

Semantically Aligned Question and Code Generation for Automated Insight Generation

Ananya Singha
t-asingha@microsoft.com
Microsoft
India

Bhavya Chopra
t-bhchopra@microsoft.com
Microsoft
India

Anirudh Khatri
t-anikhatri@microsoft.com
Microsoft
India

Sumit Gulwani
sumitg@microsoft.com
Microsoft
USA

Austin Z. Henley
azh321@gmail.com
Microsoft
USA

Vu Le
levu@microsoft.com
Microsoft
USA

Chris Parnin
chrisparnin@microsoft.com
Microsoft
USA

Mukul Singh
singhmukul@microsoft.com
Microsoft
India

Gust Verbruggen
gverbruggen@microsoft.com
Microsoft
Belgium

ABSTRACT

Automated insight generation is a common tactic for helping knowledge workers, such as data scientists, to quickly understand the potential value of new and unfamiliar data. Unfortunately, automated insights produced by large-language models can generate code that does not correctly correspond (or *align*) to the insight. In this paper, we leverage the semantic knowledge of large language models to generate targeted and insightful questions about data and the corresponding code to answer those questions. Then through an empirical study on data from Open-WikiTable, we show that embeddings can be effectively used for filtering out semantically unaligned pairs of question and code. Additionally, we found that generating questions and code together yields more diverse questions.

ACM Reference Format:

Ananya Singha, Bhavya Chopra, Anirudh Khatri, Sumit Gulwani, Austin Z. Henley, Vu Le, Chris Parnin, Mukul Singh, and Gust Verbruggen. 2024. Semantically Aligned Question and Code Generation for Automated Insight Generation. In *2024 International Workshop on Large Language Models for Code (LLM4Code '24)*, April 20, 2024, Lisbon, Portugal. ACM, New York, NY, USA, 8 pages. <https://doi.org/10.1145/3643795.3648381>

1 INTRODUCTION

In the *Age of Copilots*¹, numerous companies have developed *copilots* that leverage large-language models (LLMs) to answer and perform tasks directly in products. To facilitate user interaction

with the product copilot, a common tactic is to automatically generate suggested questions or followups. This tactic is also common in exploratory data analysis, where automated insights are generated to help data scientists initiate interaction with data [10]. More generally, a *good question* can be an effective tool for human knowledge acquisition [4], and effective problem-solving [14].

Outcome	Year	Championship	Opponent in final	Score
Winner	1990	World Snooker Ch.	Jimmy White	18–12
Winner	1990	Grand Prix (2)	Nigel Bond	10–5
Winner	1987	Grand Prix	Dennis Taylor	10–7
Winner	1990	Asian Open (2)	Dennis Taylor	9–3
Winner	1990	Dubai Classic (2)	Steve Davis	9–1
Winner	1990	UK Championship (2)	Steve Davis	16–15
Winner	1988	British Open	Mike Hallett	13–2
Runner-up	1988	UK Championship	Doug Mountjoy	12–16
Winner	1989	Asian Open	James Wattana	9–2
Winner	1989	Dubai Classic	Doug Mountjoy	9–2
Winner	1989	UK Championship	Steve Davis	16–12
Runner-up	1990	European Open	John Parrott	6–10
Runner-up	1991	Classic	Jimmy White	4–10
Runner-up	1989	International Open	Steve Davis	4–9

Figure 1: A table displaying the snooker championship results from the Open-WikiTable corpus.

In this paper, we leverage large language models [1] to generate **semantically-aligned** question and code pairs for supporting automated insight generation. Data scientists (and other users) often explore a new dataset, by using tools, writing queries, or code. Unfortunately this process takes up a significant part of data-driven workflows [3, 5]—understanding how to select and combine operations such as sorting, filtering, reshaping, grouping, joining, and pivoting data can be time-consuming. This problem can be made even worse if a data scientist is misled by using incorrect code generated by an LLM.

For example, consider a table with snooker data for a specific player, as shown in Figure 1. Following are a few examples of the automated insights generated by a LLM.

In which championship did the player face the toughest competition based on the average frames played in the finals?

¹<https://www.forbes.com/sites/forbestechcouncil/2023/09/08/the-age-of-co-pilots>

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions.acm.org.

LLM4Code '24, April 20, 2024, Lisbon, Portugal

© 2024 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 979-8-4007-0579-3/24/04...\$15.00
<https://doi.org/10.1145/3643795.3648381>

```
> (df.groupby('Championship')['Score']
    .apply(lambda x:
        sum(int(score.split('-')[0]) +
            int(score.split('-')[1])
            for score in x) / len(x))
    .idxmax())
'UK Championship (2)'
```

How many championships did the player win consecutively?

```
> df.sort_values(by=['Year'])
    .groupby(
        (df['Outcome'] != df['Outcome'].shift())
        .cumsum())['Outcome']
    .transform('size')
    .max()
5
```

The above insights, highlighted in green, are useful insights that help in understanding different segments of the tabular data and are beneficial to provide a head start with code that can be used and extended. Unfortunately, the same approach, without any additional filtering, also generates misleading and incorrect insights or code, potentially negating the value of the entire approach:

What is the overall trend in the player's performance over the years?

```
>>> 'Increasing' if
    df['Outcome'].eq('Winner').sum() >
    df['Outcome'].eq('Runner-up').sum() else
    ↪ 'Decreasing' if
    df['Outcome'].eq('Winner').sum() <
    df['Outcome'].eq('Runner-up').sum()
    else 'Fluctuating'
'Increasing'
```

→ While the insight is interesting, the code does not appropriately account for multiple championships in years.

What is the average score in the finals when the player wins compared to when they are the runner-up?

```
>>> df[df['Outcome'] == 'Winner']['Score']
    .apply(lambda x:
        sum(map(int, x.split('--'))))
    .mean()
18.0
```

→ The insight is computing interesting stats about the data but the calculation is incorrect. The calculation must take the first part of the score after splitting rather than both.

We conducted an empirical study to understand how we can correctly identify semantically-aligned question and code amongst all generated pairs. First, we conduct a user study to understand the relevance of the question and code pairs we generate using LLMs. Then, to determine if questions and code align, we compare multiple models on their ability to rate the correctness of (question, code) pairs. Finally, to measure diversity, we compute the edit distance between generated code snippets after masking constants. Based on our study we found that: (1) users found most generated code pairs to be interesting and meaningful for their work, (2) a semantic alignment classifier based on code embeddings performs on par with GPT-4 [13] on a human-annotated dataset and (3) generating questions and code together gives more diverse inspirations.

2 RELATED WORK

LLM are incrementally used for tasks that involve generating executable code given some textual context. This marks one of the challenging problems because, unlike text, sequentially generated code has to maintain its semantic and syntactical correctness. Several fine-tuning and training efforts have been carried out over the years to enhance the natural language to code generation capabilities. For example, CodeBERT (BERT), CodeT5 (T5) and Codex (GPT-3) have been fine-tuned for text-to-code tasks for multiple domains like pandas, Java etc. Huge training datasets, bigger models, expert tuning approaches, enriched training objectives specializing in code context have led to improvements from the baselines in terms of lower syntactical errors. Though semantics still remains a concern [19] and there have been fewer efforts in generating similar code and natural language together for tabular data analysis.

An integral part of the data science workflow involved analysing the tabular data. Data analysis tasks such as data cleaning [6], formatting [15], transformation [17], and visualization [9] have been explored and collectively fall under the workflow of Exploratory Data Analysis (EDA). Previously these tasks were carried out by a symbolic setup like MetaInsights [10] which mine common data pattern to show insights. Similarly, LLM are also used as an orchestrator along with the symbolic systems to rank and summaries insights [11]. Though none of these works focus on the alignment issue between the question and the generate code.

Finally, our work is closely related to the research on automatically generating question and answers in various domains, and additional challenges associated with alignment [2, 16], and diversity [18] of results. In contrast to these approaches, we specialize our approach in the table code generation domain and to handle tabular/code specific issues, leverage LLMs to derive semantic table and column information during synthesis, and leverage embeddings to train cost-effective classifiers for measuring alignment.

3 PROPOSED METHOD

Given a table T , the goal is to find question and program pairs (q_i, p_i) . We make sure that the generated pairs (q_i, p_i) are aligned and $p_i(T)$ correctly answers q_i . Additionally, we need to ensure that we maintain a balance between diversity (unique (q_i, p_i) pairs) and complexity (q_i is non-trivial and provides an interesting insight). We make no assumption of the type of the data present inside each cell of the given table and consider it to be string.

3.1 Architecture

Figure 2 shows an overview of our approach. ① We analyse the given table to ② build a prompt that contains relevant information about the table, and we ③ instruct the model to generate questions and code pairs. ④ We then clean invalid and non-executable candidates, and finally ⑤ remove those candidates where the generated code does not align with the question. These stages of our approach are detailed in the following sections.

3.2 Table Analysis and Prompt Creation

The prompt contains three main relevant pieces of information, which together allow the model to perform the right wrangling operations when providing the code for insightful questions.

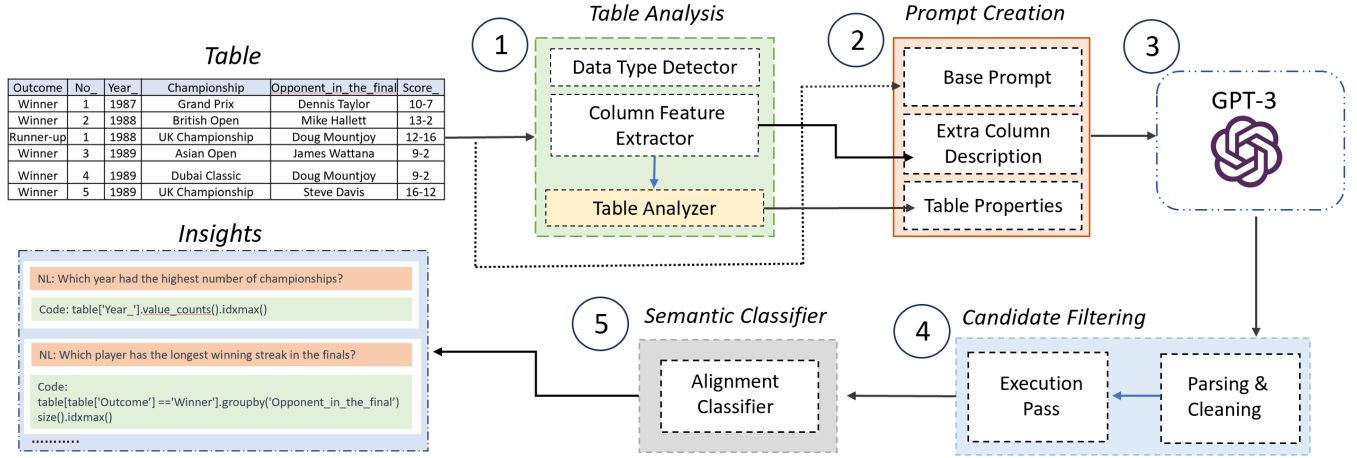


Figure 2: Overview of approach. ① The table is analysed to gather relevant information to ② build a prompt that helps the model ③ to provide interesting insights. ④ The output of the model is cleaned and invalid suggestions are filtered. ⑤ The suggestions are then screened by the Semantic Classifier to present the actionable insights.

First, a small subset of rows allows the model to understand the syntax of each column. For example, in Figure 1, an example of the score “10–7” is required to understand that splitting on “–” and integer conversion is required to compute averages.

Second, properties about columns provide the model with a complete picture about the scope of the data. These properties are the appropriate data-type information (integer, float, string, and date-time) and number of missing elements of each column, and cardinality, extrema, peak-frequency, column position and entropy for each column.

Third, we leverage these features to predict columns that are suitable candidates for grouping and aggregation. A similar approach was used by Auto-Suggest [17] for generating such operations, and we use their dataset for training the model. Instead of complex post-processing, we only use the predictor to provide a weak signal to the LLM, which then leverages its semantic understanding to refine these candidates. Our predictor is a decision tree, which achieves 88% accuracy on an 80/20 split.

3.3 Filtering and Semantic Alignment

We instruct the model to generate question and code pairs and filter pairs for which p_i cannot be executed on the input table. Inspecting the resulting code pairs, we discovered that the question and code do not always align well together.

For example, consider the generated insight for the snooker data from Figure 1: *Which player has the longest winning streak in the finals?*

```
> table[table['Outcome'] == 'Winner']
   .groupby('Opponent_in_final')
   .size().idxmax()
"Steve Davis"
```

While this is an interesting question, the corresponding code does not take consecutive years into account. Similarly, consider the generated insight for another table with columns [year, starts, wins,

poles, position, teams]: *What is the average number of starts by drivers in the top 5 positions?*

```
>table[table['Position']
   .str.contains('5th')] ['Starts']
   .mean()
"23"
```

The code only takes the 5th position in account, rather than the top 5, and this question and code pair is not semantically aligned.

We propose to predict the semantic alignment of question-code pairs (q, p) by leveraging pre-trained embeddings of text and code [12]. Instead of directly using the embeddings, we train a small, fully connected network to predict semantic alignment from the concatenated embeddings of q and p . Binary cross entropy between the predicted and the ground truth label is used as loss function.

We use a dataset of natural language descriptions and pandas expressions [7] to create a dataset of aligned and misaligned pairs. By using GPT-3 to generate code from these descriptions and using the execution result, we obtain 2209 aligned and 422 misaligned pairs. Next, we swap the code for two pairs (q_1, p_1) and (q_2, p_2) to generate additional misaligned pairs (q_1, p_2) and (q_2, p_1) . In total, this resulted in 1552 misaligned pairs. The semantic alignment classifier achieves 95.6% F1 score on 80/20 split.

4 USER STUDY

Effective insights should be perceived as useful and relevant to data analysis tasks. Prior to evaluating the semantic alignment of generated question and code pairs, we ensure the insights generated by our approach, enable productive exploration of datasets and are sufficiently complex. To measure the quality of generated questions, we performed a user study where we asked annotators to rate insights and provide feedback.

4.1 Methodology

4.1.1 Participants. We recruited 5 participants from within our organization, who have a background in computer science, and varying levels of experience with data analysis tasks.

4.1.2 Task. Participants were asked to examine a sample of 12 tables randomly selected from our dataset, and rate 6–7 insights generated for each table. To ensure that participants spent sufficient time on each table, we first asked them to describe five exploratory analysis questions they might have about the table. This enabled participants to get acquainted with the data, and identify meaningful metrics and columns of interest from the data. Next, for each of the 76 insights presented, participants answered each of the following questions on a 7-point (1–7) Likert-scale:

- (M1) This insight is relevant and enhances my understanding of the table.
- (M2) This insight saves time and improves my productivity.
- (M3) I would not have been able to independently arrive at a similar query as suggested by this insight.

Ratings provided for these questions help us assess the relevance and usefulness of generated questions. Once the participants have provided ratings for all 76 insights, we ask them for their overall feedback and additional comments.

4.1.3 Analysis. To analyze the ratings and assess which insights are considered meaningful by the participants, we use three metrics, one corresponding to each Likert-scale question: Relevance (M1), Productivity (M2), and Ingenuity (M3). We then rank all insights using a combined metric:

$$\text{Relevance} + \text{Productivity} + \text{Ingenuity}$$

We qualitatively analyze the ranked insights to identify desirable properties that contribute to increased usefulness of insights, and properties that lead to lower ratings.

4.2 Findings

4.2.1 Ratings for Questions/Insights. 76 insights generated over 12 tables by our method were inspected by five raters for perceived relevance, improvements in productivity, and ingenuity. We report the ratings obtained in Table 1. Participants positively rated the relevance of insights overall (79% agreement, median 6), and find them to be time-saving in analyzing data (77% agreement, median 6). Lastly, participants find the presented insights to be a mix of trivial and complex generations, as suggested by the low agreement on Ingenuity of insights (35.5% agreement, median 3).

4.2.2 Qualitative findings. We qualitatively inspected the ranked insights (as discussed in 4.1.3) to find desirable properties that make generated questions useful. We find three overarching categories of insights, and present representative examples for each stratum:

The **top 20 percentile** insights provide highly semantically relevant statistics about the table. These insights have high M1, M2, and M3 ratings. They present non-trivial code generations that use complex operations in moderately long chains (such as groupby, apply, and diff). These insights highlight the most critical data points and exploratory findings while having high entropy (in other words, being interesting to the raters). Following are examples from this category of insights:

How many locations have a capacity in use greater than 100% (Rank 7; Score 13.67)

```
> len(table[table['Capacity_in_use']
               .str.rstrip('%')
               > 100])
```

→ The LLM realizes that the data reports transport capacity and utilization across cities, and the insight brings to attention the cities with overburdened transportation facilities.

What is the range of municipality areas in each region? (max area - min area) (Rank 8; Score 13.5)

```
> table.groupby('Regional_County_Municipality')
      ['Area_km_2_']
      .apply(lambda x: x.max()-x.min())
```

→ The insight groups the data by regions and uncovers the range of land area. Our model is able to pick the correct columns names, while presenting a complex and meaningful insight.

The next stratum of insights lying in **20–80 percentile** have high M1 and M2 ratings, with low M3 ratings—these insights are relevant, and improve productivity, but are commonly thought of by the raters (indicated by low *Ingenuity* ratings). These insights involve computation of meaningful but commonly reported statistics (such as min, max, avg, value counts, and top-n) after performing groupby operations, or conditional filtering of data. Some examples include:

How many games have been released in all three regions? (Rank 16; Score 13.00)

```
> table.loc[(table['Europe'] != 'N/A N/A')
            & (table['North_America'] != 'N/A N/A')
            & (table['Japan'] != 'N/A N/A')]
      .shape[0]
```

→ The insight presents a sufficiently complex code snippet, but is commonly thought of by users. It identifies that the table contains release dates in three countries and leverages missing values to count games available in all countries.

What is the most common nationality among the directors who have won awards? (Rank 41; Score 11.8)

```
>>> table.groupby('Nationality_of_director')
      ['Award'].count().idxmax()
```

→ Though this question helps participants gain a highly meaningful insight about their data, it is commonly thought of.

Upon analyzing the **bottom 20 percentile** insights we find three sub-categories, exposing characteristics of insights that are undesirable for participants. These findings further motivate the need for a filtering mechanism:

- Majority of the insights in this strata compute trivial statistical measures for the data (such as min, max, mean, sum, nunique). Participants provided low M1, M2, and M3 ratings for these questions. For instance:

How many unique NFL teams are present in the data? (Rank 62; Score 10.16)

```
> table['NFL_Team'].nunique()
```

- Few insights contained highly instance-specific conditions or constraints to filter data, and are not perceived to be useful

Table 1: Study Survey Responses

	% Agree	Likert Response Counts ¹							Distribution ²
		SD	D	SwD	N	SwA	A	SA	
The insight is relevant & enhances my understanding.	79.21%	12	30	17	20	67	100	134	
The insight saves time & improves my productivity.	76.84%	14	35	15	24	62	105	125	
I would not have independently arrived at a similar query.	35.52%	77	64	59	45	48	55	32	

¹ Strongly Disagree (SD), Disagree (D), Somewhat Disagree (SwD), Neutral (N), Somewhat Agree (SwA), Agree (A), Strongly Agree (SA).

² Net stacked distribution removes the Neutral option and shows the skew between positive and negative responses.

■ Strongly Disagree, ■ Disagree, ■ Somewhat Disagree, ■ Somewhat Agree, ■ Agree, ■ Strongly Agree.

despite performing complex operations. Such insights also receive low ratings on M1, M2, and M3. For instance:

What is the highest position secured by Scissor Sisters in the chart? (Rank 76; Score 8.00)

```
> table[table['Artist']
      == 'Scissor Sisters']
      ['Highest_Position'].max()
```

→ Here, participants did not find much use in insight finding for ‘Scissor Sisters’ in particular, and expected the insight to be more generic.

- Lastly, two insights could not be drawn from the provided table, and their code generations used incorrect target column(s). For instance, one of these insights attempts to report information on population age, while the table does not contain any columns that indicate age: **Which region has municipalities with the highest and lowest average population age?** (Rank 68; Score 10.33)

4.2.3 End-of-study survey. Ratings from the end-of-survey questions show that 4 of 5 participants found it challenging to get acquainted with the data, and identify meaningful insights that can be derived — “It is difficult to think of meaningful queries without context of the data. While the tables in most cases were self-explanatory, there were a few which I just didn’t understand. I think generated insights were very helpful in this” (R3). Responses from the open-field feedback suggest that insights generated using our method help participants by (1) enhancing their understanding of the data, (2) identifying columns of interest, and (3) suggesting non-trivial and semantically relevant insights.

5 EVALUATION

We conducted an empirical study to answer the following research questions:

- RQ1** Can we leverage embeddings to select semantically aligned (question, code) pairs from model generations?
- RQ2** How do the generation task, number of insights and prompt influence the diversity of insights?

We used the Open-WikiTable corpus [8] for our experiments. The dataset contains 24,680 tables, and natural language questions. The dataset was specifically constructed for the purpose of evaluating systems that perform complex reasoning tasks over tables.

To ensure our evaluation would be tractable, we randomly selected 430 tables from the 24,680 tables available.

5.1 Semantic Alignment (RQ1)

We evaluate whether embeddings can be used to detect semantic alignment, and perform an ablation on different strategies of combining the embeddings for question and code.

5.1.1 Experimental Setup. We compare our semantic alignment classifier against data annotated by humans.

To create our annotated dataset, we first generated 10,175 insights across 430 tables—about 25 insights per table. For some tables, fewer insights were produced when token limits were exceeded. We were able to execute 8,954 out of 10,175 (88%) of the generated insights. From the executable insights, a sample of 309 (question, code) pairs—stratified on code length—was annotated by expert data scientists as semantically aligned or not.

We use the OpenAI gpt-3.5-turbo, gpt-3.5-turbo-16k, and gpt-4 models as baselines for comparison. Table 2 shows individual annotation statistics and an ensemble that only considers alignment if all other annotators (human and LLM) agree.

Table 2: Annotation dataset with 309 data points amongst different annotators.

Annotator	Positive	Negative
Human	251	58
GPT-4	222	87
GPT-3.5	230	79
GPT-3.5 16 K	228	81
Ensemble	173	136

Besides concatenating individual embeddings, we also consider two settings to test performance: (1) classification using a single embedding for question and code, and (2) applying two fully connected networks to the individual embeddings and measuring the cosine similarity between the projected embeddings. We show the performance of each technique on the annotated data.

5.1.2 Results. Figure 3 shows precision-recall (PR) curves for different variations of the embedding models on human labelled data. It highlights that concatenating embeddings (Concat) retains the

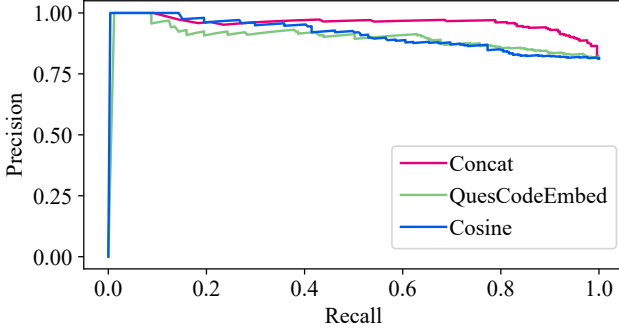


Figure 3: PR curve showing performance of three different embedding variations, on annotated dataset.

highest precision for more recall. At a recall of 80%, this variation is 5.74% and 6.8% more precise than the embedding question and code together and the cosine similarity model, respectively. We hypothesize that full attention between both embeddings allows the classifier to capture the most interactions.

Table 3 shows agreement between different models (our classifier and LLM baselines) and the human annotations, showing that our classifier (90%) performs on par with GPT-4 (89.2%). Our embedding-based classifier is 728 times more cost effective² than GPT-4.

Table 3: Agreement between models and human annotator.
** best setting (Concat).

Classifier	TP	TN	FP	FN	Acc.	F1
GPT-4	211	47	11	40	83.5%	89.2%
GPT-3.5	200	27	30	51	73.7%	83.1%
GPT-3.5 16K	198	27	31	53	72.8%	82.5%
Embed**	217	46	12	34	85.1%	90.0%

Figure 4 shows how the alignment (logits from the classification model) slightly degrades as more insights are generated. We observed that the complexity of the insights increases as they were generated, which increases the potential for semantic alignment errors. For example, an initial insight generated for the table in Figure 1 is *Which year had the highest number of championships?*

```
> table[table['Year'].value_counts().idxmax()]
```

As we generate more insights, the complexity increases. Two examples are *Which position has the highest number of goals in the league?*

```
> table.groupby(['Position_'])
    .agg({'League_Goals':sum})
    .sort_values(by=['League_Goals'],
                ascending=False)
    .iloc[0].name
```

and *What is the most common party of Presidents who served multiple terms?*

²Currently, the average embedding computation cost per question and code pair is 0.007 cents in comparison to 5.1 cents by gpt-4. (7th December 2023)

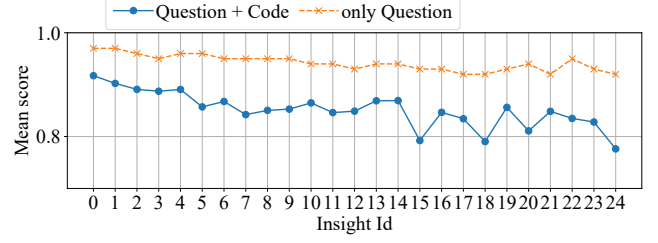


Figure 4: Distribution of semantic alignment between generated question and code pairs across different insights for tables.

```
> table[table['Name_'].isin(
    table['Name_'].value_counts()[
        table['Name_'].value_counts()>1
    ].index)
].groupby('Party')['Name_']
    .count()
    .sort_values(ascending=False)
    .iloc[0]
```

Complex questions typically require grouping, sorting and filtering.

By manually inspecting cases where our classifier disagreed with human labels, we encountered situations where the distinction between aligned and misaligned was blurry. This suggests that additional context might be necessary to identify alignment. For example: *Which universities had the highest number of students sign up in a given year?*

```
> table.loc[table.groupby('Year')['#SignedUp']
    .idxmax()]
    [['Year', 'University', '#SignedUp']]
```

The classifier identifies the above case as aligned, while humans disagree due to the presence of additional information (year and number of sign-ups).

5.1.3 Execution Prediction. Executing code generated by the language model is not desired in production settings. We therefore investigate whether the semantic alignment classifier can detect if code will execute or not.

Figure 5 shows the logits distribution of the classifier over both executable and non-executable pairs. Note that the semantic alignment is not known for executable pairs. In general, non-executable pairs are given lower scores—scores close to zero occur the most often. High alignment scores still occur for non-executable scores due to subtle syntax errors, which the embeddings do not capture. Consider the following insight: *Who played the character with the maximum duration on the show?*

```
> table.iloc[table['Duration']
    .apply(lambda x: sum(
        int(i.split('x')[1]) -
        int(i.split('x')[0]) + 1
        for i in x.split('-'))
    .idxmax())['Actor']]
```

The code looks correct, but a parenthesis is unmatched, causing an alignment score of 0.9.

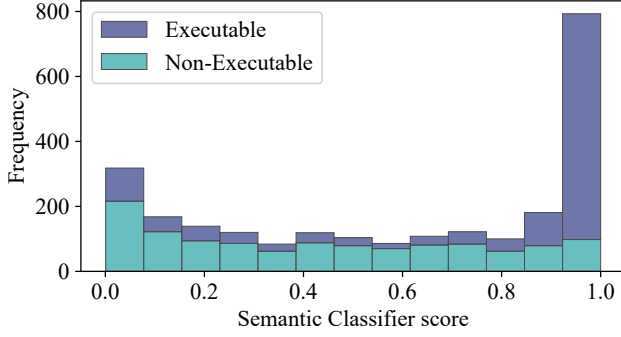


Figure 5: Distribution of semantic alignment scores among the executable and non-executable question code pairs stacked after one after another.

5.2 Diversity (RQ2)

We evaluate how different generation tasks, few-shot prompting and the number of insights influences the diversity of generations.

5.2.1 Experimental Setup. We measure the diversity of two (question, code) pairs by masking constants in the code and computing the edit distance on 8,954 executable insights generated across 430 tables. For example,

```
table[table['Outcome']=='Winner']
    .groupby('Opponent in final')
```

becomes `table[table[[]]==].groupby()`. Diversity of multiple insights is computed as the average diversity across all pairs of code.

Besides generating questions and code together, we consider three variations: (1 and 2) generating only questions or code and performing the corresponding translation task, and (3) generating code and questions (in reverse order). We use the `gpt-3` model again to generate the code for the given question and vice versa. Only executable code generations are retained.

Apart from the zero-shot setup, we also try a one-shot (due to token limitations) setup with 3 static samples of question and code pairs. The goal is twofold: showing the model what interesting questions look like, and ensuring that it follows the desired output format.

5.2.2 Results. Figure 6a shows that generating questions and code together, in that particular order, yields more diverse questions than any other setup. Generating a single modality (questions or code only) yields the least diversity. We hypothesize that having code or questions closer together biases the model towards generating similar output. Generating code only results in the least diverse questions. When generating code and questions (in that order) the model likely pays more attention to the code, and again becomes more repetitive.

Figure 6b shows that providing an example of insights reduces the diversity. We noticed that the model tries to mimic the given example for different columns, which harms the overall diversity. For example, when providing *What is the average salary amongst all participants?* as a seed insight, the model generates questions like *What is the average number of matches played by each team?*

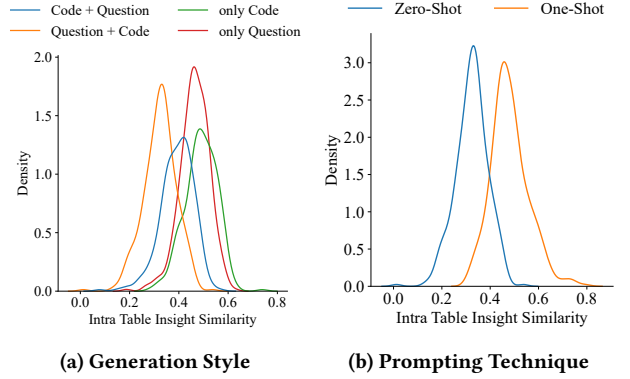


Figure 6: Variation in diversity across different (a) Generation styles (b) Prompting techniques.

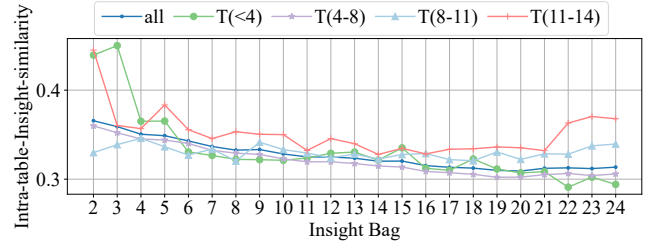


Figure 7: Evolution of diversity as more insights are generated for all tables and divided by number of columns. T(4-8) means tables with 4 to 8 columns. For tables with more columns, the model becomes repetitive.

and *What is the average difference between “Points” and “Against” for all the teams?* In total, 9 out of 25 questions are about averages.

Table 4: Token and time costs associated with different generation styles.

Generation Style	Tokens	Inference (in sec)
Only Question	56,382	121.52
Only Code	12,833	62.3
Code + Question	1,512	19.89
Question + Code	1,515	20.9

Figure 4 shows that generating questions only and then translating to code yields higher alignment (93%). Two drawbacks of this approach are diversity and cost. Table 4 shows the average number of tokens and associated inference time for different approaches. Generating questions and code together is 37 times cheaper and six times faster than generation questions and translating them.

Figure 7 shows that initial generations are diverse, but we see diminishing returns from generating more insights. After around 20 insights, the questions becomes more repetitive. Surprisingly, this effect is stronger for tables with more columns. We find that the model repeats the the exact same insight for multiple columns.

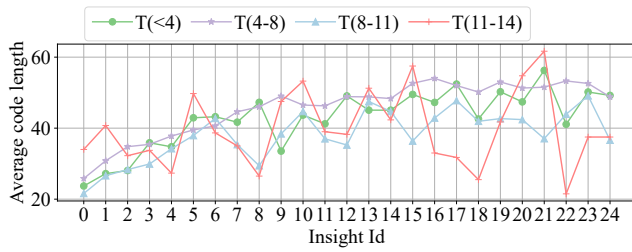


Figure 8: Evolution of code length as more insights are generated.

To verify this effect, Figure 8 shows how the length of the code (after removing column names) evolves for tables with different numbers of columns. For smaller number of columns, the model cannot just repeat itself, and it generates more complex questions.

6 LIMITATIONS

In this work, we have focused solely on English text generation, leaving the exploration of suggestions in other languages as a potential avenue for future research. Furthermore, we observed a notable drop in performance due to the data-distribution shift when testing the model from human annotated to machine-generated data. It is essential to build models that can account for such variations. Additionally, we treat the generated suggestions as an unordered set and do not rank or filter them based on diversity or relevance, which could be a promising direction for future investigation. For generation tasks, we have only considered models from the `gpt-3` and `gpt-3.5` class, as these are state-of-the-art in generation capabilities. However, the results may differ with smaller models, as they may lack in-context learning abilities.

7 CONCLUSION

In this paper, we explore leveraging LLMs to aid automated insight generation, by using them to generate aligned question and code pairs for tabular data. Obtaining diverse and interesting questions, along with the code to generate them—even when transformations are required—is not possible with traditional methods, which rely on statistical features or historical operations from other users. In our study, performed on data from Open-WikiTable, we showed that insights thus generated are meaningful and interesting to users. Further we showed that generating questions and code together yields diverse questions, and that an embedding-based semantic alignment classifier performs on par with GPT-4 for filtering cases where question and code are misaligned at a fraction of the cost.

REFERENCES

- [1] 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 (2020), 1877–1901.
- [2] Jose Cambronero, Hongyu Li, Seohyun Kim, Koushik Sen, and Satish Chandra. 2019. When deep learning met code search. In *Proceedings of the 2019 27th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering*. 964–974.
- [3] Tamraparni Dasu, Theodore Johnson, S. Muthukrishnan, and Vladislav Shkapenyuk. 2002. Mining database structure; or, how to build a data quality browser. In *ACM SIGMOD Conference*.

- [4] Arthur C Graesser, Shulan Lu, George Tanner Jackson, Heather Hite Mitchell, Mathew Ventura, Andrew Olney, and Max M Louwerse. 2004. AutoTutor: A tutor with dialogue in natural language. *Behavior Research Methods, Instruments, & Computers* 36 (2004), 180–192.
- [5] Philip Guo. 2013. Data science workflow: Overview and challenges. *Commun. ACM* (2013).
- [6] Zhipeng Huang and Yeye He. 2018. Auto-detect: Data-driven error detection in tables. In *Proceedings of the 2018 International Conference on Management of Data*. 1377–1392.
- [7] Naman Jain, Skanda Vaidyanath, Arun Iyer, Nagarajan Natarajan, Suresh Parthasarathy, Sriram Rajamani, and Rahul Sharma. 2022. Jigsaw: Large Language Models Meet Program Synthesis. In *Proceedings of the 44th International Conference on Software Engineering (Pittsburgh, Pennsylvania) (ICSE '22)*. Association for Computing Machinery, New York, NY, USA, 1219–1231. <https://doi.org/10.1145/3510003.3510203>
- [8] Sunjun Kweon, Yeonsu Kwon, Seonhee Cho, Yohan Jo, and Edward Choi. 2023. Open-WikiTable: Dataset for Open Domain Question Answering with Complex Reasoning over Table. *arXiv preprint arXiv:2305.07288* (2023).
- [9] Doris Jung-Lin Lee, Dixin Tang, Kunal Agarwal, Thyne Boonmark, Caitlyn Chen, Jake Kang, Ujjaini Mukhopadhyay, Jerry Song, Micah Yong, Marti A. Hearst, and Aditya G. Parameswaran. 2021. Lux: Always-on Visualization Recommendations for Exploratory Dataframe Workflows. *Proceedings of the VLDB Endowment* 15, 3 (nov 2021), 727–738.
- [10] Pingchuan Ma, Rui Ding, Shi Han, and Dongmei Zhang. 2021. Metainsight: Automatic discovery of structured knowledge for exploratory data analysis. In *Proceedings of the 2021 International Conference on Management of Data*. 1262–1274.
- [11] Pingchuan Ma, Rui Ding, Shuai Wang, Shi Han, and Dongmei Zhang. 2023. Demonstration of InsightPilot: An LLM-Empowered Automated Data Exploration System. *arXiv preprint arXiv:2304.00477* (2023).
- [12] Arvind Neelakantan, Tao Xu, Raul Puri, Alec Radford, Jesse Michael Han, Jerry Tworek, Qiming Yuan, Nikolas Tezak, Jong Wook Kim, Chris Hallacy, et al. 2022. Text and code embeddings by contrastive pre-training. *arXiv preprint arXiv:2201.10005* (2022).
- [13] OpenAI. 2023. GPT-4 technical report. *arXiv* (2023), 2303–08774.
- [14] Kumar Shridhar, Jakub Macina, Mennatallah El-Assady, Tanmay Sinha, Manu Kapur, and Mrinmaya Sachan. 2022. Automatic Generation of Socratic Sub-questions for Teaching Math Word Problems. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, Abu Dhabi, United Arab Emirates, 4136–4149. <https://aclanthology.org/2022.emnlp-main.277>
- [15] Mukul Singh, José Cambronero, Sumit Gulwani, Vu Le, Carina Negreanu, Mohammad Raza, and Gust Verbruggen. 2022. CORNET: A neurosymbolic approach to learning conditional table formatting rules by example. *arXiv preprint arXiv:2208.06032* (2022).
- [16] Yue Wang, Weishi Wang, Shafiq Joty, and Steven CH Hoi. 2021. Codet5: Identifier-aware unified pre-trained encoder-decoder models for code understanding and generation. *arXiv preprint arXiv:2109.00859* (2021).
- [17] Cong Yan and Yeye He. 2020. Auto-suggest: Learning-to-recommend data preparation steps using data science notebooks. In *Proceedings of the 2020 ACM SIGMOD International Conference on Management of Data*. 1539–1554.
- [18] Sen Yang, Qingyu Zhou, Dawei Feng, Yang Liu, Chao Li, Yunbo Cao, and Dongsheng Li. 2021. Diversity and Consistency: Exploring Visual Question-Answer Pair Generation. In *Findings of the Association for Computational Linguistics: EMNLP 2021*. Association for Computational Linguistics, Punta Cana, Dominican Republic, 1053–1066. <https://doi.org/10.18653/v1/2021.findings-emnlp.91>
- [19] Daoguang Zan, Bei Chen, Fengji Zhang, Dianjie Lu, Bingchao Wu, Bei Guan, Wang Yongji, and Jian-Guang Lou. 2023. Large language models meet NL2Code: A survey. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 7443–7464.