

Enhancing NLIDBs: Advancing from Text-to-SQL to Text-to-Multi-SQL

Anonymous ACL submission

Abstract

Text-to-SQL is a key part of Natural Language Interfaces to Databases (NLIDB), helping non-technical users query databases. However, it has one significant limitation: it assumes users know that SQL queries return data in a table format. This often leads to problems when users ask for data that cannot be handled by a single query. To address this, we propose Text-to-Multi-SQL, which uses multiple SQL queries to meet complex user needs. We created SpiderS, the first Text-to-Multi-SQL dataset, based on the Spider dataset. It includes 7,000 training examples, 1,024 validation examples, and 2,147 test examples, created using a mix of manual and GPT4o generated data. Tests show that current models struggle with multiple SQL generation, and even advanced models perform worse (0.167-0.327 drop in accuracy) on SpiderS. We also found that models are very sensitive to prompts—switching from “generate one SQL” to “generate one or multiple SQL” significantly reduces their performance.

1 Introduction

The research on NLIDB (Natural Language Interface to DataBases) has always been a hot topic in the fields of databases and Natural Language Processing (NLP) (Abbas et al., 2022; Pourreza and Rafiei, 2023; Hong et al., 2024; Zhang et al., 2024). In the study of NLIDB, Text-to-SQL is considered one of the key components (Abbas et al., 2022; Qin et al., 2022; Deng et al., 2022; Katsogiannis-Meimarakis and Koutrika, 2023a; Hong et al., 2024; Zhang et al., 2024). Many studies suggest that Text-to-SQL technology allows users without database knowledge to easily access data from a database (Iacob et al., 2020; Abbas et al., 2022; Qin et al., 2022; Katsogiannis-Meimarakis and Koutrika, 2023a; Gao et al., 2023; Pourreza et al., 2024; Hong et al., 2024; Zhang et al., 2024; Zhu et al., 2024).

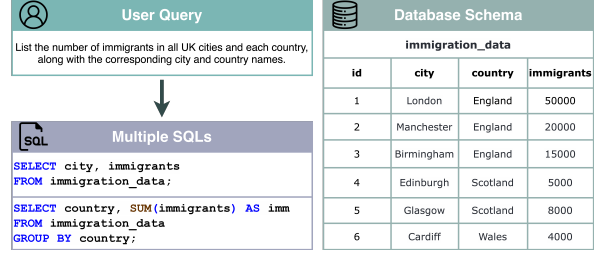


Figure 1: An example of Text-to-Multi-SQL.

However, in our practice, we discovered an implicit assumption in Text-to-SQL: users are expected to understand that the data returned by each SQL query is expected to be a tabular format, typically a rectangular table with a consistent column structure (though there may be variations, such as missing values in sparse data or more complex structures in nested data). For many users without SQL knowledge, they are unaware of this assumption, which leads to the problem that their requests often cannot be fulfilled by a single SQL query, or even by SQL queries at all. This can cause a problem: users may request data that exceeds the capabilities of a single SQL query.

Figure 1 shows a practical example from an NLIDB-based census flow data platform. The data the user needs cannot be fulfilled by a single SQL query; however, two SQL queries can perfectly meet the user’s needs. Although it’s possible to force the query to be completed with a single SQL, the returned data is neither elegant nor easy to understand. Both we and the users agree that using two SQL queries to return two separate result sets would be a better approach.

Based on this finding, we believe it is necessary to extend traditional Text-to-SQL to Text-to-Multi-SQL to better meet the diverse data needs of users. To this end, we chose the Spider (Yu et al., 2018) Text-to-SQL dataset as the basis and constructed the first Text-to-Multi-SQL dataset: SpiderS. Like

the Spider dataset, SpiderS contains 7,000 training examples, 1,024 validation examples, and 2,147 test examples. The data was constructed using a combination of manual annotation and data generated by Large Language Models (LLMs), with around 1,100 manually annotated examples used to guide the LLM in generating additional paired data, where each question could be paired with multiple SQL queries. During generation, we rewrote the original Spider questions and SQLs to ensure the diversity of the SpiderS dataset. To the best of our knowledge, SpiderS is the first Text-to-Multi-SQL dataset, and this paper is the first dedicated to the study of Text-to-Multi-SQL.

Experimental results show that many existing Text-to-SQL models are unable to generate multiple SQL queries. While LLM-based Text-to-SQL models can output multiple SQL queries by modifying the prompt, their performance drops significantly on the SpiderS dataset because they may not be trained on a Text-to-Multi-SQL dataset. The absolute drop in execution accuracy ranges from 0.167 to 0.327. Interestingly, most models show poor robustness to prompts that request one or multiple SQL queries. When we simply changed the prompt from ‘generate one SQL’ to ‘generate one or multiple SQL,’ the performance of most models on the original Spider dataset decreased to varying degrees, with the largest drop seen in the CHESS (Talaie et al., 2024) model—an SOTA open source Text-to-SQL model on the BIRD (Li et al., 2024)—whose execution accuracy dropped from 0.759 to 0.462.

Overall, the contributions of this paper are as follows:

1. We systematically study Text-to-Multi-SQL and construct the first Text-to-Multi-SQL dataset, SpiderS, based on the Spider Text-to-SQL dataset.
2. We found that Text-to-SQL models are highly vulnerable to prompt modifications: simply changing the instruction from ‘generate one SQL’ to ‘generate one or multiple SQLs’ leads to a significant performance drop.
3. We found that existing Text-to-SQL models/LLMs fail to address the need for Text-to-Multi-SQL, with some unable to generate multiple SQL queries. Even capable models show a significant accuracy drop on SpiderS compared to the Spider dataset.

2 Relate Work

2.1 Natural Language Interface to Databases

Research in Natural Language Interface to Databases (NLIDB) focuses on enabling users to query databases using natural language, with key tasks including Text2SQL, Text-to-Graph, and Text-to-SPARQL. The Text2SQL (Katsogiannis-Meimarakis and Koutrika, 2023b) task, which converts natural language questions into SQL queries, is the most prominent, addressing challenges such as query intent interpretation, schema mapping, and handling complex SQL constructs like joins and subqueries. Recent advancements in deep learning, particularly transformer (Vaswani, 2017) and large language (Devlin, 2018; Achiam et al., 2023) models, have significantly improved performance (Scholak et al., 2021; Qi et al., 2022; Li et al., 2023; Gao et al., 2024b) in this area. In addition to Text2SQL, other tasks such as Text-to-Graph (Lawrence, 2024) and Text-to-SPARQL (Luo et al., 2023) aim to extend NLIDB systems to graph and semantic web databases, each introducing unique challenges related to graph structure and ontology-based querying. There is a growing emphasis on improving related models’ ability to handle more complex domain-specific queries and to provide accurate context-sensitive responses.

2.2 Popular Text-to-SQL Datasets

As one of the most important studies in human-computer interaction systems, Text2SQL has attracted many people to explore. Several popular databases are proposed for converting natural language queries into structured query language. WikiSQL (Zhong et al., 2017) is a large-scale dataset for converting natural language questions into SQL queries, containing 80,000 question-SQL pairs over tables extracted from Wikipedia. Spider (Yu et al., 2018) is a large-scale, cross-domain dataset for SQL query generation, with a focus on multi-table queries, complex conditions, and aggregation. BIRD (Li et al., 2024) is a benchmark for evaluating multi-lingual and cross-lingual SQL query generation models, particularly focusing on database retrieval across different languages. NL2GQL (Zhou et al., 2024) is a dataset that pairs natural language questions with corresponding GQL queries, designed to help develop models for querying graph databases, including tasks like node selection, graph traversal, and pathfinding.

SQL-eval¹ is an open-source PostgreSQL evaluation dataset constructed based on Spider (Yu et al., 2018), which is used to evaluate Text-to-SQL models’ ability to generate executable SQL queries from natural language questions. However, all of the above datasets focus on convert one natural language question into one SQL query. It does not apply to many real-world scenarios where the answer to the user’s question should come from multiple SQL.

2.3 Limitations of Text-to-SQL Systems

Recent studies have advanced Text-to-SQL systems from several complementary angles. However, it’s still hard for these methods to address complex tasks. Biswal et al. (2024) integrate Retrieval-Augmented Generation (RAG) with database tables to propose a unified, general-purpose paradigm for answering natural language questions over databases, aiming to provide a better user experience than traditional text-to-SQL approaches. Likewise, CHASE (Ma et al., 2025) enhances hybrid query execution by bridging structured and unstructured data. Renggli et al. (2025) analyze the limitations of Text-to-SQL from an evaluation perspective, which can lead to both prediction and evaluation errors. However, none of these papers involved Text-to-Multi-SQL.

3 Task Definition

Traditional Text-to-SQL approaches typically assume a one-to-one mapping between question and SQL query. However, in many real-world applications, a single query is insufficient to fully capture the intent of a complex question. For example, multi-faceted questions often require retrieving information from different perspectives, applying multiple filtering conditions, or aggregating distinct sets of records.

Text-to-Multi-SQL aims to convert a natural language question Q into one or multiple SQL queries $y_1, y_2, \dots$ that retrieve the correct results from a given database D . A database is formally defined as $D = \langle C, T \rangle$, where $C = \{c_1, c_2, \dots, c_m\}$ represents the set of column names and $T = \{t_1, t_2, \dots, t_m\}$ represents the set of table names. Given D and Q , the objective is to generate the most appropriate SQL query (or queries set), such that executing $y_1, y_2, \dots$ on D yields the expected

answer. The query generation process can be expressed as

$$\{y_1, y_2, \dots\} = f(Q, D|\theta), \quad (1)$$

where $f(\cdot|\theta)$ is a model parameterized by θ , which maps the natural language question Q and database schema D to a set of SQL queries $\{y_1, y_2, \dots\}$.

For some questions that can be solved using either multiple SQL queries or a single SQL query, we prepare two sets of gold SQL queries. As long as the prediction matches any one of the gold SQL sets, we consider the task successful. For example, in all cases generated by the 6th generation strategy in Appendix A, we have prepared two gold SQL sets for evaluation. However, during training, we use the multi-SQL gold SQL set because its query results better align with the user’s intent.

4 Dataset Construction

Figure 2 illustrates our data annotation process.

4.1 Question and SQL Annotation

SQL pattern coverage. Before beginning the annotation process, our team surveyed multiple SQL queries commonly found in real-world applications and identified 14 distinct patterns (e.g., requiring an aggregated result and columns before aggregated (type 6), grouping the same table based on different attributes (type 7), and sorting the same data based on different ordering (type 9) in Appendix A). Each pattern was carefully documented in our annotation manual, detailing the associated SQL clauses and typical ways users might phrase corresponding questions. This approach ensured that the final dataset would capture a broad spectrum of query complexities and structures.

Question clarity. Consistent with the Spider dataset, we aimed to avoid ambiguous questions that either lacked sufficient details or required knowledge outside the database to answer. Instead, we focused on queries that explicitly stated the necessary conditions and returned values. For example, the question "Find all the teacher names and also all the classrooms for math courses" can be interpreted in two ways: either as a request for teachers and classrooms for math courses, or as separate requests for all teacher names and all math course classrooms. To avoid ambiguity, we use phrases like "first ..., then ..." or "separately list ..." to clarify that multiple distinct SQL queries are needed.

¹<https://github.com/defog-ai/sql-eval>

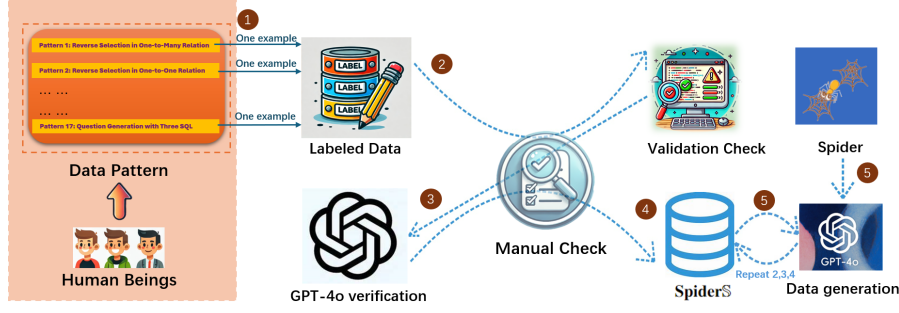


Figure 2: Data annotation process.

Annotation steps. We began by creating a detailed data annotation manual that described each of the 14 SQL patterns, including examples of both questions and their corresponding SQL statements. Two Computer science students (paper authors) studying in the US, all proficient in SQL, were then tasked with labeling one example for each pattern across different databases. These examples served as references for One-Shot Learning in later data generation stages.

After the initial round of labeling, we conducted a manual cross-check to detect obvious inconsistencies or errors. Next, the annotated data was run through an automatic validation script (described in Section 4.2) that checked structural correctness (e.g., matching pattern design, proper use of SQL keywords) and basic logical consistency. Finally, we used GPT-4o to assess whether each user question was clearly stated and matched the corresponding SQL queries accurately. Any data flagged during this process was subjected to a second manual review to resolve potential issues.

Annotation tools. In order to conduct data annotation more efficiently, we establish a website which display the database schema and provide some possibly useful examples as annotation references. After modifying or adding SQLs on the webpage, we can directly execute the SQLs and check whether the result is as expected. The website also is connected to GPT-4o, which may provide meaningful suggestions to make the modified questions without ambiguity. The website also provides data comparison, allowing users to intuitively view the differences between the original questions and the new questions, facilitating modification and review.

4.2 Data Review

Automatic validation script. To maintain annotation consistency and reduce human error, we developed and iteratively refined an automatic valida-

tion script. The script performed several checks:

- **Syntactic correctness:** Verified that all SQL queries could be parsed and executed without errors.
- **Pattern matching:** Ensured each query corresponded to one of the 15 predefined SQL patterns and did not include extraneous or missing clauses.
- **Schema alignment:** Checked that table names and column names used in the SQL were present in the database schema.

Entries that did not pass any of these checks were returned to the annotators for correction or deletion.

GPT4o-Based validation. In addition to the rule-based checks, we employed GPT-4o as a second layer of validation. GPT-4o was given a pair consisting of a user query and its corresponding SQL statement. It was then asked to:

1. Evaluate whether the question was sufficiently clear and unambiguous.
2. Verify if the SQL query logically aligned with the question.
3. Flag any potential misalignment or unclear phrasing.

Data flagged by GPT-4o for potential issues—such as vague question wording or mismatches in conditions—was forwarded to a human reviewer for final judgment. This hybrid workflow ensured a more robust validation process than either method could achieve in isolation.

4.3 Construct the Final Dataset

After one round of manual validation, one round of automatic validation, and one round of GPT-4o

SQL keyword	Spider	SpiderS
WHERE	51.37%	52.34%
GROUP	29.02%	31.66%
ORDER	15.70%	14.93%
LIMIT	39.40%	35.10%
UNION	1.02%	0.11%
INTERSECT	3.91%	0.22%
EXCEPT	3.49%	0.37%
JOIN	58.55%	39.33%

Table 1: Percentage of SQL keywords in Spider and SpiderS.

verification, we obtained a high-quality set of annotated question-SQL pairs. Then, we generated additional data using GPT-4o based on One-Shot learning if a proper example was found; otherwise, we used the Zero-Shot learning approach to generate extra data.

We expanded the dataset to about twice the size of the original Spider dataset and applied the same multi-step validation process, including automatic checks, manual review, and GPT-4o-based validation, to ensure accuracy and consistency. For most data that failed automatic validation or GPT-4o checks, we removed it directly to save annotation time. Through manual inspection and review, we ensured that the new dataset remained consistent in quality with the original Spider dataset. To maintain balance, we set a limit so that no single SQL pattern made up more than 15% of the dataset. This prevented overrepresentation of certain query types and ensured diversity. After these rigorous validation steps, we produced a well-verified, high-quality Text-to-Multiple-SQL dataset.

4.4 Dataset Statistics

Our dataset follows the same structure as the Spider dataset, with 7,000 training samples, 1,034 validation samples, and 2,147 test samples. Of these, 1,130 examples were manually created, and the remaining 9,051 were automatically generated. On average, each question requires 2.125 SQL queries to answer. Specifically, 1,266 samples need 3 SQL queries, while 8,915 samples require 2. Table 1 compares the frequency of SQL keywords between the two datasets. Compared to the original Spider, SpiderS has similar frequencies for keywords like WHERE, GROUP, ORDER, and LIMIT. However, JOIN, UNION, INTERSECT, and EXCEPT are less frequent in SpiderS because we opted for simpler SQL queries in SpiderS, aiming to avoid overly

complex queries, as multiple complex queries are less common in real user scenarios.

5 Evaluation Metrics

Execution Score (ES). In order to evaluate the correctness of each of the multiple generated SQL queries, we propose a metric called Execution Score (ES). For the i^{th} evaluation question requiring multiple SQL queries to address, given P_i and G_i as the executed result sets of the predicted SQLs and the ground-truth SQLs separately, ES_i is defined as the intersection of the result set P_i and G_i , relative to the maximum between the size of the two result sets. This metric is applied to evaluate the ability of generating meaningful and correct SQL queries matching part of the question. Therefore, we consider the correctness for each single predicted SQL and assign a higher score if the executed results of predicted SQLs match more ground-truth ones. The final ES is the accumulation of ES_i , which is shown below:

$$ES_i = \frac{|P_i \cap G_i|}{\max(|P_i|, |G_i|)}, \quad (2)$$

where:

- P_i : The set of the i^{th} predicted SQL queries.
- G_i : The set of the i^{th} ground-truth SQL queries.
- $P_i \cap G_i$: The set of SQL query pairs (p, g) for the i^{th} question such that the execution result of p and g is identical.
- $|\cdot|$: The cardinality (size) of a set.

Execution Accuracy (EX). EX is formally defined as the proportion of examples in the evaluation set for which the executed results of both the predicted and ground-truth SQLs for a single question are identical, relative to the overall number of SQLs. However, in our task definition, we require multiple SQLs to address one single question. Therefore, following the previous definition of Execution Score (ES), the i^{th} question in the evaluation set is considered solved only when ES_i is equal to 1. That is, the execution result of each of the predicted SQLs can match the one of a corresponding ground-truth SQL. The computation of EX can be expressed as follows:

$$EX = \frac{\sum_{i=1}^N \mathbb{M}(ES_i)}{N}, \quad (3)$$

Models	Spider		Spider $\mathbb{M}$		Spider $\mathbb{S}$		Spider $_{small}$		Spider $\mathbb{M}_{small}$		Spider $\mathbb{S}_{small}$	
	EX	ES	EX	ES	EX	ES	EX	ES	EX	ES	EX	ES
Closed-source Models												
GPT-4o	-	-	-	-	0.497	0.648	0.733	0.733	0.572	0.625	0.519	0.667
Claude-3.5-Sonnet	-	-	-	-	-	-	0.629	0.629	0.589	0.589	0.302	0.461
o3-mini	-	-	-	-	-	-	0.737	0.737	0.722	0.723	0.525	0.650
Open-source Models												
Qwen2.5-Coder-7B-Instruct	0.756	0.756	0.674	0.694	0.442	0.596	0.761	0.761	0.678	0.697	0.454	0.608
Qwen2.5-Coder-32B-Instruct	0.766	0.766	0.675	0.695	0.503	0.647	0.779	0.779	0.689	0.705	0.515	0.659
CodeLlama-7b-Instruct-hf	0.123	0.123	0.082	0.089	0.038	0.069	0.125	0.125	0.070	0.078	0.033	0.065
CodeLlama-34b-Instruct-hf	0.177	0.177	0.221	0.224	0.088	0.141	0.172	0.172	0.216	0.218	0.086	0.140
Llama-3.1-8B-Instruct	0.628	0.628	0.621	0.628	0.392	0.549	0.622	0.622	0.615	0.620	0.411	0.567
Llama-3.1-405B-Instruct	-	-	-	-	-	-	0.787	0.787	0.748	0.750	0.534	0.688
DeepSeek-R1-Distill-Qwen-7B	0.394	0.394	0.320	0.347	0.195	0.322	0.375	0.375	0.300	0.330	0.195	0.330
DeepSeek-R1-Distill-Llama-8B	0.435	0.435	0.388	0.404	0.125	0.283	0.415	0.415	0.374	0.387	0.137	0.293
Mistral-Large-Instruct-2411	-	-	-	-	-	-	0.671	0.671	0.454	0.522	0.453	0.634
Fine-tuned Models												
Qwen2.5-Coder-7B-Instruct	-	-	0.786	0.787	0.725	0.831	-	-	0.784	0.785	0.731	0.837
CodeLlama-7b-Instruct-hf	-	-	0.728	0.728	0.656	0.781	-	-	0.728	0.729	0.651	0.781
Llama-3.1-8B-Instruct	-	-	0.761	0.761	0.706	0.811	-	-	0.753	0.753	0.724	0.822
DeepSeek-R1-Distill-Qwen-7B	-	-	0.664	0.664	0.604	0.734	-	-	0.660	0.661	0.615	0.744
DeepSeek-R1-Distill-Llama-8B	-	-	0.728	0.729	0.665	0.788	-	-	0.730	0.730	0.672	0.793
Text2SQL Models												
CHESS	-	-	-	-	-	-	0.759	0.759	0.462	0.589	0.559	0.699
XiYanSQL-QwenCoder-32B	0.849	0.849	0.809	0.814	0.669	0.782	0.852	0.852	0.823	0.827	0.685	0.799

Table 2: Performance Comparison of Various Models on Spider and Spider $\mathbb{S}$ Datasets.

where function $\mathbb{M}(\cdot)$ returns 1 only if ES_i is equal to 1, indicating that every executed result in predicted result set P_i can match a corresponding one in ground-truth result set G_i . The expression of $\mathbb{M}(\cdot)$ is shown below:

$$\mathbb{M}(ES) = \begin{cases} 1, & \text{if } ES_i = 1, \\ 0, & \text{otherwise.} \end{cases} \quad (4)$$

6 Experiments

6.1 Experimental Setup

Prompt. In terms of prompt design, the distinction lies in generating one SQL versus one or more SQL. When evaluating Spider $\mathbb{S}$, the latter must be used, whereas for Spider, both approaches are applicable. To analyze the impact of prompts, we employ different prompts on Spider to observe how the model’s performance is affected when the generated SQL is no longer restricted to a single query.

Dataset. We conducted our experiments in both the Spider and Spider $\mathbb{S}$ test sets. The test set distribution of the original Spider dataset consists of 470 easy, 857 medium, 463 hard, and 357 extra-hard samples, totaling 2,147 queries. In our

Spider $\mathbb{S}$ dataset, the distribution has slightly shifted, with 351 easy, 837 medium, 595 hard, and 364 extra-hard samples, while still maintaining a total of 2,147 queries. Despite these minor variations—such as a decrease in easy queries (470 $\rightarrow$ 351) and an increase in hard queries (463 $\rightarrow$ 595)—the overall difficulty distribution remains similar across both datasets.

To improve evaluation efficiency and reduce evaluation costs, we randomly sampled 1,000 examples each from the original Spider and Spider-S test sets to create a small test set for low-cost evaluation.

The evaluation tasks used in this study are shown as follows:

- **Spider $\mathbb{S}$:** The complete Spider $\mathbb{S}$ test set(2,147-item).
- **Spider $\mathbb{S}_{small}$:**The small Spider $\mathbb{S}$ test set(1,000-item).
- **Spider:** The complete Spider test set(2,147-item) with a prompt for generating one SQL.
- **Spider $_{small}$:** The small Spider test set (1,000-item) with a prompt for generating one SQL.
- **Spider $\mathbb{M}$:** The complete Spider test set(2,147-item) with a prompt for generating one or multi

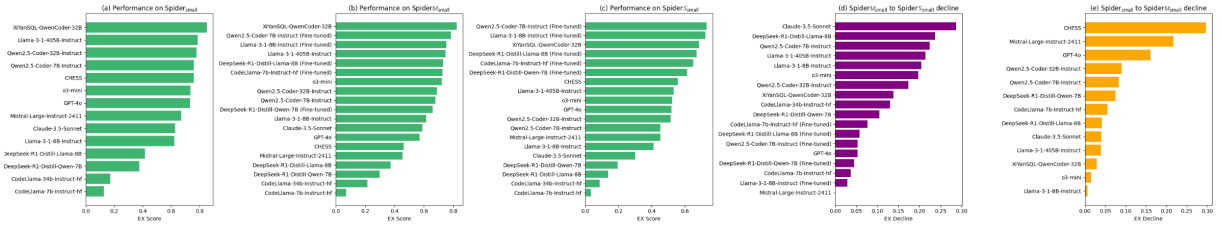


Figure 3: Model EX rankings on the Spider(a), SpiderM(b), and SpiderS(c), and the EX drop caused by task migration(d,e).

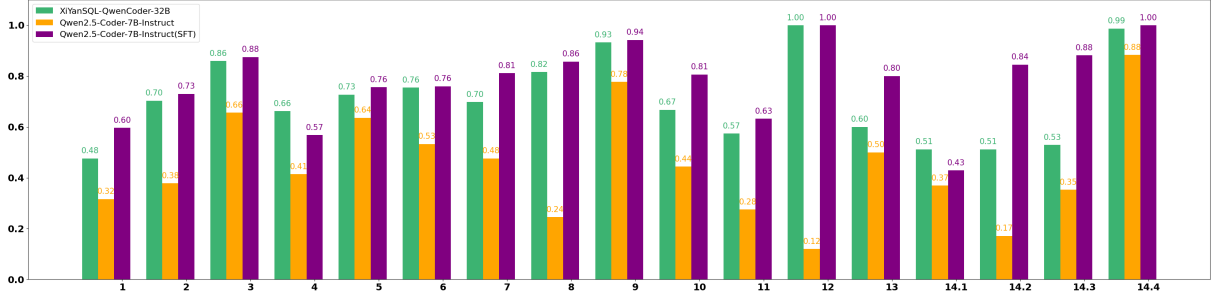


Figure 4: EX scores for different question types.

SQL.

- **SpiderM_{small}**: The small Spider test set (1,000-item) with a prompt for generating one or multi SQL.

Models. To comprehensively evaluate the Text-to-Multi-SQL capabilities of different models, we assessed a diverse set of models, including open-source, fine-tuned, closed-source, and specialized Text2SQL models. For open source models, we study Llama3.1 (Grattafiori et al., 2024), Qwen 2.5 Coder (Qwen et al., 2025), DeepSeek-R1-Distill-models(DeepSeek-AI et al., 2025), CodeLlama (Rozière et al., 2024). Among the closed-source model, we evaluate GPT-4o (OpenAI et al., 2024), Claude-3.5-Sonnet, and o3-mini. For text-to-SQL models, we choose CHESS (Talaie et al., 2024) and XiYanSQL-QwenCoder-32B (Gao et al., 2024a).

6.2 Main Results

Table 2 compares the performance of different models on the Spider and SpiderS datasets. Overall, XiYanSQL-QwenCoder-32B, which is specifically optimized for SQL generation, performs the best. Among the other models, Llama-3.1-405B-Instruct achieves the best results, with o3-mini close behind. The model with the lowest performance is CodeLlama-34B-Instruct, mainly due to its inability to generate correct JSON output and the significant hallucinations it produces.

While the highest EX score on the Spider dataset is 0.85, even models fine-tuned specifically for SpiderS only reach a maximum EX score of 0.731. This suggests that there is still room for improvement on SpiderS. More importantly, our experiments show that many models struggle with robustness when generating multiple SQL queries, which we discuss in more detail in Section 6.3.

Comparison of Small and Full Datasets. When comparing the full test set to its subset, we find that the results are very similar, with the small dataset slightly outperforming the larger one. This shows that our results on the small dataset are reliable.

Model Size and Performance. There is a clear positive correlation between model size and performance. For example, in the Llama-3.1 series, the 8B-parameter achieves an EX score of 0.411 on the SpiderS_{small}, while the 405B-parameter model achieves a score of 0.534. A similar trend is seen in the Qwen2.5 series, where the 32B model outperforms the 7B model by 0.059 in EX score on the SpiderS_{small}.

Effectiveness of Fine-Tuning. Comparing model performance before and after fine-tuning, we find that domain-specific fine-tuning leads to significant improvements. For example, CodeLlama-7B’s EX score on the SpiderS increased from 0.082 to 0.728, a 788% improvement.

6.3 Robustness Analysis

Almost all models show a drop in performance when moving from Spider to SpiderM or from SpiderM to SpiderS. The performance decreases for these models are shown in Figures 3 (d) and (e). The only exception is CHESS, which improved from SpiderM to SpiderS, likely because its performance had already dropped considerably from Spider to SpiderM.

This led to an interesting finding: simply adding the words ‘or multi’ to the prompt caused a sharp decline in the model’s ability to generate correct SQL, indicating that the model wasn’t well-prepared for handling multiple SQL queries. Among these models, Llama-3.1-8B-Instruct and o3-mini were the most robust, with their EX scores dropping by only 0.007 and 0.015, respectively. In contrast, CHESS and Mistral-Large-Instruct-2411 experienced the largest performance drops, with declines of 0.292 and 0.217, respectively.

6.4 Challenges of Text-to-SQL Models in Supporting Multi-SQL Generation

Many text-to-SQL models struggle to accommodate Text-to-Multi-SQL, which is why Table 2 only reports the results of two text-to-SQL models. Firstly, traditional sequence-to-sequence models based on grammar decoders are incapable of generating multiple SQL statements. Secondly, even modern models built on LLMs often lack support for multi SQL in their design. For instance, the MAC-SQL (Wang et al., 2024) model breaks down the SQL generation task into a process of generating multiple SQL clauses, which are then combined to form a complete SQL statement. However, this generation process is not well-suited to the requirements of producing multiple SQL statements.

6.5 Analysis on DeepSeek-R1-Distill Models

The performance of the DeepSeek-R1-Distill models was relatively poor, so we conducted a detailed analysis of the generated data to identify the causes. First, we set the maximum generation token for these models to 1024, but around 17% of the examples exceeded this limit, which prevented the tasks from being completed. One reason for the token shortage is that the model repeatedly included the prompt’s requirements in its output.

We randomly selected 30 examples where the task was completed successfully and analyzed the model’s reasoning process. We found that only 11

of these examples were correct. The model made several mistakes, such as generating hallucinations, frequently using non-existent database columns, or selecting incorrect columns. It also seemed to lose track of the original schema structure as the reasoning progressed. Based on this analysis, we conclude that the Distill models used in this experiment are not well-suited for text-to-SQL tasks. Additionally, the influence of prior knowledge on the Distill models led to suboptimal performance, even after fine-tuning, compared to traditional models.

6.6 Analysis on Question Types

To explore the question types (see Appendix A) where the models perform poorly, and to identify the reasons for their lower performance, we selected three representative models and analyzed their execution accuracy across all question types. The results are shown in Figure 4. Overall, there are 14 question types in total, with type 14 having 4 subclasses to represent all patterns requiring 3 SQLs. This distinction was made to easily differentiate them from patterns requiring only 2 SQLs.

The fine-tuned Qwen2.5-Coder-7B-Instruct shows improvements across all question types. Type 12 sees marked improvement after fine-tuning, as the questions are actually quite simple; the previous low performance was due to the model’s lack of training. XiYanSQL-QwenCoder-32B, while strong overall, struggles with multi-SQL tasks (types 14.1, 14.2, and 14.3), likely due to its prior single-SQL training. All models struggle with type 14.1, which is our most complex question type: combining one original question with one generated question that requires two SQL queries. Type 4 is also difficult for models, as it lacks common patterns and requires precise handling of each SQL.

7 Conclusion

We introduce the Text-to-Multi-SQL paradigm and the SpiderS dataset, shedding light on key limitations in current Text-to-SQL models. Our experiments reveal that even state-of-the-art models like GPT-4o and XiYanSQL experience significant performance drops when tasked with multi-SQL generation. A factor is the models’ sensitivity to prompt changes. To address these issues, future work should focus on improving model robustness and enhancing multi-query coordination to make data more accessible to non-technical users.

Limitations

Our work focuses solely on Text-to-SQL and Text-to-Multi-SQL, but even text-to-Multi-SQL cannot meet all user needs. As in the example in TAG (Biswal et al., 2024), users want to summarize the content of tables, and relying purely on SQL is currently difficult to achieve. Due to limitations in computing resources, we have not fine-tuned larger models (such as the 70B). At the moment, it is unclear how much better these models would perform after being fine-tuned.

References

Shanza Abbas, Muhammad Umair Khan, Scott Uk-Jin Lee, Asad Abbas, and Ali Kashif Bashir. 2022. [A review of nldb with deep learning: Findings, challenges and open issues](#). *IEEE Access*, 10:14927–14945.

Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.

Asim Biswal, Liana Patel, Siddarth Jha, Amog Kamsetty, Shu Liu, Joseph E. Gonzalez, Carlos Guestrin, and Matei Zaharia. 2024. [Text2sql is not enough: Unifying ai and databases with tag](#). *Preprint*, arXiv:2408.14717.

DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jiashi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaojun Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiushi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng Zhou, Shuting Pan, S. S. Li, Shuang Zhou, Shaoqing Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun, T. Wang, Wangding Zeng, Wanjia Zhao, Wen Liu,

Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin, Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxiang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang, Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yuan Ou, Yuduan Wang, Yue Gong, Yuheng Zou, Yujia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu, Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu, Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu Zhang, and Zhen Zhang. 2025. [Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning](#). *Preprint*, arXiv:2501.12948.

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.

Jacob Devlin. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*.

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](#). *Preprint*, arXiv:2308.15363.

Yingqi Gao, Yifu Liu, Xiaoxia Li, Xiaorong Shi, Yin Zhu, Yiming Wang, Shiqi Li, Wei Li, Yuntao Hong, Zhiling Luo, Jinyang Gao, Liyu Mou, and Yu Li. 2024a. [A preview of xian-sql: A multi-generator ensemble framework for text-to-sql](#). *arXiv preprint arXiv:2411.08599*.

Yingqi Gao, Yifu Liu, Xiaoxia Li, Xiaorong Shi, Yin Zhu, Yiming Wang, Shiqi Li, Wei Li, Yuntao Hong, Zhiling Luo, et al. 2024b. [Xian-sql: A multi-generator ensemble framework for text-to-sql](#). *arXiv preprint arXiv:2411.08599*.

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,

736	Charlotte Caucheteux, Chaya Nayak, Chloe Bi,	Zacharie Delpierre Coudert, Zheng Yan, Zhengxing	800
737	Chris Marra, Chris McConnell, Christian Keller,	Chen, Zoe Papakipos, Aaditya Singh, Aayushi Sri-	801
738	Christophe Touret, Chunyang Wu, Corinne Wong,	vastava, Abha Jain, Adam Kelsey, Adam Shajnfeld,	802
739	Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Al-	Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand,	803
740	lonsius, Daniel Song, Danielle Pintz, Danny Livshits,	Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei	804
741	Danny Wyatt, David Esiobu, Dhruv Choudhary,	Baevski, Allie Feinstein, Amanda Kallet, Amit San-	805
742	Dhruv Mahajan, Diego Garcia-Olano, Diego Perino,	gani, Amos Teo, Anam Yunus, Andrei Lupu, An-	806
743	Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy,	dres Alvarado, Andrew Caples, Andrew Gu, Andrew	807
744	Elina Lobanova, Emily Dinan, Eric Michael Smith,	Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchan-	808
745	Filip Radenovic, Francisco Guzmán, Frank Zhang,	dani, Annie Dong, Annie Franco, Anuj Goyal, Aparajita	809
746	Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis An-	jita Saraf, Arkabandhu Chowdhury, Ashley Gabriel,	810
747	derson, Govind Thattai, Graeme Nail, Gregoire Mi-	Ashwin Bharambe, Assaf Eisenman, Azadeh Yaz-	811
748	alon, Guan Pang, Guillem Cucurell, Hailey Nguyen,	dan, Beau James, Ben Maurer, Benjamin Leonhardi,	812
749	Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan	Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi	813
750	Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Is-	Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Han-	814
751	han Misra, Ivan Evtimov, Jack Zhang, Jade Copet,	cock, Bram Wasti, Brandon Spence, Brani Stojkovic,	815
752	Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park,	Brian Gamido, Britt Montalvo, Carl Parker, Carly	816
753	Jay Mahadeokar, Jeet Shah, Jelmer van der Linde,	Burton, Catalina Mejia, Ce Liu, Changhan Wang,	817
754	Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu,	Changkyu Kim, Chao Zhou, Chester Hu, Ching-	818
755	Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang,	Hsiang Chu, Chris Cai, Chris Tindal, Christoph Fe-	819
756	Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park,	ichtenhofer, Cynthia Gao, Damon Civin, Dana Beaty,	820
757	Joseph Rocca, Joshua Johnstun, Joshua Saxe, Jun-	Daniel Kreymer, Daniel Li, David Adkins, David	821
758	teng Jia, Kalyan Vasuden Alwala, Karthik Prasad,	Xu, Davide Testuggine, Delia David, Devi Parikh,	822
759	Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth	Diana Liskovich, Didem Foss, Dingkan Wang, Duc	823
760	Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer,	Le, Dustin Holland, Edward Dowling, Eissa Jamil,	824
761	Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal	Elaine Montgomery, Eleonora Presani, Emily Hahn,	825
762	Lakhotia, Lauren Rantala-Yearly, Laurens van der	Emily Wood, Eric-Tuan Le, Erik Brinkman, Este-	826
763	Maaten, Lawrence Chen, Liang Tan, Liz Jenkins,	ban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun,	827
764	Louis Martin, Lovish Madaan, Lubo Malo, Lukas	Felix Kreuk, Feng Tian, Filippas Kokkinos, Firat	828
765	Blecher, Lukas Landzaat, Luke de Oliveira, Madeline	Ozgenel, Francesco Caggioni, Frank Kanayet, Frank	829
766	Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar	Seide, Gabriela Medina Florez, Gabriella Schwarz,	830
767	Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew	Gada Badeer, Georgia Sweet, Gil Halpern, Grant	831
768	Oldham, Mathieu Rita, Maya Pavlova, Melanie Kam-	Herman, Grigory Sizov, Guangyi, Zhang, Guna	832
769	badur, Mike Lewis, Min Si, Mitesh Kumar Singh,	Lakshminarayanan, Hakan Inan, Hamid Shojanaz-	833
770	Mona Hassan, Naman Goyal, Narjes Torabi, Niko-	eri, Han Zou, Hannah Wang, Hanwen Zha, Haroun	834
771	lay Bashlykov, Nikolay Bogoychev, Niladri Chatterji,	Habeeb, Harrison Rudolph, Helen Suk, Henry As-	835
772	Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick	pegren, Hunter Goldman, Hongyuan Zhan, Ibrahim	836
773	Alrassy, Pengchuan Zhang, Pengwei Li, Petar Va-	Damlaj, Igor Molybog, Igor Tufanov, Ilias Leontiadis,	837
774	sic, Peter Weng, Prajjwal Bhargava, Pratik Dubal,	Irina-Elena Veliche, Itai Gat, Jake Weissman, James	838
775	Praveen Krishnan, Punit Singh Koura, Puxin Xu,	Geboski, James Kohli, Janice Lam, Japhet Asher,	839
776	Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj	Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jen-	840
777	Ganapathy, Ramon Calderer, Ricardo Silveira Cabral,	nifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy	841
778	Robert Stojnic, Roberta Raileanu, Rohan Maheswari,	Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe	842
779	Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ron-	Cummings, Jon Carvill, Jon Shepard, Jonathan Mc-	843
780	nie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan	Phie, Jonathan Torres, Josh Ginsburg, Junjie Wang,	844
781	Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sa-	Kai Wu, Kam Hou U, Karan Saxena, Kartikay Khan-	845
782	hana Chennabasappa, Sanjay Singh, Sean Bell, Seo-	delwal, Katayoun Zand, Kathy Matosich, Kaushik	846
783	hyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sha-	Veeraraghavan, Kelly Michelena, Keqian Li, Ki-	847
784	ran Narang, Sharath Raparthy, Sheng Shen, Shengye	ran Jagadeesh, Kun Huang, Kunal Chawla, Kyle	848
785	Wan, Shruti Bhosale, Shun Zhang, Simon Van-	Huang, Lailin Chen, Lakshya Garg, Lavender A,	849
786	denhende, Soumya Batra, Spencer Whitman, Sten	Leandro Silva, Lee Bell, Lei Zhang, Liangpeng	850
787	Sootla, Stephane Collot, Suchin Gururangan, Syd-	Guo, Licheng Yu, Liron Moshkovich, Luca Wehrst-	851
788	ney Borodinsky, Tamar Herman, Tara Fowler, Tarek	edt, Madian Khabsa, Manav Avalani, Manish Bhatt,	852
789	Sheasha, Thomas Georgiou, Thomas Scialom, Tobias	Martynas Mankus, Matan Hasson, Matthew Lennie,	853
790	Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal	Matthias Reso, Maxim Groshev, Maxim Naumov,	854
791	Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh	Maya Lathi, Meghan Keneally, Miao Liu, Michael L.	855
792	Ramanathan, Viktor Kerkez, Vincent Gouget, Vir-	Seltzer, Michal Valko, Michelle Restrepo, Mihir Pa-	856
793	ginie Do, Vish Vogeti, Vitor Albiero, Vladan Petro-	tel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark,	857
794	vic, Weiwei Chu, Wenhan Xiong, Wenxin Fu, Whit-	Mike Macey, Mike Wang, Miquel Jubert Hermoso,	858
795	ney Meers, Xavier Martinet, Xiaodong Wang, Xi-	Mo Metanat, Mohammad Rastegari, Munish Bansal,	859
796	aofang Wang, Xiaoqing Ellen Tan, Xide Xia, Xin-	Nandhini Santhanam, Natascha Parks, Natasha	860
797	feng Xie, Xuchao Jia, Xuwei Wang, Yaelle Gold-	White, Navyata Bawa, Nayan Singhal, Nick Egebo,	861
798	schlag, Yashesh Gaur, Yasmine Babaei, Yi Wen,	Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich	862
799	Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao,	Laptev, Ning Dong, Norman Cheng, Oleg Chernoguz,	863

864	Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin	Peter Lawrence. 2024. Text-to-graph via llm: pre-	924
865	Kent, Parth Parekh, Paul Saab, Pavan Balaji, Pe-	training, prompting, or tuning.	925
866	dro Rittner, Philip Bontrager, Pierre Roux, Piotr		
867	Dollar, Polina Zvyagina, Prashant Ratanchandani,	Haoyang Li, Jing Zhang, Cuiping Li, and Hong Chen.	926
868	Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel	2023. Resdsq: Decoupling schema linking and	927
869	Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu	skeleton parsing for text-to-sql. In <i>Proceedings of</i>	928
870	Nayani, Rahul Mitra, Rangaprabhu Parthasarathy,	<i>the AAAI Conference on Artificial Intelligence</i> , vol-	929
871	Raymond Li, Rebekkah Hogan, Robin Battey, Rocky	ume 37, pages 13067–13075.	930
872	Wang, Russ Howes, Ruty Rinott, Sachin Mehta,		
873	Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara	Jinyang Li, Binyuan Hui, Ge Qu, Jiaxi Yang, Binhua	931
874	Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov,	Li, Bowen Li, Bailin Wang, Bowen Qin, Ruiying	932
875	Satadru Pan, Saurabh Mahajan, Saurabh Verma,	Geng, Nan Huo, et al. 2024. Can llm already serve	933
876	Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lind-	as a database interface? a big bench for large-scale	934
877	say, Shaun Lindsay, Sheng Feng, Shenghao Lin,	database grounded text-to-sqls. <i>Advances in Neural</i>	935
878	Shengxin Cindy Zha, Shishir Patil, Shiva Shankar,	<i>Information Processing Systems</i> , 36.	936
879	Shuqiang Zhang, Shuqiang Zhang, Sinong Wang,		
880	Sneha Agarwal, Soji Sajuyigbe, Soumith Chintala,	Haoran Luo, Zichen Tang, Shiyao Peng, Yikai Guo,	937
881	Stephanie Max, Stephen Chen, Steve Kehoe, Steve	Wentai Zhang, Chenghao Ma, Guanting Dong, Meina	938
882	Satterfield, Sudarshan Govindaprasad, Sumit Gupta,	Song, Wei Lin, Yifan Zhu, et al. 2023. Chatkbqa: A	939
883	Summer Deng, Sungmin Cho, Sunny Virk, Suraj	generate-then-retrieve framework for knowledge base	940
884	Subramanian, Sy Choudhury, Sydney Goldman, Tal	question answering with fine-tuned large language	941
885	Remez, Tamar Glaser, Tamara Best, Thilo Koehler,	models. <i>arXiv preprint arXiv:2310.08975</i> .	942
886	Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim		
887	Matthews, Timothy Chou, Tzook Shaked, Varun	Rui Ma, Kai Zhang, Zhenying He, Yinan Jing, X Sean	943
888	Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai	Wang, and Zhenqiang Chen. 2025. Chase: A	944
889	Mohan, Vinay Satish Kumar, Vishal Mangla, Vlad	native relational database for hybrid queries on	945
890	Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu,	structured and unstructured data. <i>arXiv preprint</i>	946
891	Vladimir Ivanov, Wei Li, Wenchen Wang, Wen-	<i>arXiv:2501.05006</i> .	947
892	wen Jiang, Wes Bouaziz, Will Constable, Xiaocheng		
893	Tang, Xiaojian Wu, Xiaolan Wang, Xilun Wu, Xinbo	OpenAI, :, Aaron Hurst, Adam Lerer, Adam P. Goucher,	948
894	Gao, Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia,	Adam Perelman, Aditya Ramesh, Aidan Clark,	949
895	Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi,	AJ Ostrow, Akila Welihinda, Alan Hayes, Alec	950
896	Youngjin Nam, Yu, Wang, Yu Zhao, Yuchen Hao,	Radford, Aleksander Mądry, Alex Baker-Whitcomb,	951
897	Yundi Qian, Yunlu Li, Yuzi He, Zach Rait, Zachary	Alex Beutel, Alex Borzunov, Alex Carney, Alex	952
898	DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang,	Chow, Alex Kirillov, Alex Nichol, Alex Paino, Alex	953
899	Zhiwei Zhao, and Zhiyu Ma. 2024. The llama 3 herd	Renzin, Alex Tachard Passos, Alexander Kirillov,	954
900	of models . <i>Preprint</i> , arXiv:2407.21783.	Alexi Christakis, Alexis Conneau, Ali Kamali, Allan	955
		Jabri, Allison Moyer, Allison Tam, Amadou Crookes,	956
901	Zijin Hong, Zheng Yuan, Qinggang Zhang, Hao Chen,	Amin Tootoochian, Amin Tootoonchian, Ananya	957
902	Junnan Dong, Feiran Huang, and Xiao Huang. 2024.	Kumar, Andrea Vallone, Andrej Karpathy, Andrew	958
903	Next-generation database interfaces: A survey of llm-	Braunstein, Andrew Cann, Andrew Codispoli, An-	959
904	based text-to-sql . <i>Preprint</i> , arXiv:2406.08426.	drew Galu, Andrew Kondrich, Andrew Tulloch, An-	960
		drey Mishchenko, Angela Baek, Angela Jiang, An-	961
905	Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan	toine Pelisse, Antonia Woodford, Anuj Gosalia, Arka	962
906	Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and	Dhar, Ashley Pantuliano, Avi Nayak, Avital Oliver,	963
907	Weizhu Chen. 2021. Lora: Low-rank adaptation of	Barret Zoph, Behrooz Ghorbani, Ben Leimberger,	964
908	large language models . <i>Preprint</i> , arXiv:2106.09685.	Ben Rossen, Ben Sokolowsky, Ben Wang, Benjamin	965
		Zweig, Beth Hoover, Blake Samic, Bob McGrew,	966
909	Radu Cristian Alexandru Iacob, Florin Brad,	Bobby Spero, Bogo Gertler, Bowen Cheng, Brad	967
910	Elena-Simona Apostol, Ciprian-Octavian Tru-	Lightcap, Brandon Walkin, Brendan Quinn, Brian	968
911	ică, Ionel Alexandru Hosu, and Traian Rebedea.	Guarraci, Brian Hsu, Bright Kellogg, Brydon East-	969
912	2020. Neural approaches for natural language	man, Camillo Lugaresi, Carroll Wainwright, Cary	970
913	interfaces to databases: A survey . In <i>Proceedings of</i>	Bassin, Cary Hudson, Casey Chu, Chad Nelson,	971
914	<i>the 28th International Conference on Computational</i>	Chak Li, Chan Jun Shern, Channing Conger, Char-	972
915	<i>Linguistics</i> , pages 381–395, Barcelona, Spain	lotte Barette, Chelsea Voss, Chen Ding, Cheng Lu,	973
916	(Online). International Committee on Computational	Chong Zhang, Chris Beaumont, Chris Hallacy, Chris	974
917	Linguistics.	Koch, Christian Gibson, Christina Kim, Christine	975
		Choi, Christine McLeavey, Christopher Hesse, Clau-	976
918	George Katsogiannis-Meimarakis and Georgia Koutrika.	dia Fischer, Clemens Winter, Coley Czarnecki, Colin	977
919	2023a. A survey on deep learning approaches for	Jarvis, Colin Wei, Constantin Koumouzelis, Dane	978
920	text-to-sql . <i>The VLDB Journal</i> , 32(4):905–936.	Sherburn, Daniel Kappler, Daniel Levin, Daniel Levy,	979
		David Carr, David Farhi, David Mely, David Robin-	980
921	George Katsogiannis-Meimarakis and Georgia Koutrika.	son, David Sasaki, Denny Jin, Dev Valladares, Dim-	981
922	2023b. A survey on deep learning approaches for	itris Tsipras, Doug Li, Duc Phong Nguyen, Duncan	982
923	text-to-sql . <i>The VLDB Journal</i> , 32(4):905–936.		

983	Findlay, Edede Oiwoh, Edmund Wong, Ehsan As-	jan Troll, Randall Lin, Rapha Gontijo Lopes, Raul	1047
984	dar, Elizabeth Proehl, Elizabeth Yang, Eric Antonow,	Puri, Reah Miyara, Reimar Leike, Renaud Gaubert,	1048
985	Eric Kramer, Eric Peterson, Eric Sigler, Eric Wal-	Reza Zamani, Ricky Wang, Rob Donnelly, Rob	1049
986	lace, Eugene Brevdo, Evan Mays, Farzad Khorasani,	Honsby, Rocky Smith, Rohan Sahai, Rohit Ramchan-	1050
987	Felipe Petroski Such, Filippo Raso, Francis Zhang,	dani, Romain Huet, Rory Carmichael, Rowan Zellers,	1051
988	Fred von Lohmann, Freddie Sulit, Gabriel Goh,	Roy Chen, Ruby Chen, Ruslan Nigmatullin, Ryan	1052
989	Gene Oden, Geoff Salmon, Giulio Starace, Greg	Cheu, Saachi Jain, Sam Altman, Sam Schoenholz,	1053
990	Brockman, Hadi Salman, Haiming Bao, Haitang	Sam Toizer, Samuel Miserendino, Sandhini Agar-	1054
991	Hu, Hannah Wong, Haoyu Wang, Heather Schmidt,	wal, Sara Culver, Scott Ethersmith, Scott Gray, Sean	1055
992	Heather Whitney, Heewoo Jun, Hendrik Kirchner,	Grove, Sean Metzger, Shamez Hermani, Shantanu	1056
993	Henrique Ponde de Oliveira Pinto, Hongyu Ren,	Jain, Shengjia Zhao, Sherwin Wu, Shino Jomoto, Shi-	1057
994	Huiwen Chang, Hyung Won Chung, Ian Kivlichan,	rong Wu, Shuaiqi, Xia, Sonia Phene, Spencer Papay,	1058
995	Ian O’Connell, Ian O’Connell, Ian Osband, Ian Sil-	Srinivas Narayanan, Steve Coffey, Steve Lee, Stew-	1059
996	ber, Ian Sohl, Ibrahim Okuyucu, Ikai Lan, Ilya	art Hall, Suchir Balaji, Tal Broda, Tal Stramer, Tao	1060
997	Kostrikov, Ilya Sutskever, Ingmar Kanitscheider,	Xu, Tarun Gogineni, Taya Christianson, Ted Sanders,	1061
998	Ishaan Gulrajani, Jacob Coxon, Jacob Menick, Jakub	Tejal Patwardhan, Thomas Cunningham, Thomas	1062
999	Pachocki, James Aung, James Betker, James Crooks,	Degry, Thomas Dimson, Thomas Raoux, Thomas	1063
1000	James Lennon, Jamie Kiros, Jan Leike, Jane Park,	Shadwell, Tianhao Zheng, Todd Underwood, Todor	1064
1001	Jason Kwon, Jason Phang, Jason Teplitz, Jason	Markov, Toki Sherbakov, Tom Rubin, Tom Stasi,	1065
1002	Wei, Jason Wolfe, Jay Chen, Jeff Harris, Jenia Var-	Tomer Kaftan, Tristan Heywood, Troy Peterson, Tyce	1066
1003	avva, Jessica Gan Lee, Jessica Shieh, Ji Lin, Jiahui	Walters, Tyna Eloundou, Valerie Qi, Veit Moeller,	1067
1004	Yu, Jiayi Weng, Jie Tang, Jieqi Yu, Joanne Jang,	Vinnie Monaco, Vishal Kuo, Vlad Fomenko, Wayne	1068
1005	Joaquin Quinero Candela, Joe Beutler, Joe Lan-	Chang, Weiyi Zheng, Wenda Zhou, Wesam Manassra,	1069
1006	ders, Joel Parish, Johannes Heidecke, John Schul-	Will Sheu, Wojciech Zaremba, Yash Patil, Yilei Qian,	1070
1007	man, Jonathan Lachman, Jonathan McKay, Jonathan	Yongjik Kim, Youlong Cheng, Yu Zhang, Yuchen	1071
1008	Uesato, Jonathan Ward, Jong Wook Kim, Joost	He, Yuchen Zhang, Yujia Jin, Yunxing Dai, and	1072
1009	Huizinga, Jordan Sitkin, Jos Kraaijeveld, Josh Gross,	Yury Malkov. 2024. Gpt-4o system card . <i>Preprint</i> ,	1073
1010	Josh Kaplan, Josh Snyder, Joshua Achiam, Joy Jiao,	arXiv:2410.21276.	1074
1011	Joyce Lee, Juntang Zhuang, Justyn Harriman, Kai		
1012	Fricke, Kai Hayashi, Karan Singhal, Katy Shi, Kavin	Mohammadreza Pourreza, Hailong Li, Ruoxi Sun,	1075
1013	Karthik, Kayla Wood, Kendra Rimbach, Kenny Hsu,	Yeounoh Chung, Shayan Talaei, Gaurav Tarlok	1076
1014	Kenny Nguyen, Keren Gu-Lemberg, Kevin Button,	Kakkar, Yu Gan, Amin Saberi, Fatma Ozcan, and	1077
1015	Kevin Liu, Kiel Howe, Krithika Muthukumar, Kyle	Sercan O. Arik. 2024. Chase-sql: Multi-path reason-	1078
1016	Luther, Lama Ahmad, Larry Kai, Lauren Itow, Lau-	ing and preference optimized candidate selection in	1079
1017	ren Workman, Leher Pathak, Leo Chen, Li Jing, Lia	text-to-sql . <i>Preprint</i> , arXiv:2410.01943.	1080
1018	Guy, Liam Fedus, Liang Zhou, Lien Mamitsuka, Lil-		
1019	ian Weng, Lindsay McCallum, Lindsey Held, Long	Mohammadreza Pourreza and Davood Rafiei. 2023.	1081
1020	Ouyang, Louis Feuvrier, Lu Zhang, Lukas Kon-	Din-sql: Decomposed in-context learning of text-to-	1082
1021	draciuk, Lukasz Kaiser, Luke Hewitt, Luke Metz,	sql with self-correction . <i>Preprint</i> , arXiv:2304.11015.	1083
1022	Lyric Doshi, Mada Aflak, Maddie Simens, Madelaine		
1023	Boyd, Madeleine Thompson, Marat Dukhan, Mark	Jiexing Qi, Jingyao Tang, Ziwei He, Xiangpeng Wan,	1084
1024	Chen, Mark Gray, Mark Hudnall, Marvin Zhang,	Yu Cheng, Chenghu Zhou, Xinbing Wang, Quanshi	1085
1025	Marwan Aljubei, Mateusz Litwin, Matthew Zeng,	Zhang, and Zhouhan Lin. 2022. Rasat: Integrating	1086
1026	Max Johnson, Maya Shetty, Mayank Gupta, Meghan	relational structures into pretrained seq2seq model	1087
1027	Shah, Mehmet Yatbaz, Meng Jia Yang, Mengchao	for text-to-sql . <i>arXiv preprint arXiv:2205.06983</i> .	1088
1028	Zhong, Mia Glaese, Mianna Chen, Michael Jan-		
1029	ner, Michael Lampe, Michael Petrov, Michael Wu,	Bowen Qin, Binyuan Hui, Lihan Wang, Min Yang,	1089
1030	Michele Wang, Michelle Fradin, Michelle Pokrass,	Jinyang Li, Binhua Li, Ruiying Geng, Rongyu Cao,	1090
1031	Miguel Castro, Miguel Oom Temudo de Castro,	Jian Sun, Luo Si, Fei Huang, and Yongbin Li. 2022.	1091
1032	Mikhail Pavlov, Miles Brundage, Miles Wang, Mil-	A survey on text-to-sql parsing: Concepts, methods,	1092
1033	nal Khan, Mira Murati, Mo Bavarian, Molly Lin,	and future directions . <i>Preprint</i> , arXiv:2208.13629.	1093
1034	Murat Yesildal, Nacho Soto, Natalia Gimelshein, Na-		
1035	talie Cone, Natalie Staudacher, Natalie Summers,	Qwen, :, An Yang, Baosong Yang, Beichen Zhang,	1094
1036	Natan LaFontaine, Neil Chowdhury, Nick Ryder,	Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li,	1095
1037	Nick Stathas, Nick Turley, Nik Tezak, Niko Felix,	Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin,	1096
1038	Nithanth Kudige, Nitish Keskar, Noah Deutsch, Noel	Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang,	1097
1039	Bundick, Nora Puckett, Ofir Nachum, Ola Okelola,	Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang,	1098
1040	Oleg Boiko, Oleg Murk, Oliver Jaffe, Olivia Watkins,	Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li,	1099
1041	Olivier Godement, Owen Campbell-Moore, Patrick	Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji	1100
1042	Chao, Paul McMillan, Pavel Belov, Peng Su, Pe-	Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang	1101
1043	ter Bak, Peter Bakkum, Peter Deng, Peter Dolan,	Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang	1102
1044	Peter Hoeschele, Peter Welinder, Phil Tillet, Philip	Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru	1103
1045	Pronin, Philippe Tillet, Prafulla Dhariwal, Qiming	Zhang, and Zihan Qiu. 2025. Qwen2.5 technical	1104
1046	Yuan, Rachel Dias, Rachel Lim, Rahul Arora, Ra-	report . <i>Preprint</i> , arXiv:2412.15115.	1105

1106	Cedric Renggli, Ihab F Ilyas, and Theodoros Rekatsinas. 2025. Fundamental challenges in evaluating text2sql solutions and detecting their limitations. <i>arXiv preprint arXiv:2501.18197</i> .	1161
1107		1162
1108		1163
1109		1164
1110	Baptiste Rozière, Jonas Gehring, Fabian Gloeckle, Sten Sootla, Itai Gat, Xiaoqing Ellen Tan, Yossi Adi, Jingyu Liu, Romain Sauvestre, Tal Remez, Jérémy Rapin, Artyom Kozhevnikov, Ivan Evtimov, Joanna Bitton, Manish Bhatt, Cristian Canton Ferrer, Aaron Grattafiori, Wenhan Xiong, Alexandre Défossez, Jade Copet, Faisal Azhar, Hugo Touvron, Louis Martin, Nicolas Usunier, Thomas Scialom, and Gabriel Synnaeve. 2024. Code llama: Open foundation models for code . <i>Preprint</i> , arXiv:2308.12950.	1165
1111		1166
1112		1167
1113		
1114		
1115		
1116		
1117		
1118		
1119		
1120	Torsten Scholak, Nathan Schucher, and Dzmitry Bahdanau. 2021. Picard: Parsing incrementally for constrained auto-regressive decoding from language models. <i>arXiv preprint arXiv:2109.05093</i> .	1168
1121		1169
1122		1170
1123		
1124	Shayan Talaei, Mohammadreza Pourreza, Yu-Chen Chang, Azalia Mirhoseini, and Amin Saberi. 2024. Chess: Contextual harnessing for efficient sql synthesis. <i>arXiv preprint arXiv:2405.16755</i> .	
1125		
1126		
1127		
1128	A Vaswani. 2017. Attention is all you need. <i>Advances in Neural Information Processing Systems</i> .	
1129		
1130	Bing Wang, Changyu Ren, Jian Yang, Xinnian Liang, Jiaqi Bai, Linzheng Chai, Zhao Yan, Qian-Wen Zhang, Di Yin, Xing Sun, and Zhoujun Li. 2024. Mac-sql: A multi-agent collaborative framework for text-to-sql . <i>Preprint</i> , arXiv:2312.11242.	
1131		
1132		
1133		
1134		
1135	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. In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing</i> , pages 3911–3921.	1168
1136		1169
1137		
1138		
1139		
1140		
1141		
1142		
1143	Weixu Zhang, Yifei Wang, Yuanfeng Song, Victor Junqiu Wei, Yuxing Tian, Yiyan Qi, Jonathan H. Chan, Raymond Chi-Wing Wong, and Haiqin Yang. 2024. Natural language interfaces for tabular data querying and visualization: A survey . <i>Preprint</i> , arXiv:2310.17894.	
1144		
1145		
1146		
1147		
1148		
1149	Yaowei Zheng, Richong Zhang, Junhao Zhang, Yanhan Ye, Zheyang Luo, Zhangchi Feng, and Yongqiang Ma. 2024. Llamafactory: Unified efficient fine-tuning of 100+ language models . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 3: System Demonstrations)</i> , Bangkok, Thailand. Association for Computational Linguistics.	
1150		
1151		
1152		
1153		
1154		
1155		
1156		
1157	Victor Zhong, Caiming Xiong, and Richard Socher. 2017. Seq2sql: Generating structured queries from natural language using reinforcement learning. <i>arXiv preprint arXiv:1709.00103</i> .	
1158		
1159		
1160		
	Yuhang Zhou, Yu He, Siyu Tian, Yuchen Ni, Zhangyue Yin, Xiang Liu, Chuanjun Ji, Sen Liu, Xipeng Qiu, Guangnan Ye, et al. 2024. r3-nl2gql: A model coordination and knowledge graph alignment approach for nl2gql. In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 13679–13692.	1168
		1169
		1170
	Xiaohu Zhu, Qian Li, Lizhen Cui, and Yongkang Liu. 2024. Large language model enhanced text-to-sql generation: A survey . <i>Preprint</i> , arXiv:2410.06011.	
	A Data Generation Patterns	1171
	To enhance the diversity and complexity of text-to-SQL datasets, we propose a systematic approach to data generation. Our methodology is based on modifying existing text-to-SQL pairs, ensuring that a single natural language query can correspond to multiple SQL queries while preserving the semantics of the original question. This approach allows us to construct a dataset with richer linguistic and structural variations, facilitating robust model training and evaluation.	1172
		1173
		1174
		1175
		1176
		1177
		1178
		1179
		1180
		1181
	We adopt multiple patterns to transform single-query data into multi-query pairs. These patterns fall into 14 different categories, each designed to introduce different types of complexity while maintaining the integrity of the original intent. Below is a detailed introduction of our generation patterns.	1182
		1183
		1184
		1185
		1186
		1187
	1. Reverse Selection in One-to-Many Relationships	1188
		1189
	To introduce compositional complexity into text-to-SQL pairs, we employ reverse selection on one-to-many relationships. This approach ensures that a given natural language query requires multiple SQL statements to retrieve the requested information, as opposed to a single SQL query.	1190
		1191
		1192
		1193
		1194
		1195
		1196
	In relational databases, a one-to-many relationship exists when an entity in one table (e.g., departments) is associated with multiple entities in another table (e.g., teachers). The department table serves as the parent (one), while the teacher table acts as the child (many), with the foreign key in the teacher table referencing the primary key in the department table.	1197
		1198
		1199
		1200
		1201
		1202
		1203
		1204
		1205
	Consider the following example, where the original query is:	1206
		1207
	<i>Find the names of teachers who are older than 40.</i>	1208
		1209

1210
1211
1212
1213
1214
1215
1216
1217
1218
1219
1220
1221
1222
1223
1224
1225
1226
1227
1228
1229
1230
1231
1232
1233
1234
1235
1236
1237
1238
1239
1240
1241
1242
1243
1244
1245
1246
1247
1248
1249
1250
1251
1252
1253
1254
1255
1256
1257

This query can be answered using a single SQL statement:

```
SELECT name FROM teacher WHERE age > 40;
```

To enforce multi-query execution, we introduce a reverse selection constraint by modifying the query as follows:

Find the names of teachers who are older than 40 and the departments of the remaining teachers.

Here, the teacher name and the department name of other teachers are semantically incompatible in a single SQL query, necessitating at least two queries:

```
SELECT name FROM teacher WHERE age > 40;
```

```
SELECT department_name FROM teacher WHERE age <= 40;
```

This transformation ensures that a single query is insufficient to retrieve all requested information, thereby enriching the complexity of text-to-Multiple-SQLs pairs.

2. **Reverse Selection in One-to-One Relationships**

Unlike the one-to-many relationships discussed previously, a one-to-one relationship occurs when each record in a table corresponds to a single record in another table. In this case, reverse selection cannot involve direct inclusion of attributes from the same table, as such queries can always be answered using a single SQL statement. Therefore, to enforce compositional complexity, reverse selection in one-to-one relationships must ensure that the two queries retrieve different types of information.

A critical constraint of one-to-one reverse selection is that the secondary query must retrieve different attributes. If the modified query asks for the same attribute from different records, it can always be rewritten as a single SQL query, rendering the transformation invalid. So we need to ensure the queries ask for logically distinct attributes. Consider the following transformation:

Find the names of teachers who are older than 40 and retrieve the average age of the remaining teachers.

This results in the following SQL queries:

```
SELECT name FROM teacher WHERE age > 40;
```

```
SELECT AVG(age) FROM teacher WHERE age <= 40;
```

Here, the first query retrieves individual teacher names, while the second query retrieves an aggregate statistic, ensuring they cannot be merged into a single SQL statement.

3. **Reverse Selection in One-to-One Relationships Without Aggregation**

One-to-one reverse selection can be further constrained by eliminating aggregation functions, making the compositional SQL queries more structurally distinct while still requiring multiple SQL queries to fully answer the question. This ensures that query decomposition is necessary but does not rely on aggregate statistics to enforce multi-query generation.

4. **Distinguishing Separate Entities for Independent Queries**

In the process of constructing diverse SQL queries, one effective approach is to generate questions that target distinct entities within the database. At first glance, it may seem intuitive that retrieving different types of information always necessitates multiple SQL queries. However, this is not necessarily the case. When two entities are inherently linked through a foreign key relationship, it is often possible to retrieve both pieces of information using a single query. For instance, if the original query asks for the names of all teachers, it can be extended to also request their department names without requiring an additional SQL statement. Since teachers are linked to departments via a foreign key, a simple join operation suffices to bring both sets of information together within a single query.

However, to ensure that two distinct queries are required rather than one combined query, the entities being retrieved must not only belong to different tables but also be subject to independent constraints. Consider the case where we modify the original question from "Find the names of teachers older than 40" to "Find the names of teachers older than 40 and the names of departments with more than 40 members." This slight modification forces the

1258
1259
1260
1261
1262
1263
1264
1265
1266
1267
1268
1269
1270
1271
1272
1273
1274
1275
1276
1277
1278
1279
1280
1281
1282
1283
1284
1285
1286
1287
1288
1289
1290
1291
1292
1293
1294
1295
1296
1297
1298
1299
1300
1301
1302
1303
1304
1305
1306

1307	retrieval of two independent sets of information—one about teachers filtered by age, and	remain separate rather than being merged into	1355
1308	another about departments filtered by the number of members. Since these conditions apply	a single SQL statement.	1356
1309	to separate tables and do not share a common	The process of generating such queries relies	1357
1310	join path that would allow them to be merged	on selecting an appropriate example from existing data, ensuring that one query retains a	1358
1311	into a single SQL query, it becomes necessary	condition while the other does not. The system systematically scans the dataset to identify	1359
1312	to generate two distinct SQL statements.	candidate queries that fit this asymmetric	1360
1313		pattern, filtering out those that contain multiple	1361
1314		constraints or already require multiple	1362
1315	<i>Find the names of teachers who are older than</i>	SQL statements. Once an appropriate pair	1363
1316	<i>40 and the names of departments with more</i>	of queries is selected, the prompt is carefully	1364
1317	<i>than 40 people.</i>	constructed to emphasize the independence of	1365
1318	This question introduces conditions on different	the two queries while maintaining clarity in	1366
1319	attributes from separate tables, making it	natural language.	1367
1320	impossible to construct a single SQL query		1368
1321	that answers both parts. Instead, two independent		1369
1322	queries must be generated:		
1323	<i>SELECT name FROM teacher WHERE age ></i>	6. Blending Aggregates with Individual Attributes: Merging Statistical Insights with	1370
1324	<i>40;</i>	Specifics	1371
1325	<i>SELECT name FROM department WHERE</i>	In natural language database queries, there is	1372
1326	<i>num_people > 40;</i>	often a need to balance between querying individual	1373
1327		records and summarizing trends. A	1374
1328	5. Asymmetric Query Expansion: One Condition, One Open Query	common way to expand a query is by introducing	1375
1329	In contrast to the previous method, which ensured	an aggregation function alongside a standard	1376
1330	that two distinct objects were queried with their own	selection. This method ensures that both	1377
1331	independent constraints, this method introduces an	specific details and overarching insights are	1378
1332	asymmetry: one of the SQL queries retains a specific	retrieved simultaneously, making the query	1379
1333	condition while the other remains unrestricted. This	richer and more informative.	1380
1334	subtle shift in design changes the way queries	For example, consider the original query:	1381
1335	are structured and how they interact with the	<i>Find the names of teachers who are older than</i>	1382
1336	database schema. Instead of ensuring that	<i>40.</i>	1383
1337	both queries are filtered by independent	This query retrieves a list of individual teachers	1384
1338	constraints, one query maintains a filtering	who meet the condition, focusing on granular	1385
1339	condition while the other remains a general	details. However, by expanding the question	1386
1340	retrieval of all records from a related table.	to include an aggregation function, we	1387
1341		can introduce a statistical perspective. The	1388
1342	To illustrate, consider an original query asking	modified query becomes:	1389
1343	for "Find the names of all teachers." A	<i>Find the names of teachers who are older than</i>	1390
1344	straightforward extension would be to also	<i>40 and their average age.</i>	1391
1345	retrieve the names of all departments, a task	Here, in addition to retrieving individual	1392
1346	that could easily be achieved using a single	teacher names, the query also calculates the	1393
1347	SQL query through a simple join operation. However,	overall average age of teachers above 40.	1394
1348	if we modify this to "Find the names of	This transformation creates a dual perspective—	1395
1349	teachers who are older than 40 and all department	while one part of the query extracts direct	1396
1350	names," we introduce a structural asymmetry—	information from the database, the other	1397
1351	one query applies a condition (teachers	part summarizes the information using an	1398
1352	older than 40), while the other query	aggregation function.	1399
1353	retrieves all departments without filtering. This	To construct such queries correctly, it is	1400
1354	distinction ensures that the two queries must	crucial to ensure that both components operate	1401
			1402

1403	on the same filtered dataset. This prevents in-	<i>Tell me which states have at least 3 heads</i>	1449
1404	consistencies, such as computing an average	<i>born there and which age groups have at least</i>	1450
1405	age across all teachers instead of just those	<i>3 heads?</i>	1451
1406	older than 40. The corresponding SQL query	The result is that we have to use 2 different	1452
1407	would look like this:	SQL queries to answer this question:	1453
1408	<i>SELECT name FROM teachers WHERE age</i>	<i>SELECT age FROM head GROUP BY age</i>	1454
1409	<i>> 40;</i>	<i>HAVING COUNT(*) >= 3;</i>	1455
1410	<i>SELECT AVG(age) FROM teachers WHERE</i>	<i>SELECT born_state FROM head GROUP BY</i>	1456
1411	<i>age > 40;</i>	<i>born_state HAVING COUNT(*) >= 3;</i>	1457
1412	Alternatively, a single SQL statement can of-	8. Switching Between HAVING and LIMIT	1458
1413	ten achieve the same result using a GROUP	for Diverse Filtering	1459
1414	BY clause:	This method modifies SQL queries by ensur-	1460
1415	<i>SELECT name, (SELECT AVG(age) FROM</i>	ing that one query filters aggregated results us-	1461
1416	<i>teachers WHERE age > 40) AS avg_age</i>	ing HAVING, while the other applies ORDER	1462
1417	<i>FROM teachers WHERE age > 40;</i>	BY and LIMIT. Unlike the previous approach,	1463
1418	Since this data generation method can also be	which changed the GROUP BY column to	1464
1419	completed with a SQL query, we introduced	introduce variation, this method focuses on	1465
1420	the sql_0 statement to store all queries inte-	switching between filtering patterns to gener-	1466
1421	grated into one SQL.	ate distinct yet complementary queries.	1467
1422	7. Querying Different Grouping Dimensions	If the original query includes LIMIT, the sec-	1468
1423	Grouping data is a fundamental aspect of	ond query must replace it with a HAVING	1469
1424	SQL queries, allowing us to derive meaning-	clause to introduce a different filtering method.	1470
1425	ful insights by categorizing and summarizing	Conversely, if the original query already uses	1471
1426	records. While most queries focus on a sin-	HAVING, the new query must avoid it and in-	1472
1427	gle way to group data, an alternative approach	stead adopt an ORDER BY and LIMIT struc-	1473
1428	introduces multiple GROUP BY clauses, ef-	ture. This ensures that the two queries cannot	1474
1429	fectively broadening the scope of the query.	be merged into one, as they use fundamentally	1475
1430	By shifting the lens through which data is ag-	different ways of selecting data.	1476
1431	gregated, we can uncover deeper relationships	For example, if the original query retrieves	1477
1432	within the dataset that a single GROUP BY	the department with the most teachers using	1478
1433	might miss.	ORDER BY and LIMIT:	1479
1434	This method is especially useful when deal-	<i>SELECT department FROM teachers GROUP</i>	1480
1435	ing with datasets where different categorical	<i>BY department ORDER BY COUNT(*) DESC</i>	1481
1436	attributes can provide valuable insights when	<i>LIMIT 1;</i>	1482
1437	analyzed separately. For example, consider	The second query should instead filter using	1483
1438	a dataset of department management. The	HAVING, such as selecting age groups where	1484
1439	original question:	the number of teachers exceeds a threshold:	1485
1440	<i>Tell me which states have at least 3 heads born</i>	<i>SELECT age FROM teachers GROUP BY age</i>	1486
1441	<i>there?</i>	<i>HAVING COUNT(*) > 2;</i>	1487
1442	The corresponding SQL query would look like	This variation ensures that both queries serve	1488
1443	this:	different analytical purposes while maintain-	1489
1444	<i>SELECT born_state FROM head GROUP BY</i>	ing structural diversity. By alternating be-	1490
1445	<i>born_state HAVING COUNT(*) >= 3;</i>	tween HAVING and LIMIT, we create queries	1491
1446	However, what if we want to explore a differ-	that offer different perspectives on the dataset,	1492
1447	ent way of grouping the data? We could also	preventing them from being collapsed into a	1493
1448	ask:	single SQL statement.	1494
		9. Reordering the Same Data from Different	1495
		Perspectives	1496

1497	This method transforms a query that sorts data	To introduce variation while maintaining	1544
1498	using ORDER BY into two separate queries,	meaningful structure, we modify the query	1545
1499	each sorting the same dataset by a different	to introduce an additional unordered dataset.	1546
1500	attribute. Unlike cases where LIMIT is in-	Instead of just listing teachers sorted by expe-	1547
1501	involved, which restricts the number of results,	rience, we now also retrieve the departments	1548
1502	this approach ensures that both queries return	they belong to, but without any explicit sorting	1549
1503	the full dataset, just presented in different or-	applied to them.	1550
1504	ders.		
1505	Consider the original query:	Modified Question:	1551
1506	<i>Find the names of all teachers, sorted by age</i>	<i>Find the names of all teachers, sorted by their</i>	1552
1507	<i>in descending order.</i>	<i>ages in descending order, and separately list</i>	1553
1508	Original SQL Query:	<i>all department names.</i>	1554
1509	<i>SELECT name FROM teachers ORDER BY</i>	Corresponding SQL Queries:	1555
1510	<i>age DESC;</i>	<i>SELECT name FROM teachers ORDER BY</i>	1556
1511	To introduce variety, we modify the query by	<i>age DESC;</i>	1557
1512	adding an alternative sorting criterion. Instead	<i>SELECT department_name FROM depart-</i>	1558
1513	of only ordering by age, we ask for two sep-	<i>ments;</i>	1559
1514	arate orderings—one by age and another by	This modification ensures that two separate	1560
1515	salary. The question can be transformed into:	results are produced, but unlike the previous	1561
1516	<i>Give me two separate lists of teacher names:</i>	approach (where both queries applied sort-	1562
1517	<i>one sorted by age in descending order and the</i>	ing in different ways), here only one result is	1563
1518	<i>other sorted by salary in descending order.</i>	ordered while the other remains unaffected.	1564
1519	Corresponding SQL Queries:	This distinction is important because it em-	1565
1520	<i>SELECT name FROM teachers ORDER BY</i>	phasizes that ORDER BY is not necessarily	1566
1521	<i>age DESC;</i>	required for every SQL statement—it depends	1567
1522	<i>SELECT name FROM teachers ORDER BY</i>	on the nature of the query and what needs to	1568
1523	<i>salary DESC;</i>	be compared or presented.	1569
1524	This transformation ensures that both results		
1525	remain distinct and cannot be merged into a		
1526	single SQL query.		
1527	10. Sorting One Object While Leaving the	11. Sorting One Object with LIMIT Condition	1570
1528	Other Unordered	While Leaving the Other Unordered	1571
1529	This method builds upon the previous ap-	This method is quite similar to the previous	1572
1530	proach of modifying ORDER BY queries but	method (10) in that it requires an SQL query	1573
1531	introduces a crucial difference: instead of ap-	with an ORDER BY clause and an extra query	1574
1532	plying sorting criteria to both SQL queries,	remaining unsorted. However, the ORDER	1575
1533	one query retains the ORDER BY clause	BY query also includes an additional LIMIT	1576
1534	while the other remains unsorted. The key	condition. Furthermore, we intentionally se-	1577
1535	idea is to introduce contrast—one dataset fol-	lect the additional SQL query to include a	1578
1536	lows an explicit order while the other remains	WHERE clause while minimizing the use of	1579
1537	in its natural or default arrangement.	any aggregation functions. This is set to make	1580
1538	Consider the original query:	the problem more challenging and to ensure	1581
1539	<i>Find the names of all teachers, sorted by age</i>	the conditions of both 2 queries are clearly	1582
1540	<i>in descending order.</i>	described to avoid any potential ambiguity.	1583
1541	Original SQL Query:	Consider the original query:	1584
1542	<i>SELECT name FROM teachers ORDER BY</i>	<i>Find the name of the oldest teacher</i>	1585
1543	<i>age DESC;</i>	Original SQL Query:	1586
		<i>SELECT name FROM teachers ORDER BY</i>	1587
		<i>age DESC LIMIT 1;</i>	1588
		To introduce variation while maintaining	1589
		meaningful structure, we modify the query to	1590

1591	introduce an additional unordered query with	Based on the the detailed explained pattern	1637
1592	a WHERE clause. We intentionally select an	above, we introduce an additional query to se-	1638
1593	additional query with a WHERE clause and	lect another attribute in the same table sorted	1639
1594	avoid using any aggregation function.	by exactly the same attribute (e.g, salary)	1640
1595	Modified Question:	but in the reverse order (e.g, ASC). We can	1641
1596	<i>Find the name of the oldest teacher and the</i>	also change the number of results returned by	1642
1597	<i>ages of all teachers whose names contain the</i>	changing the LIMIT clause (e.g, the lowest	1643
1598	<i>letter "A".</i>	three salaries).	1644
1599	Corresponding SQL Queries:	Modified Question:	1645
1600	<i>SELECT name FROM teachers ORDER BY</i>	<i>Show the name of the teacher with the highest</i>	1646
1601	<i>age DESC LIMIT 1;</i>	<i>salary and the office numbers of the teachers</i>	1647
1602	<i>SELECT age FROM teachers WHERE name</i>	<i>who have the lowest three salaries.</i>	1648
1603	<i>LIKE "%A%";</i>	Corresponding SQL Queries:	1649
1604	This modification is similar to the previous	<i>SELECT name FROM teachers ORDER BY</i>	1650
1605	method, but introduces more conditions: One	<i>salary DESC LIMIT 1;</i>	1651
1606	query is required to select only some results	<i>SELECT office_no FROM teachers ORDER</i>	1652
1607	from the whole sorted list, while another query	<i>BY salary ASC LIMIT 3;</i>	1653
1608	should return the unsorted results under the	This modification requires the agent to output	1654
1609	WHERE condition. This type of modification	two distinct SQL queries to select two sub-	1655
1610	further challenges agents' ability to compre-	sets of different attributes based on the sorted	1656
1611	hend the nature of the questions.	result of the same attribute but in reverse or-	1657
1612	12. Dual Attribute Selection Sorted by the	ders. This type of modification could pose	1658
1613	Same Attribute but in Reverse Orders with	challenges for some agents, as they must se-	1659
1614	LIMIT Condition	lect two distinct attributes based on a condi-	1660
1615	This method also requires an original SQL	tion and its nearly opposite form.	1661
1616	query to return results constrained by the	13. Dual Attribute Selection Sorted by the	1662
1617	LIMIT clause sorted by some attributes. We	GROUP BY clause of the Primary Key but	1663
1618	then add another query to select another at-	in Reverse Orders with LIMIT Condition	1664
1619	tribute sorted by exactly the same attribute but	This method can be considered as a more chal-	1665
1620	in reverse order. Notice that the results are still	lenging variant of the previous method (12) as	1666
1621	constrained by the LIMIT clause. Intuitively,	the main logic is quite similar. The only dif-	1667
1622	one simple way to achieve this type of modi-	ference is that we use the GROUP BY clause	1668
1623	fication is to transform the statements in the	of the primary key of the table to do sorting	1669
1624	original question that use superlative forms	(e.g, GROUP BY id ORDER BY COUNT(*)	1670
1625	into their opposites (e.g., changing 'maxi-	DESC LIMIT 1). This task is considered more	1671
1626	mum' to 'minimum' and 'most' to 'least').	challenging as the original question requires	1672
1627	Also, we intentionally include some questions	the agent to understand the structure of the	1673
1628	to change the number of results required by	table and use the GROUP BY clause properly.	1674
1629	the LIMIT clause to make this type of modifi-	A more complex task in the same setting to se-	1675
1630	cations more variant and more challenging.	lect two attributes sorted by some attributes of	1676
1631	Consider the original query:	groups in reverse orders can result in greater	1677
1632	<i>Show the name of the teacher with the highest</i>	confusion for agents.	1678
1633	<i>salary.</i>	Consider the original query:	1679
1634	Original SQL Query:	<i>Show the status of the city that has hosted the</i>	1680
1635	<i>SELECT name FROM teachers ORDER BY</i>	<i>greatest number of competitions.</i>	1681
1636	<i>salary DESC LIMIT 1;</i>	Original SQL Query:	1682
		<i>SELECT T1.Status FROM city AS T1</i>	1683
		<i>JOIN farm_competition AS T2 ON</i>	1684

```
T1.City_ID = T2.Host_city_ID GROUP
BY T2.Host_city_ID ORDER BY COUNT(*)
DESC LIMIT 1;
```

Based on the additional explanation of this variant, we select another proper attribute (e.g, the primary key itself) in the table and only changes the sorting order to the reverse one.

Modified Question:

Show the status of the city that has hosted the greatest number of competitions and the city ID with the fewest competitions.

Corresponding SQL Queries:

```
SELECT T1.Status FROM city AS T1
JOIN farm_competition AS T2 ON
T1.City_ID = T2.Host_city_ID GROUP
BY T2.Host_city_ID ORDER BY COUNT(*)
DESC LIMIT 1;
```

```
SELECT T1.City_ID FROM city AS
T1 JOIN farm_competition AS T2 ON
T1.City_ID = T2.Host_city_ID GROUP BY
T2.Host_city_ID ORDER BY COUNT(*) ASC
LIMIT 1;
```

14. Question Generation with Three SQL Queries

We also introduce a small proportion of questions that require three SQL queries to address. These questions are generated based on some of the previous methods. As this could be a more challenging and difficult task for agents, we only select some relatively simple methods to generate questions.

14.1 Combining One Original Question with One Already Generated Question Requiring Two SQL Queries

An intuitive way to generate questions requiring three SQL queries to address is to combine one original question from the dataset with one generated and well-evaluated question requiring two queries. To ensure that the final generated question is reasonable, relatively concise, and unambiguous, we intentionally select those original SQL queries that include a WHERE clause while excluding certain keywords that may form complex questions (such as LIMIT, GROUP, EXCEPT, INTERSECT, UNION, ORDER) and control the length of the selected

SQL queries. Then we sample a well-evaluated question from our generated dataset and ask the agent to merge the two questions. The agent is also suggested to use some keywords like "separately" or "first, then," to avoid any potential ambiguous statements.

14.2 Querying Different Grouping Dimensions to Generate three queries

This method is quite similar to the previous method 7 as we still ask to group the table data based on different attributes separately, but now three different attributes are used instead of only two.

Two example generated questions are shown below:

1. *Show the number of teachers for each department, age, and country separately.*
2. *Find the department, the age group and the country group with the highest number of teachers.*

14.3 Reordering the Same Data from Three Different Perspectives

This method is quite similar to the previous method 9 as we still ask to provide lists of one single attribute sorted separately by different attributes, but this time we require three lists. Notice that each query sorted by each attribute can still be in ascending or descending order randomly.

An example generated questions are shown below:

Provide me with three lists of teacher names, sorted by age, department id and salary in descending order, respectively.

14.4 Two or Three Queries with DISTINCT

Previously our methods didn't consider the keyword DISTINCT to avoid any potential ambiguous statement. Here we ask the agent to generate a small subset of questions to ask for distinct values for selected columns only. The generated questions may require two or three queries to return only distinct/unique values for mentioned attributes.

A simple example question is shown below:

Provide me with distinct teacher name, unique department names, and distinct teacher ages, respectively.

B Implementation Details

All experiments were conducted using LLAMA Factory (Zheng et al., 2024), with LoRA-based (Hu et al., 2021) fine-tuning utilizing the bf16 precision. The relevant hyperparameters were configured as follows:

- preprocessing_num_workers: 16
- per_device_train_batch_size: 1
- gradient_accumulation_steps: 4
- learning_rate: 1.0e-4
- num_train_epochs: 3.0
- lr_scheduler_type: cosine
- warmup_ratio: 0.1

Other parameters were set to the default values provided by LLAMA Factory (Zheng et al., 2024). The experiments were conducted on single or multiple A100 GPUs (40GB/80GB). For multi-GPU setups, DeepSpeed was used to optimize parallel training.