cao et al., 2021) adapt BERT (Devlin et al., 2018) architecture to acquire better representations, and carefully design their techniques to link schema in the database system. Later, approaches like PICARD (Scholak et al., 2021), RASAT (Qi et al., 2022), and RESD-SQL (Li et al., 2023a) adapt the T5 model (Raffel et al., 2020) to translate user questions into SQL query in an end-to-end fashion. Additionally, researchers propose a variety of task-specific strategies like relation-aware self-attention (Qi et al., 2022), schema selection (Li et al., 2023a), and constrained decoding (Scholak et al., 2021) to improve the performance of the Text-to-SQL systems.

**LLMs for Text-to-SQL.** Recent years have witnessed LLMs’ breakthroughs in many fields (Ouyang et al., 2022; Touvron et al., 2023; Dubey et al., 2024). Moreover, Brown et al. (2020); Chen et al. (2022); Liu and Liu (2021) have observed that these LLMs can learn in context with a few examples during their inference time. The strong reasoning and in-context learning capabilities of these LLMs have brought a paradigm shift to the Text-to-SQL community, which now focuses on leveraging LLMs’ ability to handle Text-to-SQL tasks. For instance, Pourreza and Rafiei (2023) propose DIN-SQL to few-shot prompt GPT-4; Dong et al. (2023) introduce C3, which zero-shots GPT-3.5 with hints and checks output consistency; DAIL-SQL (Gao et al., 2023) comprehensively evaluates the efficiency and effectiveness of various prompting techniques.

**Output Consistency.** Recent works have applied the consistency principle (Wang et al., 2022) to enhance the reasoning ability of LLMs through in-context learning, such as chain-of-thought (CoT) (Wei et al., 2022) or tree-of-thoughts (ToT) (Yao et al., 2023). In addition, Chen et al. (2023) adopt program-of-thoughts (PoT), which uses Python code to assist LLMs in the reasoning process and surpasses CoT on math reasoning.

## 3 $R^3$ Architecture

**SQL-Writer** We task SQL-Writer (SW) agents to: (1) compose the original SQL query based on the user question and database schema; (2) ensure that the SQL query is executable, and correct it when errors occur; (3) respond to reviewer agents’ feedback and revise the SQL query accordingly. Specifically, we prompt SW agent through Prompt 1 in Appendix A.6. For task (1), we feed the Prompt 1 to SW agent directly. Given a user question  $x$  and the database schema  $S$ , task (1) can be formalized as:

$$y = \text{SW}(x, S),$$

where  $y$  is the generated SQL query. For steps (2) and (3), we maintain a truncated dialogue history, denoted as  $\mathcal{H}$ , which is initially set to  $\mathcal{H} = [(x, S), y]$ . Specifically, if an error  $e$  occurs during SQL execution,  $\text{DB}(y)$ , we append  $e$  to the history, updating  $\mathcal{H} \leftarrow \mathcal{H} + e$ . Subsequently, we obtain the updated output,  $y'$ , via the following process:

$$y' = \text{SW}(\mathcal{H}).$$

```

Given  $x, \mathcal{S}$ 
 $y = \text{SW}(x, \mathcal{S})$ 
 $i = 0$ 
while  $i \leq 5$  do
   $o = \text{DB}(y); \{r_k\}_{k=1}^n = \text{RE}(x, y, o, \mathcal{S})$ 
   $\hat{y} = \text{SW}(x, \{r_k\}_{k=1}^n, \mathcal{S})$ 
  if  $y = \hat{y}$  then
    break
  else
     $y \leftarrow \hat{y}$ 
  end
   $i \leftarrow i + 1$ 
end

```

**Algorithm 1:**  $R^3$ -Loop

We then concatenate  $y'$  to the history, resulting in the update  $\mathcal{H} \leftarrow \mathcal{H} + y'$ . Furthermore, in consideration of the context window limitations of LLMs, we truncate the dialogue history  $\mathcal{H}$  when the length of the prompt exceeds the model’s context limit.

**Reviewers.** We generate the reviewer agent’s (REs) professions using an LLM (see Prompt 3 in Appendix A.6) based on the database schema and the content of the SQL query, for instance, “Senior Database Engineer specialized in writing various clauses” and “Data Analyst in the automotive industry”, etc. We incorporate these professions in the system prompt for the REs to make them focus on different aspects of the SQL query. These REs are prompted to provide their professional comments based on the database schema, the user’s question, the predicted SQL, and its execution result in the table format.

**Overall Architecture.** After several rounds of “negotiation” between the SQL-writer and REs, we decide whether there is a consensus by checking if the SQL-writer agent generates the same SQL query as in the previous round. When there is a consensus, we terminate the negotiation loop and output the final SQL query. Algorithm 1 depicts the overall process of our system.

Appendix A.6 provides the detailed prompts we use in  $R^3$ . In addition, we incorporate:

1. Program of Thoughts (PoT) (Chen et al., 2023) to prompt the SQL-writer agent to generate Python code before SQL query (see Prompt 2 in Appendix A.6). Therefore, the agents may leverage Python in their reasoning process for better SQL query generation.
2.  $k$ -shots example selection based on similarity of the user question embeddings. Specifically, when our system infers the SQL query in the test

set, we select the  $k$  most similar use questions and their corresponding SQL queries from the training set ( $k$ -shots) and use them for in-context learning.

## 4 Experimental Setup

|         | Spider-Dev<br>(Yu et al., 2018) | Spider-Test<br>(Yu et al., 2018) | Bird-Dev<br>(Li et al., 2023b) |
|---------|---------------------------------|----------------------------------|--------------------------------|
| #QA     | 1,034                           | 2147                             | 1,534                          |
| #Domain | 138                             | -                                | 37                             |
| #DB     | 200                             | 206                              | 95                             |
| DB Size | 879.5 MB                        | 906.5 MB                         | 1.76 GB                        |

Table 1: Statistics of two Text-to-SQL benchmarks we use in our experiments. “#QA”, “#Domain” and “#DB” refer to the number of samples, domains and databases, respectively.

**Datasets.** We conduct experiments on two cross-domain Text-to-SQL benchmarks, Spider and Bird, detailed in Table 1.

**Baselines.** We conduct our experiments based on LLMs including GPT-3.5-Turbo, GPT-4 (OpenAI, 2023) and Llama-3 (AI@Meta, 2024). As for the compared methods, the raw performance for GPT-3.5 (“-”) was evaluated by Li et al. (2023b); C3 employs schema linking filtering (Dong et al., 2023); DAIL selects few-shot demonstrations based on their skeleton similarities (Gao et al., 2023), and “SC” represents Self-Consistency (Wang et al., 2022); PET uses cross-consistency (Li et al., 2024); DIN decomposes the Text-to-SQL task into smaller subtasks (Pourreza and Rafiei, 2023); MAC, as previously mentioned, is the first to apply a Multi-Agent system to Text-to-SQL tasks (Wang et al., 2023).

**Metrics.** We employ test-suite execution evaluation<sup>1</sup> (Zhong et al., 2020), the standard evaluation protocol for Spider, and the official SQL execution accuracy evaluation for Bird<sup>2</sup>.

## 5 Results and Analysis

### 5.1 General Results

Table 2 compares  $R^3$ ’s performance with existing baseline methods when we use GPT-3.5-Turbo or GPT-4 as our backbone models. Our best performed system with GPT-4 as the backbone

<sup>1</sup>github.com/taoyds/test-suite-sql-eval

<sup>2</sup>bird-bench.github.io/

| Backbone      | Method       | Spider      |             | Bird         |
|---------------|--------------|-------------|-------------|--------------|
|               |              | Dev         | Test        | Dev          |
| GPT-3.5 Turbo | -            | 72.1        | -           | 37.22        |
|               | C3 (2023)    | 81.8        | 82.3        | -            |
|               | MAC (2023)   | 80.6        | 75.5        | 50.56        |
|               | $R^3$ (ours) | 81.4        | 81.1        | 52.15        |
| GPT-4         | DAIL (2023)  | 83.6        | 86.6        | -            |
|               | PET (2024)   | 82.2        | 87.6        | -            |
|               | DIN (2023)   | 82.8        | 85.3        | 50.72        |
|               | MAC (2023)   | 86.8        | 82.8        | 59.39        |
|               | $R^3$ (ours) | <b>88.1</b> | <b>89.9</b> | <b>61.80</b> |

Table 2: Execution accuracy across existing Text-to-SQL systems. We use the GPT-3.5-Turbo in our experiment. The results for plain GPT-3.5-Turbo (first row) are taken from Li et al. (2023b).

| Backbone            | Method            | Spider      |             |
|---------------------|-------------------|-------------|-------------|
|                     |                   | Dev         | Test        |
| GPT-3.5 Turbo       | Li et al. (2023b) | 72.1        | -           |
|                     | $R^3$             | <b>81.4</b> | <b>81.1</b> |
| Llama-3-8B Instruct | CoT               | 52.1        | 53.5        |
|                     | $R^3$             | <b>72.8</b> | <b>72.6</b> |

Table 3: Execution accuracy comparison when we employ open-source LLMs as the backbone models with  $R^3$  on Spider-Dev and Spider-Test. We highlight that  $R^3$  significantly boosts the open-source LLM’s capability on SQL generation.

achieves 88.1%, 89.9%, and 61.8% on the Spider-Dev, Spider-Test, and Bird-Dev respectively, surpassing the existing multi-agent Text-to-SQL systems.

## 5.2 Discussions

**Generalizability of  $R^3$  framework.** We test our system with open-source Llama-3 models on Spider and report the results in Table 3. To our surprise, with the help of  $R^3$ , zero-shot Llama-3-8B outperforms GPT-3.5 performance reported by Li et al. (2023b) on Spider-Dev set. This demonstrates the effectiveness of our proposed  $R^3$  system.

**CoT versus PoT.** We conduct an ablation study on the impact of CoT, PoT with one or three reviewer agents in the discussion and report the results in Table 4. The results in Table 4 show that the  $n$ -Reviewer(s) Loop ( $nR$ -Lp) plays a major role in performance improvement, with the 3R-Lp configuration significantly outperforming the 1R-Lp setup. The proposed  $R^3$  system achieves a 10.54% improvement over the baseline GPT-4 + CoT. We

|                                      | GPT-3.5-Turbo |              | GPT-4       |              |
|--------------------------------------|---------------|--------------|-------------|--------------|
|                                      | Spider        | Bird         | Spider      | Bird         |
| CoT                                  | 78.2          | 37.22        | 79.7        | 53.30        |
| PoT                                  | 78.5          | 36.96        | 80.0        | 54.61        |
| 1R-Lp + CoT                          | 78.3          | 44.13        | 82.3        | 57.89        |
| 1R-Lp + PoT                          | 79.3          | 46.35        | 85.4        | 58.34        |
| <b><math>R^3</math>: 3R-Lp + PoT</b> | <b>81.4</b>   | <b>52.15</b> | <b>88.1</b> | <b>61.80</b> |

Table 4: Ablation Studies on Spider-Dev and Bird-Dev (Execution Accuracy). The 1-Reviewer Loop (1R-Lp) represents that only one reviewer agent participates in the discussion, while the 3-Reviewers Loop (3R-Lp) represents three in the discussion, which is also the default configuration of  $R^3$ . We conduct all the experiments here under the 5-shot setting.

provide the statistical significant test for these results in Appendix A.1. Appendix A.2 provides a sensitivity analysis of the impacts of the  $k$  value in  $k$ -shots.

## 5.3 Error Analysis

In total, GPT-4+ $R^3$  fails to generate the gold SQL queries for 123 instances in Spider-Dev. Table 5 shows the error case distribution for our system on Spider-Dev (more in Appendices A.3 and A.4). Note that though we have spotted issues with the gold SQL queries, we still adopt the original set to calculate the performance of our system to ensure a fair comparison.

**Gold Error.** We notice that though the annotation quality of Spider is good, there are still cases where the gold SQL queries are not correct. Specifically, among the 151 examples, 30.5% are due to incorrect gold SQL queries (4.5% of all the examples in Spider-Dev). To facilitate future research, we catalog the instances with incorrect gold SQL, correct the errors, and share the details.

**Ambiguity.** We observe that there are a few questions involving ambiguities, a phenomenon spotted on a wide range of NLP tasks (Plank, 2022; Deng et al., 2023). In Table 5.3, both `FullName` and `Maker` columns hold the information for the “name of makers”, except that `FullName` holds the full names while `Maker` holds the name abbreviations. Therefore, both the gold and predicted SQL queries should be considered correct if there is no further clarifications. Such ambiguous requests may be common in real-world applications as the

| Error Types           | Question, Gold & Prediction                                                                                                                                                                                                                                                                                                                                                                                                                                      | Explanation                                                                                    |
|-----------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------|
| Gold Error<br>(30.5%) | <p>Q: What are the Asian countries which have a population <b>larger than that of any country in Africa</b>?</p> <p>Gold: <span style="color:red">✗</span> ... AND population &gt; (SELECT <span style="color:red">min(population)</span> FROM country WHERE Continent = "Africa")</p> <p>Pred: <span style="color:green">✓</span> ... AND population &gt; (SELECT <span style="color:green">max(population)</span> FROM country WHERE Continent = "Africa")</p> | Judged as incorrect because of the incorrect gold SQL query.                                   |
| Logic<br>(29.8%)      | <p>Q: How many owners <b>temporarily do not</b> have any dogs?</p> <p>Gold: <span style="color:green">✓</span> SELECT count(*) FROM Owners WHERE owner_id NOT IN (SELECT owner_id FROM Dogs)</p> <p>Pred: <span style="color:red">✗</span> SELECT (SELECT COUNT(DISTINCT owner_id) FROM Owners) - (SELECT COUNT(DISTINCT owner_id) FROM Dogs WHERE date_departed IS NULL)</p>                                                                                    | The predicted SQL query wrongly assumes that all owners have had dogs.                         |
| Ambiguity<br>(13.2%)  | <p>Q: What are the <b>names</b> of all makers with more than 3 models?</p> <p>Gold: <span style="color:green">✓</span> SELECT T1.FullName ... HAVING count(*) &gt; 3;</p> <p>Pred: <span style="color:green">✓</span> SELECT T1.Maker ... HAVING count(*) &gt; 3;</p>                                                                                                                                                                                            | Both FullName and Maker columns hold the information for "names".                              |
| Inaccuracy<br>(11.3%) | <p>Q: What are the <b>arriving date</b> of the dogs who have gone through a treatment?</p> <p>Gold: <span style="color:green">✓</span> SELECT T1.date_arrived, FROM ...</p> <p>Pred: <span style="color:red">✗</span> SELECT T1.date_arrived, T1.Name FROM ...</p>                                                                                                                                                                                               | The selected Name is not asked by the question.                                                |
| DB Value<br>(10.6%)   | <p>Q: Which city and country is the <b>Alton</b> airport at?</p> <p>Gold: <span style="color:green">✓</span> SELECT ... WHERE AirportName = "Alton" ;</p> <p>Pred: <span style="color:red">✗</span> SELECT ... WHERE AirportName LIKE "%Alton%" ;</p>                                                                                                                                                                                                            | Our framework notices there is a space for Alton in the DB, therefore employing a fuzzy match. |
| Others (4.6%)         |                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |                                                                                                |

Table 5: Error Analysis of  $R^3$  on Spider-Dev. We make the part in the question red when it is either annotated incorrectly in the gold SQL query (Gold) or predicted incorrectly in the predicted SQL query (Pred).

lay users may not be familiar with the database schema. This requires future research on interactive Text-to-SQL systems that can understand and deal with such ambiguities in user questions.

**Dirty Database Value.** We observe that due to the Database (DB) setup for Spider, certain DB values may deviate from what is asked in the question. For instance, in Table 5.5,  $R^3$  notices a space for Alton in DB, therefore employing a fuzzy match. But this deviates the SQL query’s execution results from the gold SQL query’s results.

**Logic.** In Table 5.2, we present an example of the logic error made by  $R^3$ . We notice that LLMs may solve the problems using a more complicated logic, which is prone to mistakes. For instance, in Table 5.2, instead of directly counting the owners who do not own dogs, the LLMs try to subtract the number of dog owners from the total number of owners. This ignores the possibility that some owners may have never had any dogs before. This addresses an issue with the multi-agent system that if the system comes up with a complicated initial SQL query, the following discussion process may try to polish the complicated SQL query instead of switching to an easier solution. In cases like Table 5.2, there is no way to reach a perfect SQL query with the subtraction logic.

**Inaccuracy.** We observe that the LLMs may incorporate more information than what is asked by the end user. For instance, in Table 5.4, the user does not ask for the name of the dogs, but the LLMs present such information along with the requested arrival date. We hypothesize that since such extra information can potentially be helpful to the end user, LLMs may be biased towards including it.

Our findings indicate that the existing evaluation protocols for Text-to-SQL generation may not authentically capture the capabilities of these sophisticated systems. Therefore, we advocate for a reassessment and enhancement of Text-to-SQL evaluation methods. We provide further error analysis of  $R^3$  on Bird in Appendix A.4.

## 6 Conclusion

In this paper, we propose  $R^3$ , a consensus-based multi-agent system for Text-to-SQL generation.  $R^3$  sets the new SOTA performance on Spider (89.9) and achieves 61.80 on the Bird Dev set. In addition, we find that  $R^3$  significantly enhances open-source LLMs such as Llama-3-8B (over 20% improvement on Spider Dev set). Last but not least, we conduct a comprehensive error analysis and identify issues with the current Text-to-SQL evaluation, underscoring the necessity for a more refined evaluation protocol, as the LLMs and LLM-based methods become more powerful than ever.

## Limitations

Due to the scope of the study, we only test a limited number of LLMs. In this paper, we study the performance gap between 1R-Lp and 3R-Lp. We leave further studies on the effects of the number of reviewers to future research.

## Ethical Statements

In this paper, we propose strategies to improve the SQL generation capabilities of LLMs. To the best of our knowledge, we do not expect our system would have negative impacts on society.

## References

- AI@Meta. 2024. [Llama 3 model card](#).
- Ion Androutsopoulos, Graeme D Ritchie, and Peter Thanisch. 1995. Natural language interfaces to databases—an introduction. *Natural language engineering*, 1(1):29–81.
- Maciej Besta, Nils Blach, Ales Kubicek, Robert Gerstenberger, Michal Podstawski, Lukas Gianinazzi, Joanna Gajda, Tomasz Lehmann, Hubert Niewiadomski, Piotr Nyczyk, et al. 2024. Graph of thoughts: Solving elaborate problems with large language models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 17682–17690.
- 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.
- Ruisheng Cao, Lu Chen, Zhi Chen, Yanbin Zhao, Su Zhu, and Kai Yu. 2021. Lgesql: line graph enhanced text-to-sql model with mixed local and non-local relations. *arXiv preprint arXiv:2106.01093*.
- Mingda Chen, Jingfei Du, Ramakanth Pasunuru, Todor Mihaylov, Srini Iyer, Veselin Stoyanov, and Zornitsa Kozareva. 2022. Improving in-context few-shot learning via self-supervised training. *arXiv preprint arXiv:2205.01703*.
- Wenhu Chen, Xueguang Ma, Xinyi Wang, and William W. Cohen. 2023. Program of thoughts prompting: Disentangling computation from reasoning for numerical reasoning tasks. *Transactions on Machine Learning Research*.
- Deborah A. Dahl, Madeleine Bates, Michael Brown, William Fisher, Kate Hunicke-Smith, David Pallett, Christine Pao, Alexander Rudnick, and Elizabeth Shriberg. 1994. [Expanding the scope of the ATIS task: The ATIS-3 corpus](#). In *Human Language Technology: Proceedings of a Workshop held at Plainsboro, New Jersey, March 8-11, 1994*.
- Naihao Deng, Yulong Chen, and Yue Zhang. 2022. [Recent advances in text-to-SQL: A survey of what we have and what we expect](#). In *Proceedings of the 29th International Conference on Computational Linguistics*, pages 2166–2187, Gyeongju, Republic of Korea. International Committee on Computational Linguistics.
- Naihao Deng, Xinliang Zhang, Siyang Liu, Winston Wu, Lu Wang, and Rada Mihalcea. 2023. [You are what you annotate: Towards better models through annotator representations](#). In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 12475–12498, Singapore. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*.
- Xuemei Dong, Chao Zhang, Yuhang Ge, Yuren Mao, Yunjun Gao, Jinshu Lin, Dongfang Lou, et al. 2023. C3: Zero-shot text-to-sql with chatgpt. *arXiv preprint arXiv:2307.07306*.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.
- Dawei Gao, Haibin Wang, Yaliang Li, Xiuyu Sun, Yichen Qian, Bolin Ding, and Jingren Zhou. 2023. Text-to-sql empowered by large language models: A benchmark evaluation. *arXiv preprint arXiv:2308.15363*.
- Sirui Hong, Xiwu Zheng, Jonathan Chen, Yuheng Cheng, Jinlin Wang, Ceyao Zhang, Zili Wang, Steven Ka Shing Yau, Zijuan Lin, Liyang Zhou, et al. 2023. Metagtpt: Meta programming for multi-agent collaborative framework. *arXiv preprint arXiv:2308.00352*.
- Zijin Hong, Zheng Yuan, Qinggang Zhang, Hao Chen, Junnan Dong, Feiran Huang, and Xiao Huang. 2024. Next-generation database interfaces: A survey of llm-based text-to-sql. *arXiv preprint arXiv:2406.08426*.
- George Katsogiannis-Meimarakis and Georgia Koutrika. 2023. A survey on deep learning approaches for text-to-sql. *The VLDB Journal*, 32(4):905–936.
- Haoyang Li, Jing Zhang, Cuiping Li, and Hong Chen. 2023a. Resdsql: Decoupling schema linking and skeleton parsing for text-to-sql. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pages 13067–13075.
- Jinyang Li, Binyuan Hui, Ge Qu, Binhua Li, Jiaxi Yang, Bowen Li, Bailin Wang, Bowen Qin, Rongyu Cao, Ruiying Geng, et al. 2023b. Can llm already serve as a database interface. *A big bench for large-scale database grounded text-to-sqls*. *CoRR abs/2305.03111*.

- Zhishuai Li, Xiang Wang, Jingjing Zhao, Sun Yang, Guoqing Du, Xiaoru Hu, Bin Zhang, Yuxiao Ye, Ziyue Li, Rui Zhao, et al. 2024. Pet-sql: A prompt-enhanced two-stage text-to-sql framework with cross-consistency. *arXiv preprint arXiv:2403.09732*.
- Yixin Liu and Pengfei Liu. 2021. Simcls: A simple framework for contrastive learning of abstractive summarization. *arXiv preprint arXiv:2106.01890*.
- Henry B Mann and Donald R Whitney. 1947. On a test of whether one of two random variables is stochastically larger than the other. *The annals of mathematical statistics*, pages 50–60.
- R OpenAI. 2023. Gpt-4 technical report. *arXiv*, pages 2303–08774.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. 2022. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744.
- 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.
- Barbara Plank. 2022. The “problem” of human label variation: On ground truth in data, modeling and evaluation. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 10671–10682, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Ana-Maria Popescu, Oren Etzioni, and Henry Kautz. 2003. Towards a theory of natural language interfaces to databases. In *Proceedings of the 8th international conference on Intelligent user interfaces*, pages 149–157.
- Mohammadreza Pourreza and Davood Rafiei. 2023. Din-sql: Decomposed in-context learning of text-to-sql with self-correction. *arXiv preprint arXiv:2304.11015*.
- Jiexing Qi, Jingyao Tang, Ziwei He, Xiangpeng Wan, Yu Cheng, Chenghu Zhou, Xinbing Wang, Quanshi Zhang, and Zhouhan Lin. 2022. Rasat: Integrating relational structures into pretrained seq2seq model for text-to-sql. *arXiv preprint arXiv:2205.06983*.
- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J Liu. 2020. Exploring the limits of transfer learning with a unified text-to-text transformer. *The Journal of Machine Learning Research*, 21(1):5485–5551.
- Torsten Scholak, Nathan Schucher, and Dzmitry Bahdanau. 2021. Picard: Parsing incrementally for constrained auto-regressive decoding from language models. *arXiv preprint arXiv:2109.05093*.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.
- Bailin Wang, Richard Shin, Xiaodong Liu, Oleksandr Polozov, and Matthew Richardson. 2019. Rat-sql: Relation-aware schema encoding and linking for text-to-sql parsers. *arXiv preprint arXiv:1911.04942*.
- Bing Wang, Changyu Ren, Jian Yang, Xinnian Liang, Jiaqi Bai, Qian-Wen Zhang, Zhao Yan, and Zhoujun Li. 2023. Mac-sql: Multi-agent collaboration for text-to-sql. *arXiv preprint arXiv:2312.11242*.
- 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.
- Zelai Xu, Chao Yu, Fei Fang, Yu Wang, and Yi Wu. 2023. Language agents with reinforcement learning for strategic play in the werewolf game. *arXiv preprint arXiv:2310.18940*.
- Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Thomas L Griffiths, Yuan Cao, and Karthik Narasimhan. 2023. Tree of thoughts: Deliberate problem solving with large language models. *arXiv preprint arXiv:2305.10601*.
- Tao Yu, Rui Zhang, Kai Yang, Michihiro Yasunaga, Dongxu Wang, Zifan Li, James Ma, Irene Li, Qingning Yao, Shanelle Roman, et al. 2018. Spider: A large-scale human-labeled dataset for complex and cross-domain semantic parsing and text-to-sql task. *arXiv preprint arXiv:1809.08887*.
- John M. Zelle and Raymond J. Mooney. 1996. Learning to parse database queries using inductive logic programming. In *Proceedings of the Thirteenth National Conference on Artificial Intelligence - Volume 2*, pages 1050–1055.
- Ruiqi Zhong, Tao Yu, and Dan Klein. 2020. Semantic evaluation for text-to-sql with distilled test suites. *arXiv preprint arXiv:2010.02840*.
- Victor Zhong, Caiming Xiong, and Richard Socher. 2017. Seq2sql: Generating structured queries from natural language using reinforcement learning. *arXiv preprint arXiv:1709.00103*.

## A Appendix

### A.1 Significance Test

We divided the generated SQL by several strategies in Table 4 into 10 equal parts and calculated the execution accuracy for each. To test whether our strategy can indeed improve execution accuracy, we conduct a significance test between the “CoT” and “3R-Lp+PoT” strategies. The null hypothesis of the test is that the median execution accuracy obtained by the two strategies is the same. The Mann-Whitney U Test (Mann and Whitney, 1947) is a non-parametric statistical method used to compare whether there is a significant difference in the medians of two independent samples. Compared to the Analysis of Variance (ANOVA), it does not require the data to be normally distributed, making it suitable for small samples or data with unknown distribution.

The  $p$ -value of the test is 0.0024, which is below the commonly accepted significance level of 0.05. Therefore, we have reason to reject the null hypothesis, indicating that the “3R-Lp+PoT” strategy leads to a significant performance improvement.

**Effects of the number of “reviewer” agents.**

### A.2 Effects of $k$ in $k$ -shot.

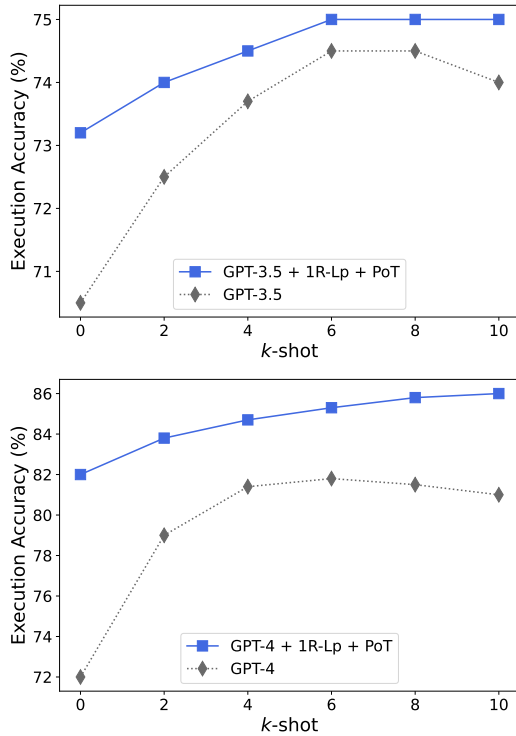


Figure 2:  $k$ -shot Sensitivity Analysis.

We test various  $k$  values on 200 random samples

from Spider-Dev. As shown in Figure 2, compared to CoT, the performance of the  $R^3$  system remains relatively stable regardless of the number of examples, which corroborates our previous findings from the 0-shot experiments with Llama-3.

### A.3 Spider Error Cases

| Error Type | Question, Gold & Prediction                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 | Reason                                                                                                                                                                           |
|------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| DB Value   | <p><b>Q:</b> Find the last name of the students who currently live in the state of North Carolina but have not registered in any degree program.</p> <p><b>Gold:</b> SELECT ... WHERE<br/>T2.state_province_county<br/>= 'NorthCarolina' EXCEPT ...</p> <p><b>Pred:</b> SELECT ... WHERE<br/>T2.state_province_county<br/>= 'North Carolina' EXCEPT ...</p>                                                                                                                                                                                                                                                 | The filtering condition in the question does not match the database value, string "NorthCalifornia" in database do not have a space in between.                                  |
| Gold Error | <p><b>Q:</b> What are the first names of all players, and their average rankings?</p> <p><b>Gold:</b> SELECT avg(ranking), T1.first_name<br/>FROM players AS T1 JOIN rankings AS T2 ON<br/>T1.player_id = T2.player_id<br/>GROUP BY T1.first_name</p> <p><b>Pred:</b> SELECT avg(ranking), T1.first_name<br/>FROM players AS T1 JOIN rankings AS T2 ON<br/>T1.player_id = T2.player_id<br/>GROUP BY T1.player_id</p>                                                                                                                                                                                        | The individuals in the table can be uniquely determined by column player_id not first_name, when GROUP BY.                                                                       |
| Gold Error | <p><b>Q:</b> Find the id and cell phone of the professionals who operate two or more types of treatments.</p> <p><b>Gold:</b> SELECT T1.professional_id,<br/>T1.cell_number FROM Professionals AS T1<br/>JOIN Treatments AS T2 ON<br/>T1.professional_id = T2.professional_id<br/>GROUP BY T1.professional_id<br/>HAVING count(*) &gt;= 2</p> <p><b>Pred:</b> SELECT T1.professional_id,<br/>T1.cell_number FROM Professionals AS T1<br/>JOIN Treatments AS T2 ON<br/>T1.professional_id = T2.professional_id<br/>GROUP BY T1.professional_id HAVING<br/>COUNT(DISTINCT T2.treatment_type_code) &gt;= 2</p> | The gold only finds professionals who have two or more records in the treatment table does not ensure that the records are for different types of treatments                     |
| Ambiguity  | <p><b>Q:</b> What are the names and ids of all makers with more than 3 models?</p> <p><b>Gold:</b> SELECT T1.FullName, T1.Id FROM<br/>CAR_MAKERS AS T1 JOIN MODEL_LIST AS T2 ON<br/>T1.Id = T2.Maker GROUP BY T1.Id HAVING<br/>count(*) &gt; 3;</p> <p><b>Pred:</b> SELECT T1.Maker, T1.Id FROM<br/>CAR_MAKERS AS T1 JOIN MODEL_LIST AS T2 ON<br/>T1.Id = T2.Maker GROUP BY T1.Id HAVING<br/>count(*) &gt; 3;</p>                                                                                                                                                                                           | Both column "Maker" and column "FullName" can answer the question about the "names of makers" in the query.                                                                      |
| Imprecise  | <p><b>Q:</b> What are the arriving date and the departing date of the dogs who have gone through a treatment?</p> <p><b>Gold:</b> SELECT DISTINCT T1.date_arrived,<br/>T1.date_departed FROM Dogs AS T1 JOIN<br/>Treatments AS T2 ON T1.dog_id = T2.dog_id</p> <p><b>Pred:</b> SELECT DISTINCT T1.date_arrived,<br/>T1.date_departed, T1.Name FROM Dogs AS T1<br/>JOIN Treatments AS T2 ON T1.dog_id =<br/>T2.dog_id</p>                                                                                                                                                                                    | The question do not require listing the specific names of the students, but only ask to list the students' arrival and departure dates. This falls under information redundancy. |

Table 6: Spider error cases.

## A.4 BIRD Error Cases

| Error Type | Question, Gold & Prediction                                                                                                                                                                                                                                                                                                                                                                                                                | Reason                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| DB Value   | <b>Q:</b> How many cards with unknown power that can't be found in foil is in duel deck A?<br><b>Gold:</b> SELECT SUM(CASE WHEN power LIKE '%%' OR power IS NULL THEN 1 ELSE 0 END) FROM cards WHERE hasFoil = 0 AND duelDeck = 'a'<br><b>Pred:</b> SELECT COUNT(*) FROM cards WHERE (power IS NULL OR power = '*') AND hasFoil = 0 AND duelDeck = 'a'                                                                                     | Values in database cannot exact match with the question and evidence.                                                                                                                                                                                                                                                                                                                                                                                                |
| Gold Error | <b>Q:</b> How many artists have designed a card with a black border color and is available in both "arena" and "mtgo" printing type?<br><b>Gold:</b> SELECT COUNT(CASE WHEN availability LIKE '%arena,mtgo%' THEN 1 ELSE NULL END) FROM cards<br><b>Pred:</b> SELECT COUNT(DISTINCT artist) FROM cards WHERE borderColor = 'black' AND availability LIKE '%arena,mtgo%'                                                                    | The gold missed one filtering condition.                                                                                                                                                                                                                                                                                                                                                                                                                             |
| Ambiguity  | <b>Q:</b> Among black card borders, which card has full artwork?<br><b>Gold:</b> SELECT id FROM cards WHERE borderColor = 'black' AND isFullArt = 1<br><b>Pred:</b> SELECT name FROM cards WHERE borderColor = 'black' AND isFullArt = 1;                                                                                                                                                                                                  | Both column "name" and column "id" can answer the question about the "the card has full artwork" in the query.                                                                                                                                                                                                                                                                                                                                                       |
| Logic      | <b>Q:</b> Which cities have the top 5 lowest enrollment number for students in grades 1 through 12?<br><b>Gold:</b> SELECT T2.City FROM frpm AS T1 INNER JOIN schools AS T2 ON T1.CDSCode = T2.CDSCode GROUP BY T2.City ORDER BY SUM(T1.'Enrollment (K-12)') ASC LIMIT 5<br><b>Pred:</b> SELECT T2.City, T1."Enrollment (K-12)" FROM frpm AS T1 JOIN schools AS T2 ON T1.CDSCode = T2.CDSCode ORDER BY T1."Enrollment (K-12)" ASC LIMIT 5; | The Pred query is incorrect because it lacks grouping and aggregation. It directly selects cities and their corresponding "Enrollment (K-12)" without summing up the enrollment numbers across all schools within each city. As a result, it returns the cities of the five schools with the lowest enrollment numbers, which does not accurately answer the question seeking the top 5 cities with the lowest total enrollment for students in grades 1 through 12. |

Table 7: Bird error cases.

## A.5 Interactive Correction Cases

| Before                                                                                                                                                                                                                    | After                                                                                                                                                                                                          | Reason                                                                                                                                                                                                                                                                                                                                                                                                                                        |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Q:</b> What is the phone number of the school that has the highest number of test takers with an SAT score of over 1500?                                                                                               |                                                                                                                                                                                                                |                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| <b>Before:</b><br>SELECT schools.Phone<br>FROM satscores<br>JOIN schools ON<br>satscores.cds =<br>schools.CDSCode<br>WHERE satscores.NumGE1500<br>> 1500<br>GROUP BY schools.School<br>ORDER BY COUNT(*) DESC<br>LIMIT 1; | <b>After:</b><br>SELECT schools.Phone<br>FROM satscores<br>JOIN schools ON<br>satscores.cds =<br>schools.CDSCode<br>ORDER BY<br>satscores.NumGE1500 DESC<br>LIMIT 1;                                           | The WHERE clause is checking if the number of test takers with an SAT score over 1500 is greater than 1500, which is not what the question is asking. The question is asking for the school with the highest number of test takers who scored over 1500. Additionally, the GROUP BY clause is grouping by the school name, which may not be unique. It would be better to group by the school code (CDSCode) which is unique for each school. |
| <b>Q:</b> List the top five schools, by descending order, from the highest to the lowest, the most number of Enrollment (Ages 5-17). Please give their NCES school identification number.                                 |                                                                                                                                                                                                                |                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| <b>Before:</b><br>SELECT T1."School Name",<br>T2.NCESSchool<br>FROM frpm AS T1<br>JOIN schools AS T2 ON<br>T1.CDSCode = T2.CDSCode<br>ORDER BY T1."Enrollment<br>(Ages 5-17)" DESC<br>LIMIT 5;                            | <b>After:</b><br>SELECT T2.NCESSchool<br>FROM frpm AS T1<br>JOIN schools AS T2 ON<br>T1.CDSCode = T2.CDSCode<br>ORDER BY T1."Enrollment<br>(Ages 5-17)" DESC<br>LIMIT 5;                                       | The question asks for the top five schools with the most number of Enrollment (Ages 5-17), and their NCES school identification number. However, the SQL query provided returns the school name, which is not required by the question.                                                                                                                                                                                                       |
| <b>Q:</b> find the package option of the tv channel that do not have any cartoon directed by Ben Jones.                                                                                                                   |                                                                                                                                                                                                                |                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| <b>Before:</b><br>SELECT<br>TV_Channel.Package_Option<br>FROM TV_Channel<br>LEFT JOIN Cartoon ON<br>TV_Channel.id =<br>Cartoon.Channel<br>WHERE Cartoon.Directed_by<br>IS NOT 'Ben Jones';                                | <b>After:</b><br>SELECT<br>TV_Channel.Package_Option<br>FROM TV_Channel<br>WHERE TV_Channel.id NOT IN<br>(<br>SELECT Cartoon.Channel<br>FROM Cartoon<br>WHERE Cartoon.Directed_by<br>= 'Ben Jones');           | The current query is selecting all channels where the director is not Ben Jones, but it should exclude channels that have any cartoons directed by Ben Jones. This can be achieved by using a subquery to find all channels that have cartoons directed by Ben Jones, and then excluding these channels in the main query. Here is the corrected SQL:                                                                                         |
| <b>Q:</b> List the earnings of poker players in descending order.                                                                                                                                                         |                                                                                                                                                                                                                |                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| <b>Before:</b><br>SELECT * FROM poker_player<br>ORDER BY Earnings DESC;                                                                                                                                                   | <b>After:</b><br>SELECT Earnings FROM<br>poker_player ORDER BY<br>Earnings DESC;                                                                                                                               | The question only asks for the earnings of poker players. Therefore, we should only select the 'Earnings' column to make the result more concise and relevant to the question.                                                                                                                                                                                                                                                                |
| <b>Q:</b> How many countries speak both English and Dutch?                                                                                                                                                                |                                                                                                                                                                                                                |                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| <b>Before:</b><br>SELECT COUNT(DISTINCT<br>CountryCode)<br>FROM countrylanguage<br>WHERE Language = 'English'<br>AND CountryCode IN (<br>SELECT CountryCode<br>FROM countrylanguage<br>WHERE Language = 'Dutch')          | <b>After:</b><br>SELECT COUNT(*)<br>FROM (<br>SELECT CountryCode<br>FROM countrylanguage<br>WHERE Language = 'English'<br>INTERSECT<br>SELECT CountryCode<br>FROM countrylanguage<br>WHERE Language = 'Dutch') | The SQL query provided uses the 'IN' clause, which selects countries that speak either English or Dutch. However, the question asks for countries speaking both languages. Using IN results in a union of sets (either English or Dutch), not an intersection (both English and Dutch). This could be done with two separate queries joined by an 'INTERSECT' clause.                                                                         |

Table 8: Interactive correction cases.

## A.6 Prompt Examples

Describe how you understand the question based on the evidence, and help me write an SQL to answer the question.

### EVIDENCE: {evidence}

### USER\_QUESTION: {question}

### RELATED SQL:

{related\_sql}

### DATABASE STRUCTURE:

{schema}

Prompt 1: CoT-SQL-Writer

Write an to answer the question.

Program of Thoughts (PoT) is a variant of Chain of Thought (CoT), pre-generating Python code to assist in the creation of SQL. Please apply PoT (and PoT only) before generating an SQL.

In your python code, `Table %s` is stored in `db\_dict['%s']`, `db\_dict` is of type dict[pandas.DataFrame].

### RELATED SQL:

{related\_sqls}

### DATABASE STRUCTURE:

{schema}

### EXAMPLES:

QUESTION: What is %s in the earliest year and what year was it?

SQL:

```
earliest_year = db_dict[%s]['Year'].min()
```

```
year_filtered_data = step1_result[step1_result['Year'] ==  
earliest_year]
```

```
result = year_filtered_data[[%s, 'Year']]
```

```
```sql
```

```
SELECT T1.%s, T2.Year FROM %s AS T1 JOIN %s AS T2 ON T1.Id = T2.Id  
WHERE T2.Year = (SELECT min(YEAR) FROM %s);
```

```
```
```

QUESTION: Show names for all %s except for %s having a %s in year 2023.

SQL:

```
%s_2023 = db_dict['%s'][db_dict['%s']['year'] == '2023']
```

```
result = db_dict[%s][~db_dict[%s][%s].isin(%ss_2023[%s])]
```

```
```sql
```

```
SELECT name FROM %s EXCEPT SELECT T2.name FROM %s AS T1 WHERE T1.  
year = 2023
```

```
```
```

QUESTION: Find the %s that %s is A and B?

SQL:

```

condition_a_data = db_dict[%s][db_dict['Cartoon'][%s] == 'A']
condition_b_data = db_dict[%s][db_dict['Cartoon'][%s] == 'B']
result = pd.merge(condition_a_data, condition_b_data, how='inner')
```sql
SELECT T1.%s FROM %s AS T1 WHERE %s = 'A'
INTERSECT
SELECT T1.%s FROM %s AS T1 WHERE %s = 'B'
```

### EVIDENCE: {evidence}
### USER_QUESTION: {question}
### SQL:

```

#### Prompt 2: PoT-SQL-Writer

You are the manager of a Database project. You are going to invite {n} experts to review an SQL query.  
Who would you invite?

considering:  
(1) the domain of this database;  
(2) the structure of this SQL.  
Please write your invitation as a JSON format dictionary, Enclose the JSON within ```json...```.

### DATABASE STRUCTURE:

{schema}

### QUESTION: {question}

### SQL:

{pred\_sql}

### EXAMPLES:

```json

```

{
  "Reviewer PVsg": "Data Analyst in automotive industry",
  "Reviewer 2KtR": "Senior Database Engineer specialized in writing
    various clauses",
  "Reviewer LmN3": "Senior Database Engineer specialized in writing
    filtering conditions"
}
```

```

### INVITATION:

#### Prompt 3: Invitation