

Table as Thought: Exploring Structured Thoughts in LLM Reasoning

Zhenjie Sun[🍏], Naihao Deng[🍏], Haofei Yu[🍏], Jiaxuan You[🍏]

[🍏]University of Illinois Urbana-Champaign, [🍏]University of Michigan

Abstract

Large language models’ reasoning abilities benefit from methods that organize their thought processes, such as chain-of-thought prompting, which employs a sequential structure to guide the reasoning process step-by-step. However, existing approaches focus primarily on organizing the sequence of thoughts, leaving structure in individual thought steps underexplored. To address this gap, we propose Table as Thought, a framework inspired by cognitive neuroscience theories on human thought. Table as Thought organizes reasoning within a tabular schema, where rows represent sequential thought steps and columns capture critical constraints and contextual information to enhance reasoning. The reasoning process iteratively populates the table until self-verification ensures completeness and correctness. Our experiments show that Table as Thought excels in planning tasks and demonstrates a strong potential for enhancing LLM performance in mathematical reasoning compared to unstructured thought baselines. This work provides a novel exploration of refining thought representation within LLMs, paving the way for advancements in reasoning and AI cognition.

1 Introduction

Recent advancements in reasoning have demonstrated that the reasoning capabilities of large language models (LLMs) can be enhanced by introducing structure into the reasoning process (Wei et al., 2023; Yao et al., 2023; Besta et al., 2024). For instance, the chain-of-thought approach organizes textual reasoning in a step-by-step manner using a linear chain structure (Wei et al., 2023). Building on this, following works have shown that incorporating more complex organizational structures further improves reasoning performance (Besta et al., 2024; Yao et al., 2023). However, these approaches structure reasoning only at the level of connections between distinct reasoning steps

(inter-thought level) and leave the content of individual steps (thought level) unstructured. This raises the critical question: *Can LLMs’ reasoning abilities be further enhanced by introducing structure within individual thoughts?*

To address this question, we draw inspiration from cognitive neuroscience theories of human thought. Neuroscientists have found that humans think in a structured way, with the brain’s organization facilitating sequential and goal-oriented reasoning. Christoff and Gabrieli (2000) provided early evidence that the prefrontal cortex supports structured reasoning through a rostrocaudal hierarchy, enabling the processing of increasingly abstract concepts and complex goal-directed behavior. Later, Friston (2005)’s predictive coding framework demonstrated how structured cognition emerges from the brain’s ability to build hierarchical models, combining experiences with current input to predict results. More recently, Jeff Hawkins (Hawkins, 2021) argued that humans think in a structured manner, with the neocortex organizing knowledge in certain structures, and thinking arises from neurons activating sequential locations in these frames. Building on these insights, we propose investigating whether similarly structured representations can be incorporated into LLMs to enhance their reasoning and planning capabilities.

In this work, we adopt a simple yet effective structural format—a tabular schema—to approximate the structured nature of human thinking processes. In our approach, the schema of a table serves as a framework for organizing and navigating knowledge. Inspired by the sequential processes described in neuroscience—where neurons activate specific patterns step by step (Hawkins, 2021)—we model these processes as the sequential population of rows in a table, moving across columns according to a predefined schema. A single table can encapsulate one or more such structured thought processes, providing a coherent con-

tainer for organizing and connecting thinking steps and associated information. Tables not only represent step-by-step processes for achieving specific goals but also serve as robust frameworks for planning tasks. Moreover, utilizing tables as structured representations enables schema design that ensures organization and data integrity, thereby facilitating efficient verification and analysis.

The contributions of our paper are as follows:

- Motivated by insights from cognitive neuroscience regarding the structured nature of human thinking, we propose a novel framework, Table as Thought, that injects structure at the thought level. To the best of our knowledge, this is the first exploration and demonstration of integrating structured representations directly into the reasoning process of large language models.
- We demonstrate the advantages of Table as Thought in tasks requiring planning, highlighting its potential to enhance performance on tasks that demand sequential and goal-oriented thought processes.
- We provide a detailed and comprehensive analysis of Table as Thought, offering insights into its functionality and strengths, and comparing the benefits of structured versus unstructured thought representations. We hope these findings inspire future research into the nature and representation of thought processes in artificial intelligence and computational linguistics.

2 Related Work

Structures in LLM Reasoning. Recent advancements in large language models (LLMs) have increasingly focused on integrating structured processes to enhance reasoning capabilities. Chain-of-Thought prompting (Wei et al., 2023) introduces a step-by-step framework that organizes thoughts in a sequential manner, enabling more coherent reasoning. Building on this, Tree of Thoughts (Yao et al., 2023) and Graph of Thoughts (Besta et al., 2024) employ hierarchical and networked structures to further enhance problem-solving, leveraging branching and interconnected paths. Moreover, self-consistency (Wang et al., 2023) improves reliability by sampling multiple reasoning paths and selecting the most consistent outcome, thereby addressing variability in generated responses.

While these methods excel at organizing reasoning at a macro level—such as through chain-

ing, branching, or aggregating thought paths—they do not address the internal structure of individual thoughts. Our work is distinct in that it introduces structure directly at the thought level, refining the granularity of reasoning processes in LLMs. By focusing on the internal organization of individual reasoning steps, we provide a novel perspective on enhancing the depth and precision of structured reasoning in LLMs.

Representations of Tables in LLM Inference.

Tables have traditionally played a significant role in LLMs for tasks involving the understanding and processing of tabular data, such as knowledge retrieval (Cong et al., 2024), question answering over structured data (Yin et al., 2020; Zhang et al., 2024b), and tabular reasoning (Herzig et al., 2020; Deng et al., 2024). In these tasks, tables are leveraged only as input for interpretation and manipulation.

The Chain-of-Table framework (Wang et al., 2024) extends the application of tables by employing them as proxies for intermediate thoughts in reasoning tasks involving tabular data. In this framework, LLMs iteratively update a table, forming a dynamic reasoning chain where the table evolves based on intermediate results. While this approach has proven effective on tabular-specific datasets, it remains inherently tied to tasks where tabular data is part of the input or reasoning context.

In contrast, our work redefines the role of tables by utilizing them as a universal framework for structuring and representing the internal thought processes of LLMs in non-table-specific tasks, such as planning and mathematical reasoning. Unlike prior approaches that depend on pre-existing tabular inputs, we employ tables as dynamic containers to organize and manipulate thoughts step by step. This approach enables structured reasoning even in tasks where no tabular data is initially present, bridging the gap between unstructured text-based reasoning and structured problem-solving paradigms. By generalizing the utility of tables beyond table-specific reasoning tasks, our work marks a significant departure from previous methods and demonstrates the versatility of this novel framework.

3 Table as Thought

We present the design of the Table as Thought framework, which introduces a novel approach to reasoning in large language models by leveraging tables as structured representations of

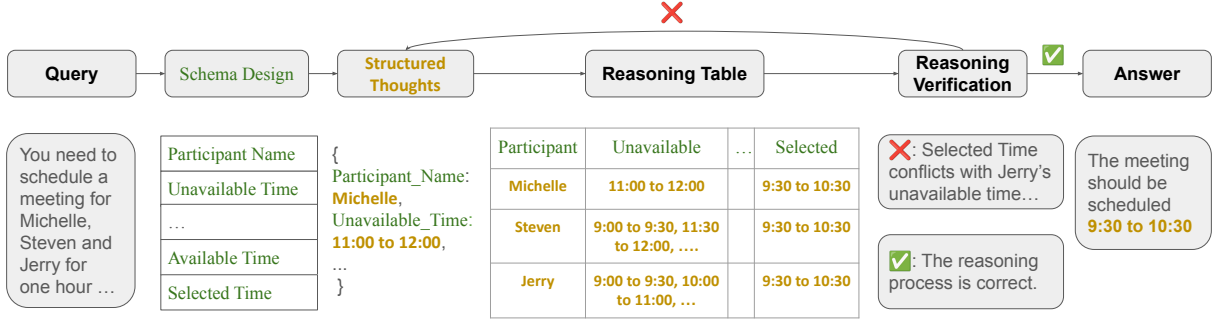


Figure 1: The Overall Pipeline for Table as Thought Reasoning. The figure illustrates how Table as Thought structures reasoning by iteratively populating a reasoning table based on the schema, verifying consistency, and updating the table until the final answer is achieved.

thoughts.

Table as Thought. Table as Thought employs a table as a container to represent one or more structured thoughts. These tables, referred to as "**reasoning tables**", encapsulate thoughts and provide a transparent representation of the reasoning process. A reasoning table T is initialized with an original table schema S , which is defined by the LLM for a given query Q . Structured thoughts Θ are then generated based on S , with each thought corresponding to a row in the reasoning table T . The table T is subsequently populated and updated according to these structured thoughts Θ .

The overall reasoning workflow using the reasoning table is illustrated in Figure 1 and formalized in Algorithm 1.

Algorithm 1 Table as Thought

Require: Query Q

Ensure: A table T that satisfies Q

- 1: $S \leftarrow \text{DESIGNSCHEMA}(Q)$ // Define table schema
 - 2: Initialize an empty table T with schema S .
 - 3: **while** not $\text{SUFFICIENT}(T, Q)$ **do**
 - 4: $\Theta \leftarrow \text{REFLECT}(T, Q)$ // Generate possible updates
 - 5: $T \leftarrow \text{UPDATETABLE}(T, \Theta)$ // Apply updates if needed
 - 6: **end while**
 - 7: **return** T
-

Schema Development Module. The Schema Development Module dynamically adapts table schemas to accommodate various queries across different reasoning tasks. For constraint-planning tasks, where the primary objective is to satisfy

constraints, we prompt LLMs to identify the constraints explicitly before designing the schema. This ensures that both explicit and implicit constraints are addressed in the reasoning process. For mathematical reasoning tasks, the schema is tailored to reflect the logical progression of the reasoning steps, enabling systematic organization of critical information.

The headers in the table schemas are designed to represent essential reasoning steps and key information pertinent to the task. These headers act as anchors for organizing and verifying intermediate and final reasoning outputs.

For example, consider the travel planning query:

I plan to travel alone, and my planned budget for the trip is around \$1,100.

In this case, a key constraint is that the total cost should not exceed \$1,400. To address this constraint, the schema must include a header such as Cost, with the type Number, ensuring that the relevant information is captured and evaluated against the budgetary constraint.

For a mathematical reasoning task, such as a question from the GSM8K dataset:

A robe takes 2 bolts of blue fiber and half that much white fiber. How many bolts in total does it take?

Here, the reasoning process requires consideration of the quantities of blue and white fibers. The schema should therefore include keys such as Blue Fiber and White Fiber, ensuring that all relevant elements are systematically tracked and calculated.

Reasoning Verification Module. Our preliminary experiments reveal that existing LLMs may fail to generate a consistent reasoning path. Therefore, we introduce a verification module to verify the completeness and correctness of the reasoning process.

For constrained reasoning tasks, such a module verifies whether the constraints identified in the schema development phase are satisfied. Constraint checking is typically performed by the LLM through reflective reasoning. The structured nature of thoughts in Table as Thought brings a natural benefit: **Auto-Check Constraints**, which are constraints set that can be externally verified. In Table as Thought, Auto-Check Constraints facilitate the systematic validation of intermediate steps and final outputs.

For math reasoning tasks, such a module ensures that the table reflects an accurate and logically correct reasoning path toward solving the problem. This involves checking whether the intermediate and final outputs align with the expected reasoning steps outlined in the schema.

Table Construction Module. The Table Construction Module iteratively generates structured thoughts and constructs the reasoning table by incorporating the schema and feedback from the reasoning verification module. This process involves dynamically adding new thoughts to the table, modifying existing entries, or removing entries that do not align with the schema or query requirements.

The iterative process terminates under one of the following conditions:

1. The reasoning table is verified as complete and correct by the reasoning verification module.
2. The maximum number of iterations, which we set empirically as 10 in all our experiments, is reached.

4 Experiments

4.1 Tasks and Language Models

For all tasks, we adopt the original evaluation methods to ensure consistency and comparability.

Constraint Planning Tasks. The goal of constraint planning tasks is to generate plans that satisfy both explicit and implicit constraints. We evaluate our approach on two datasets, each presenting different levels of complexity in the expected plans.

The TravelPlanner dataset (Xie et al., 2024) requires LLMs to generate detailed travel plans that adhere to explicit constraints provided in the query, such as budget limitations, as well as implicit constraints derived from common sense. The expected travel plans are highly complex, encompassing multi-day agendas that include transportation, accommodations, and daily attractions. Due to the exceptionally long context required for this task, which results in substantial token costs, we conduct experiments exclusively with GPT-4-o-mini. The calendar scheduling task from the NaturalPlan benchmark (Zheng et al., 2024) focuses on generating single-object plans. In this task, LLMs must determine an appropriate meeting time based on explicit constraints, such as the company’s working hours and the unavailable time slots of each participant.

Math Reasoning Tasks. We evaluate LLMs using GSM-8K and MATH500 to assess structured mathematical reasoning. GSM-8K (Cobbe et al., 2021) contains 8,000 grade-school-level word problems, testing multi-step reasoning and numerical precision. MATH500 (Lightman et al., 2023) features 500 advanced problems from the MATH dataset (Hendrycks et al., 2021), covering algebra, calculus, and geometry. It challenges models with tasks requiring symbolic manipulation and deep mathematical understanding. These datasets help evaluate our approach across diverse scenarios, from simple arithmetic to complex problems.

Language Models. The schema design and table construction modules in Table as Thought require LLMs capable of generating complex, structured outputs that conform to intricate schemas. This capability is natively supported by OpenAI’s Structured Outputs Mode, which allows for precise alignment with defined schema requirements. Consequently, our experiments are conducted exclusively on OpenAI’s GPT-4-o-mini and GPT-4-o-2024-08-06 models (OpenAI et al., 2024). Expanding the evaluation to include open-source models with similar capabilities remains an area for future work.

4.2 Text Thought Baselines

Direct Prompting. Direct Prompting involves solving queries by directly generating an answer from the input, without prompting for any intermediate reasoning steps.

CoT Prompting. Chain-of-Thought (CoT) Prompting organizes reasoning as a sequential chain of thoughts.

Text as Thought. This approach differs from Table as Thought only in its use of unstructured representations for thoughts. **Text as Thought** employs text as the medium for reasoning. This method extends CoT prompting by iteratively updating the reasoning process based on reflection. Each iteration involves generating intermediate reasoning steps, reflecting on their correctness, and refining the reasoning path as needed. The streamlined process is formalized in Algorithm 2.

Algorithm 2 Text as Thought

Require: Query Q

Ensure: A text T that satisfies Q

- 1: Initialize an empty text T .
 - 2: **while** not SUFFICIENT(T, Q) **do**
 - 3: $\Theta \leftarrow \text{REFLECT}(T, Q)$ // Generate possible updates
 - 4: $T \leftarrow \text{UPDATETEXT}(T, \Theta)$ // Apply updates if needed
 - 5: **end while**
 - 6: **return** T
-

4.3 Variations of Table as Thought

To fully explore and understand the boundaries of Table as Thought, we introduce two variations to the TravelPlanner task. These variations include Table as Thought with auto check constraint, which adds complexity to schema design, and Table as Thought with given schema, which simplifies the task by providing a predefined schema.

Table as Thought with Auto-Check Constraint. This variation builds on the vanilla Table as Thought by requiring the LLM to add additional constraints during schema design to ensure data integrity and reflect the constraints present in the query. For instance, if a TravelPlanner query includes budget constraints, the LLM is expected to design a schema with headers like Cost and enforce a rule ensuring that the sum of the column does not exceed the specified budget. By introducing this variation, we aim to explore the boundaries of LLMs in designing complex reasoning structures and handling intricate schema requirements.

Table as Thought with Given Schema. In this variation, the LLM is provided with a predefined

schema, as shown in Table 7, rather than designing the schema independently. The given schema is derived from the evaluation pipeline of the TravelPlanner task (Xie et al., 2024), where answers are processed into tables following this schema before evaluation. This variation serves as a comparative baseline to assess the effectiveness and adaptability of schemas designed by LLMs compared to fixed, predefined schemas.

5 Results

5.1 Calendar Scheduling Task

Table as Thought achieves the highest performance among all prompting methods on the Calendar Scheduling Task, as shown in Table 2. On GPT-4o, Table as Thought improves performance by 10.8% over Direct Prompting and achieves a 5.4% improvement compared to the Text as Thought baseline. This highlights the advantage of using tables as structured representations for planning over unstructured text-based representations. A similar trend is observed with GPT-4o-mini, where Table as Thought outperforms other methods, suggesting the robustness of table-based reasoning for simpler constraint reasoning tasks like calendar scheduling.

For GPT-4o, the improvement from Direct Prompting to CoT Prompting is minimal (0.5%). In contrast, incorporating self-verification through Text as Thought yields a 4.9% improvement. When transitioning from unstructured thoughts to structured tables, there is a substantial performance boost (5.4%), underscoring the benefits of structured representations in reasoning tasks.

For GPT-4o-mini, CoT Prompting achieves a moderate 2.2% improvement over Direct Prompting, but Text as Thought fails to provide any additional gains. In contrast, Table as Thought demonstrates a significant 4.4% improvement over CoT Prompting, demonstrating the effectiveness of introducing structure at the thought level over chain-like structures at the reasoning level.

5.2 TravelPlanner Task

Table 1 shows that Table as Thought with a given schema achieves the best performance on metrics for commonsense and hard constraint in the TravelPlanner task. The results reveal an important trend: on a challenging task like TravelPlanner, which demands complex reasoning, introducing increasingly sophisticated structures into the reasoning process can lead to performance degrada-

Metric	Direct	CoT	Text as Thought	Table as Thought		
				Vanilla	w/ Auto-Check constraint	w/ Given Schema
Delivery Rate (%)	100.0	100.0	100.0	100.0	99.4	100.0
Commonsense Constraint Micro Pass Rate (%)	68.3	69.0	68.3	64.4	63.8	70.1
Commonsense Constraint Macro Pass Rate (%)	2.22	2.22	0.556	0.0	0.0	3.33
Hard Constraint Micro Pass Rate (%)	7.62	6.19	3.81	3.33	1.90	5.95
Hard Constraint Macro Pass Rate (%)	4.44	4.44	2.78	1.67	0.556	5.00
Final Pass Rate (%)	0.556	0.556	0.0	0.0	0.0	1.11

Table 1: Evaluation results for different models and prompt methods on TraverPlanner Tasks on GPT4o-mini

	Direct	CoT	Text as Thought	Table as Thought
GPT-4o	64.0	64.5	69.4	74.8
GPT-4o-mini	36.2	38.4	38.4	42.3

Table 2: Performance of GPT-4o and GPT-4o-mini models under different prompting methods for calendar scheduling.

	Direct	CoT	Text as Thought	Table as Thought
MATH500				
GPT-4o	75.0	72.2	72.6	64.2
GPT-4o-mini	65.4	65.2	63.4	47.8
GSM8K				
GPT-4o	95.4	95.9	95.7	94.1
GPT-4o-mini	93.9	93.6	92.9	92.4

Table 3: Performance of GPT-4o and GPT-4o-mini models under different prompting methods for MATH500 and GSM8K.

	Direct	CoT	Text as Thought
MATH500			
GPT-4o	4.4/25.0	5.4/27.8	4.4/27.4
GPT-4o-mini	2.0/36.6	2.4/34.6	2.8/34.8
GSM8K			
GPT-4o	1.59/4.62	1.29/4.09	1.60/4.33
GPT-4o-mini	1.59/6.14	2.12/6.37	2.50/7.13

Table 4: The Percentage of Questions that Table as Thought successfully work out while other prompting methods failed vs failed rate of other prompting methods.

tion. Specifically, methods that incorporate additional complexity—such as chain-of-thought (CoT) prompting, self-reflection in Text as Thought, and rule-constrained structured thoughts in Table as Thought with Auto-Check constraint—tend to perform worse compared to simpler approaches. The exception is Table as Thought with a given schema, which avoids this degradation by relieving the LLM of the need to design its own schema, allowing it to focus solely on reasoning within a predefined structure.

5.3 Math Reasoning Tasks

Table 3 highlights a general trend in the MATH500 and GSM8K tasks: introducing additional complexity into the reasoning process often leads to a performance drop, particularly for GPT-4o-mini. For instance, on MATH500, the performance of both GPT-4o and GPT-4o-mini decreases as the reasoning structures become more complicated, from Direct Prompting to Text as Thought to Table as Thought. This effect is especially pronounced for GPT-4o-mini, where the performance of Table as Thought falls to 47.8%, compared to 65.4% with Direct Prompting. A similar trend is observed on GSM8K, where the addition of more structured reasoning methods results in marginal performance degradation. These results suggest that LLMs may already be overfitted to math reasoning tasks, as noted in recent studies (Mirzadeh et al., 2024; Zhang et al., 2024a).

Despite this general trend, Table as Thought demonstrates its potential to improve performance by successfully solving questions that text-thought-based methods fail to address, particularly with more capable models like GPT-4o. Table 4 provides a detailed breakdown of the percentage of questions that Table as Thought solves, which were missed by other methods. On MATH500, Table as Thought resolves approximately 20% of such questions, while on GSM8K, this figure exceeds 30%. These findings underscore the utility of structured reasoning in identifying alternative pathways to solutions that text-based reasoning methods may overlook.

6 Analysis

6.1 Effect of Schema Design

Schema design plays a pivotal role in structuring the reasoning paths of Calendar Scheduling tasks. Different schemas determine the granularity of the reasoning process, which in turn affects model performance.

Schema Example	
One Row	Time Slot, Jesse Availability Kathryn Availability, Megan Availability All Participants Available, Earliest Availability
Multi Row	Participant Name, Availability Start Time Availability End Time, Meeting Duration Work Hours Constraint, Schedule Constraint Preference Constraint, Proposed Meeting Time

Table 5: Schema examples for Multi Row Thought and One Row Thought.

	GPT-4o-mini	GPT-4o
One Row	45.05	72.93
Multi Row	43.46	80.28

Table 6: Performance Comparison of Multi Row and One Row Schemas for GPT-4o-mini and GPT-4o on Calendar Scheduling.

Table 5 shows that in the **one-row schema**, the reasoning process is concise: the LLM identifies all available time slots for participants in a single step and selects a suitable meeting time. This schema produces a single-row table, encapsulating the reasoning process in a compact form. In contrast, the **multi-row schema** divides the process into finer-grained steps. The LLM first extracts unavailable and preferred time slots for each participant. It then computes available time slots before aggregating this information to finalize the meeting time. This approach results in a table with multiple rows, each representing an intermediate reasoning step, and provides a more detailed reasoning path.

In Table 6, for GPT-4o, the multi-row schema outperforms the one-row schema, achieving 80.28% accuracy compared to 72.93%. In contrast, GPT-4o-mini performs better with the simpler one-row schema (45.05% vs. 43.46% for the multi-row schema). This highlights that schema complexity impacts performance differently for the two models.

6.2 LLM Struggles to Design Effective Schema for Complex Planning

Unlike Calendar Scheduling, which focuses on selecting a single time slot, TravelPlanner involves generating a comprehensive travel itinerary, which is much more complex. Our findings indicate that tasking the LLM with designing a table schema results in a notable performance drop compared to using direct prompting with a pre-defined schema.

Schema Example	
Given Schema	days, current_city, attraction, transportation, breakfast, lunch, dinner, accommodation
LLM Developed Schema	Day, Date, Location, Transportation Details, Accommodation Details, Activities/Attractions, Dining Options, Estimated Cost, Notes/Preferences

Table 7: Given Schema and Example of GPT-4o developed Schema.

Schema Designing	Reasoning Verification	ACC(%)
✓	✓	42.3
✓	×	38.5 (↓ 3.8)
×	✓	36.2 (↓ 6.1)
×	×	32.7 (↓ 9.6)

Table 8: Ablation study results for GPT-4o-mini with schema designing and reasoning verification effects on performance of calendar scheduling.

This suggests that the insufficient capability of LLM in designing table schemas may hinder its performance on complex planning tasks.

Although the provided schema is not perfect—omitting some critical columns, such as "cost" for budget constraints—it is generally more effective than most LLM-designed schemas. For instance, as shown in Table 7, the LLM-developed schema and the given schema are structurally similar. However, a key difference is the use of "Dining Options" in the LLM-designed schema, as opposed to separating dining into specific categories like "breakfast," "lunch," and "dinner." In practice, this simplification often leads the LLM to allocate only a single meal per day, which contradicts common-sense expectations for travel planning.

6.3 Ablation Study

We conducted an ablation study using GPT-4o-mini on the Calendar Scheduling task to evaluate the individual contributions of schema design and reasoning verification. Table 8 shows that when reasoning verification is removed, accuracy drops from 42.3% to 38.5% (↓ 3.8%). This indicates that without explicitly verifying constraints, the LLM may overlook key restrictions in the query, leading to false positives during self-checking. The absence of schema design leads to a larger performance drop, from 42.3% to 36.2% (↓ 6.1%), and further to 32.7% (↓ 9.6%) when both schema design and reasoning verification are removed. This highlights

Column Headers	
w/o Schema Design	Participant, Available Time Slots, Selected Meeting Time
w/ Schema Design	Participant Name, Participant Availability, Meeting Duration, Meeting Day, Proposed Meeting Time, Work Hours Start, Work Hours End, Conflict Check, Final Meeting Time , Notes/Comments

Table 9: Example of Column Headers of Table Thoughts w/wo Schema Design.

the critical role of schema design in structuring the reasoning process. Table 9 shows that without a schema, the LLM tends to create tables with fewer columns, omitting key information necessary for constraint checking. While the table without schema design contains basic headers such as Participant and Selected Meeting Time, the schema-designed table includes additional headers like Conflict Check, Work Hours Start/End, and Notes/Comments. These additional columns capture critical reasoning steps and constraints, enabling more effective verification and selection of a valid meeting time.

7 Conclusion

We proposed Table as Thought, a novel framework that introduces structured reasoning at the thought level. The framework centers on the design and utilization of table schemas, where the LLM is tasked with constructing a schema and generating structured thoughts based on it. Our results demonstrate that Table as Thought excels in constraint planning tasks, showcasing its ability to manage complex constraints effectively. Moreover, the framework exhibits significant potential for further improving performance in math reasoning tasks, particularly in addressing unsolved problems through structured reasoning.

Additionally, we conducted detailed analyses of the results, exploring the interplay between schema design, reasoning complexity, and model capabilities. These insights pave the way for future research into the nature and representation of thought processes, offering a promising direction for the development of more robust reasoning frameworks in LLMs.

Limitations

Our proposed methods are currently supported only by models capable of generating structured data

with complex schemas. This limitation restricts our experiments to a small set of closed-source models, such as those provided by OpenAI. Consequently, the generalizability of our findings to open-source LLMs remains unexplored. Future work should investigate approaches for adapting Table as Thought to a broader range of models, including those with limited native support for structured data generation.

Ethical Statement

This research was conducted using publicly available datasets (e.g., GSM-8K, MATH500, TravelPlanner) in compliance with their terms of use, ensuring no personally identifiable information (PII) was processed. While our proposed framework, Table as Thought, aims to enhance structured reasoning in LLMs, we acknowledge the potential risks of misuse in harmful applications, such as deceptive planning or adversarial reasoning. To mitigate this, we advocate for responsible deployment with appropriate safeguards.

Acknowledgement

The GPT experiments are supported by credit from OpenAI through OpenAI Researcher Access assigned to Naihao Deng. We thank Hanchen Xia and Ruiqi He for their assistance with reviewing and offering helpful feedback during the development process.

References

- Maciej Besta, Nils Blach, Ales Kubicek, Robert Gerstenberger, Lukas Gianinazzi, Joanna Gajda, Tomasz Lehmann, Michał Podstawski, Hubert Niewiadomski, Piotr Nyczyk, and Torsten Hoefer. 2024. [Graph of Thoughts: Solving Elaborate Problems with Large Language Models](#). *Proceedings of the AAAI Conference on Artificial Intelligence*, 38(16):17682–17690.
- Kalina Christoff and John D. E. Gabrieli. 2000. [The frontopolar cortex and human cognition: Evidence for a rostrocaudal hierarchical organization within the human prefrontal cortex](#). *Psychobiology*, 28:168–186.
- Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. 2021. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*.
- Tianji Cong, Madelon Hulsebos, Zhenjie Sun, Paul Groth, and H. V. Jagadish. 2024. [Observatory: Char-](#)

- acterizing embeddings of relational tables. *Preprint*, arXiv:2310.07736.
- Naihao Deng, Zhenjie Sun, Ruiqi He, Aman Sikka, Yulong Chen, Lin Ma, Yue Zhang, and Rada Mihalcea. 2024. [Tables as texts or images: Evaluating the table reasoning ability of LLMs and MLLMs](#). In *Findings of the Association for Computational Linguistics: ACL 2024*, pages 407–426, Bangkok, Thailand. Association for Computational Linguistics.
- Karl Friston. 2005. [A theory of cortical responses](#). *Philosophical Transactions of the Royal Society of London. Series B, Biological Sciences*, 360(1456):815–836.
- Jeff Hawkins. 2021. *A Thousand Brains: A New Theory of Intelligence*, first edition edition. Basic Books, Hachette Book Group, Inc., New York, NY.
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. 2021. [Measuring mathematical problem solving with the math dataset](#). *Preprint*, arXiv:2103.03874.
- Jonathan Herzig, Pawel Krzysztof Nowak, Thomas Müller, Francesco Piccinno, and Julian Eisenschlos. 2020. [TaPas: Weakly supervised table parsing via pre-training](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 4320–4333, Online. Association for Computational Linguistics.
- Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. 2023. [Let’s verify step by step](#). *Preprint*, arXiv:2305.20050.
- Iman Mirzadeh, Keivan Alizadeh, Hooman Shahrokhi, Oncel Tuzel, Samy Bengio, and Mehrdad Farajtabar. 2024. [Gsm-symbolic: Understanding the limitations of mathematical reasoning in large language models](#). *Preprint*, arXiv:2410.05229.
- OpenAI, Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, Red Avila, Igor Babuschkin, Suchir Balaji, Valerie Balcom, Paul Baltescu, Haiming Bao, Mohammad Bavarian, Jeff Belgum, Irwan Bello, Jake Berdine, Gabriel Bernadett-Shapiro, Christopher Berner, Lenny Bogdonoff, Oleg Boiko, Madelaine Boyd, Anna-Luisa Brakman, Greg Brockman, Tim Brooks, Miles Brundage, Kevin Button, Trevor Cai, Rosie Campbell, Andrew Cann, Brittany Carey, Chelsea Carlson, Rory Carmichael, Brooke Chan, Che Chang, Fotis Chantzis, Derek Chen, Sully Chen, Ruby Chen, Jason Chen, Mark Chen, Ben Chess, Chester Cho, Casey Chu, Hyung Won Chung, Dave Cummings, Jeremiah Currier, Yunxing Dai, Cory Decareaux, Thomas Degry, Noah Deutsch, Damien Deville, Arka Dhar, David Dohan, Steve Dowling, Sheila Dunning, Adrien Ecoffet, Atty Eleti, Tyna Eloundou, David Farhi, Liam Fedus, Niko Felix, Simón Posada Fishman, Juston Forte, Isabella Fulford, Leo Gao, Elie Georges, Christian Gibson, Vik Goel, Tarun Gogineni, Gabriel Goh, Rapha Gontijo-Lopes, Jonathan Gordon, Morgan Grafstein, Scott Gray, Ryan Greene, Joshua Gross, Shixiang Shane Gu, Yufei Guo, Chris Hallacy, Jesse Han, Jeff Harris, Yuchen He, Mike Heaton, Johannes Heidecke, Chris Hesse, Alan Hickey, Wade Hickey, Peter Hoeschele, Brandon Houghton, Kenny Hsu, Shengli Hu, Xin Hu, Joost Huizinga, Shantanu Jain, Shawn Jain, Joanne Jang, Angela Jiang, Roger Jiang, Haozhun Jin, Denny Jin, Shino Jomoto, Billie Jonn, Heewoo Jun, Tomer Kaftan, Łukasz Kaiser, Ali Kamali, Ingmar Kanitscheider, Nitish Shirish Keskar, Tabarak Khan, Logan Kilpatrick, Jong Wook Kim, Christina Kim, Yongjik Kim, Jan Hendrik Kirchner, Jamie Kiros, Matt Knight, Daniel Kokotajlo, Łukasz Kondraciuk, Andrew Kondrich, Aris Konstantinidis, Kyle Kopic, Gretchen Krueger, Vishal Kuo, Michael Lampe, Ikai Lan, Teddy Lee, Jan Leike, Jade Leung, Daniel Levy, Chak Ming Li, Rachel Lim, Molly Lin, Stephanie Lin, Mateusz Litwin, Theresa Lopez, Ryan Lowe, Patricia Lue, Anna Makanju, Kim Malfacini, Sam Manning, Todor Markov, Yaniv Markovski, Bianca Martin, Katie Mayer, Andrew Mayne, Bob McGrew, Scott Mayer McKinney, Christine McLeavey, Paul McMillan, Jake McNeil, David Medina, Aalok Mehta, Jacob Menick, Luke Metz, Andrey Mishchenko, Pamela Mishkin, Vinnie Monaco, Evan Morikawa, Daniel Mossing, Tong Mu, Mira Murati, Oleg Murk, David Mély, Ashvin Nair, Reiichiro Nakano, Rajeev Nayak, Arvind Neelakantan, Richard Ngo, Hyeonwoo Noh, Long Ouyang, Cullen O’Keefe, Jakub Pachocki, Alex Paino, Joe Palermo, Ashley Pantuliano, Giambattista Parascandolo, Joel Parish, Emy Parparita, Alex Passos, Mikhail Pavlov, Andrew Peng, Adam Perelman, Filipe de Avila Belbute Peres, Michael Petrov, Henrique Ponde de Oliveira Pinto, Michael, Pokorny, Michelle Pokrass, Vitchyr H. Pong, Tolly Powell, Alethea Power, Boris Power, Elizabeth Proehl, Raul Puri, Alec Radford, Jack Rae, Aditya Ramesh, Cameron Raymond, Francis Real, Kendra Rimbach, Carl Ross, Bob Rotsted, Henri Roussez, Nick Ryder, Mario Saltarelli, Ted Sanders, Shibani Santurkar, Girish Sastry, Heather Schmidt, David Schnurr, John Schulman, Daniel Selsam, Kyla Sheppard, Toki Sherbakov, Jessica Shieh, Sarah Shoker, Pranav Shyam, Szymon Sidor, Eric Sigler, Maddie Simens, Jordan Sitkin, Katarina Slama, Ian Sohl, Benjamin Sokolowsky, Yang Song, Natalie Staudacher, Felipe Petroski Such, Natalie Summers, Ilya Sutskever, Jie Tang, Nikolas Tezak, Madeleine B. Thompson, Phil Tillet, Amin Tootoonchian, Elizabeth Tseng, Preston Tuggle, Nick Turley, Jerry Tworek, Juan Felipe Cerón Uribe, Andrea Vallone, Arun Vijayvergiya, Chelsea Voss, Carroll Wainwright, Justin Jay Wang, Alvin Wang, Ben Wang, Jonathan Ward, Jason Wei, CJ Weinmann, Akila Welihinda, Peter Welinder, Jiayi Weng, Lilian Weng, Matt Wiethoff, Dave Willner, Clemens Winter, Samuel Wolrich, Hannah Wong, Lauren Workman, Sherwin Wu, Jeff Wu, Michael Wu, Kai Xiao, Tao Xu, Sarah Yoo, Kevin Yu, Qim-

ing Yuan, Wojciech Zaremba, Rowan Zellers, Chong Zhang, Marvin Zhang, Shengjia Zhao, Tianhao Zheng, Juntang Zhuang, William Zhuk, and Barret Zoph. 2024. [Gpt-4 technical report](#). *Preprint*, arXiv:2303.08774.

Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le, Ed Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. 2023. [Self-consistency improves chain of thought reasoning in language models](#). *Preprint*, arXiv:2203.11171.

Zilong Wang, Hao Zhang, Chun-Liang Li, Julian Martin Eisenschlos, Vincent Perot, Zifeng Wang, Lesly Miculicich, Yasuhisa Fujii, Jingbo Shang, Chen-Yu Lee, and Tomas Pfister. 2024. [Chain-of-table: Evolving tables in the reasoning chain for table understanding](#). *Preprint*, arXiv:2401.04398.

Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed Chi, Quoc Le, and Denny Zhou. 2023. [Chain-of-thought prompting elicits reasoning in large language models](#). *Preprint*, arXiv:2201.11903.

Jian Xie, Kai Zhang, Jiangjie Chen, Tinghui Zhu, Renze Lou, Yuandong Tian, Yanghua Xiao, and Yu Su. 2024. [Travelplanner: A benchmark for real-world planning with language agents](#). *Preprint*, arXiv:2402.01622.

Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Thomas L. Griffiths, Yuan Cao, and Karthik Narasimhan. 2023. [Tree of thoughts: Deliberate problem solving with large language models](#). *Preprint*, arXiv:2305.10601.

Pengcheng Yin, Graham Neubig, Wen-tau Yih, and Sebastian Riedel. 2020. [TaBERT: Pretraining for joint understanding of textual and tabular data](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 8413–8426, Online. Association for Computational Linguistics.

Hugh Zhang, Jeff Da, Dean Lee, Vaughn Robinson, Catherine Wu, Will Song, Tiffany Zhao, Pranav Raja, Charlotte Zhuang, Dylan Slack, Qin Lyu, Sean Hendryx, Russell Kaplan, Michele Lunati, and Summer Yue. 2024a. [A careful examination of large language model performance on grade school arithmetic](#). *Preprint*, arXiv:2405.00332.

Tianshu Zhang, Xiang Yue, Yifei Li, and Huan Sun. 2024b. [Tablellama: Towards open large generalist models for tables](#). *Preprint*, arXiv:2311.09206.

Huaixiu Steven Zheng, Swaroop Mishra, Hugh Zhang, Xinyun Chen, Minmin Chen, Azade Nova, Le Hou, Heng-Tze Cheng, Quoc V. Le, Ed H. Chi, and Denny Zhou. 2024. [Natural plan: Benchmarking llms on natural language planning](#). *Preprint*, arXiv:2406.04520.

A Completion Rate Analysis of Table as Thought on Closed and Open Source LLMs

We investigate the completion rate of Table as Thought when using both closed-source and open-source LLMs in calendar scheduling tasks. Despite advancements in open-source models such as LLaMA 3.1-8B and LLaMA 3.1-70B, these models encounter significant challenges in adhering to the tool schema required for the schema design and table construction modules. Specifically, they fail to generate the expected outputs due to the complexity of the tool schema, leading to errors in the intermediate pipeline.

In a set of 100 test examples, we observe that Table as Thought fails 100% and 63% of the time when using LLaMA 3.1-8B and LLaMA 3.1-70B, respectively. These high failure rates make experiments with open-source models infeasible because:

1. **High Computational Cost:** The need for repeated retries significantly increases computational overhead.
2. **Result Bias:** Retrying may introduce bias, as incorrect reasoning processes may be verified as correct, prematurely terminating the pipeline and artificially increasing the success rate.

In contrast, closed-source models such as GPT-4 variants demonstrate much higher success rates. The results are summarized in Table 10.

Table 10: Success Rates of Table as Thought with Different LLMs on Calendar Scheduling Tasks.

Model	Completion Rate (%)
LLaMA 3.1-8B	0
LLaMA 3.1-70B	37
GPT-4o	100
GPT-4o-mini	99

B Example of Table as Thought on TraverPlanner Task

B.1 Table as Thought w/ Given Schema

Query: Please draw up a 3-day travel itinerary for one person, beginning in Oakland and heading to Tucson from March 15th to March 17th, 2022, with a budget of \$1,400.

Reasoning table:

```
1 "day1": {
2   "days": 1,
3   "current_city": "Oakland",
4   "transportation": "Flight Number: F4002752, from Oakland to Tucson,
5     Departure Time: 15:07, Arrival Time: 17:00",
6   "breakfast": "-",
7   "attraction": "-",
8   "lunch": "-",
9   "dinner": "Pizza Street, Tucson",
10  "accommodation": "Private room with private bathroom, Tucson",
11  "total_cost": 270,
12  "id": "day1"
13 },
14 "day2": {
15   "days": 2,
16   "current_city": "Tucson",
17   "transportation": "-",
18   "breakfast": "Mocha, Tucson",
19   "attraction": "Pima Air & Space Museum, Tucson",
20   "lunch": "Pizza Street, Tucson",
21   "dinner": "Canteen Till I Die, Tucson",
22   "accommodation": "Room for rent shared bathroom, Tucson",
23   "total_cost": 61,
24   "id": "day2"
25 },
26 "day3": {
27   "days": 3,
28   "current_city": "Tucson",
29   "transportation": "Self-driving from Tucson to Oakland, Duration: 12 hours
30     42 mins, Cost: $68",
31   "breakfast": "-",
32   "attraction": "-",
33   "lunch": "-",
34   "dinner": "-",
35   "accommodation": "-",
36   "total_cost": 68,
37   "id": "day3"
38 }
```

B.2 Table as Thought w/ Self-designed Schema

Query: Please draw up a 3-day travel itinerary for one person, beginning in Oakland and heading to Tucson from March 15th to March 17th, 2022, with a budget of \$1,400.

Reasoning table:

```
1 "1": {
2   "Day": "Day 1",
3   "Date": "2022-03-15",
4   "Travel Mode": "Flight",
5   "Departure Location": "Oakland",
6   "Destination": "Tucson",
7   "Travel Duration": "1 hour 53 minutes",
8   "Travel Cost": 190,
9   "Accommodation Name": "Private room with private bathroom",
10  "Accommodation Cost": 58,
11  "Room Type": "Private room",
12  "House Rules": "No smoking",
13 }
```

```

13   "Attractions": "Pima Air & Space Museum (Cost: $15), Reid Park Zoo (Cost:
14       $10)",
15   "Dining Options": "Villa Tevere (Cost: $37), Magic Spice Wok (Cost: $31)",
16   "Total Trip Cost": 341,
17   "Budget Remaining": 1059,
18   "id": "1",
19   "Maximum Occupancy": "2 people"
20 },
21 "2": {
22   "Day": "Day 2",
23   "Date": "2022-03-16",
24   "Travel Mode": "Self-driving",
25   "Departure Location": "Tucson",
26   "Destination": "Oakland",
27   "Travel Duration": "12 hours 40 minutes",
28   "Travel Cost": 68,
29   "Accommodation Name": "Private room with private bathroom",
30   "Accommodation Cost": 58,
31   "Room Type": "Private room",
32   "House Rules": "No smoking",
33   "Attractions": "Tucson Botanical Gardens (Cost: $15), Old Tucson (Cost:
34       $20)",
35   "Dining Options": "La Plage (Cost: $93), Ooh Lala! (Cost: $70)",
36   "Total Trip Cost": 324,
37   "Budget Remaining": 735,
38   "id": "2",
39   "Maximum Occupancy": "2 people"
40 },
41 "3": {
42   "Day": "Day 3",
43   "Date": "2022-03-17",
44   "Travel Mode": "Self-driving",
45   "Departure Location": "Tucson",
46   "Destination": "Oakland",
47   "Travel Duration": "12 hours 42 minutes",
48   "Travel Cost": 68,
49   "Accommodation Name": "None",
50   "Accommodation Cost": 0,
51   "Room Type": "Private room",
52   "House Rules": "No smoking",
53   "Attractions": "Arizona-Sonora Desert Museum (Cost: $15), Children's
54       Museum Tucson (Cost: $10)",
55   "Dining Options": "Pirates of Grill (Cost: $52), Mood 4 Food (Cost: $20)",
56   "Total Trip Cost": 165,
57   "Budget Remaining": 570,
58   "id": "3",
59   "Maximum Occupancy": "N/A"
60 }

```

C Example of Table as Thought on Calendar Scheduling Task

C.1 Table as Thought w/ One Row Schema

Query: You need to schedule a meeting for Roy, Kathryn and Amy for half an hour between the work hours of 9:00 to 17:00 on Monday.

Here are the existing schedules for everyone during the day: Roy has blocked their calendar on Monday during 9:00 to 9:30, 10:00 to 10:30, 11:00 to 11:30, 12:30 to 13:00; Kathryn has meetings on Monday during 9:30 to 10:00, 16:30 to 17:00; Amy has blocked their calendar on Monday during 9:00 to 14:30, 15:00 to 16:00, 16:30 to 17:00;

Amy would rather not meet on Monday after 15:30. Find a time that works for everyone's schedule and constraints.

Raesoning table:


```

1  "1": {
2    "Meeting Duration": "30 minutes",
3    "Work Hours Start": "9:00",
4    "Work Hours End": "17:00",
5    "Roy's Unavailable Times": "9:00-9:30, 10:00-10:30, 11:00-11:30,
6      12:30-13:00",
7    "Kathryn's Unavailable Times": "9:30-10:00, 16:30-17:00",
8    "Amy's Unavailable Times": "9:00-14:30, 15:00-16:00, 16:30-17:00",
9    "Amy's Preference": "Prefers not to meet after 15:30",
10   "Proposed Meeting Time": "14:30-15:00"
11 }

```

C.2 Table as Thought w/ Multi Row Schema

Query: You need to schedule a meeting for Kathryn, Charlotte and Lauren for half an hour between the work hours of 9:00 to 17:00 on Monday.

Here are the existing schedules for everyone during the day: Kathryn has blocked their calendar on Monday during 9:00 to 9:30, 10:30 to 11:00, 11:30 to 12:00, 13:30 to 14:30, 16:30 to 17:00; Charlotte has blocked their calendar on Monday during 12:00 to 12:30, 16:00 to 16:30; Lauren has blocked their calendar on Monday during 9:00 to 10:00, 12:00 to 12:30, 13:30 to 14:30, 15:00 to 16:00, 16:30 to 17:00;

Charlotte do not want to meet on Monday after 13:30. Find a time that works for everyone's schedule and constraints.

Reasoning table:

```

1  "1": {
2    "Participant Name": "Kathryn",
3    "Unavailable Time Slots": "9:00-9:30, 10:30-11:00, 11:30-12:00,
4      13:30-14:30, 16:30-17:00",
5    "Preferred Meeting Time": "None",
6    "Meeting Duration": "30 minutes",
7    "Work Hours": "9:00-17:00",
8    "Available Time Slots": "9:30-10:30, 11:00-11:30, 12:00-13:30,
9      14:30-16:30",
10   "Constraints Satisfied": true,
11   "Proposed Meeting Time": "10:00-10:30"
12 },
13 "2": {
14   "Participant Name": "Charlotte",
15   "Unavailable Time Slots": "12:00-12:30, 16:00-16:30",
16   "Preferred Meeting Time": "Before 13:30",
17   "Meeting Duration": "30 minutes",
18   "Work Hours": "9:00-17:00",
19   "Available Time Slots": "9:00-12:00, 12:30-16:00, 16:30-17:00",
20   "Constraints Satisfied": true,
21   "Proposed Meeting Time": "10:00-10:30"
22 },
23 "3": {
24   "Participant Name": "Lauren",
25   "Unavailable Time Slots": "9:00-10:00, 12:00-12:30, 13:30-14:30,
26     15:00-16:00, 16:30-17:00",
27   "Preferred Meeting Time": "None",
28   "Meeting Duration": "30 minutes",
29   "Work Hours": "9:00-17:00",
30   "Available Time Slots": "10:00-12:00, 12:30-13:30, 14:30-15:00",
31   "Constraints Satisfied": true,
32   "Proposed Meeting Time": "10:00-10:30"
33 }

```

D Comparison of Table as Thought and Direct Prompting on GSM-8k

D.1 Example where Direct Prompting succeeds while Table as Thought Fails

Query: Stephen placed an online order for groceries. His final bill came to \$40.00. Because this was through a delivery vendor, they tacked on a 25% fee to his final total and charged him \$3.00 in delivery fees. Stephen also added a \$4.00 tip. After the extra fees, what was the final price of Stephen's groceries?

Raesoning table:

```
1 "1": {  
2   "Base_Price": 40.0,  
3   "Delivery_Fee_Percentage": 0,  
4   "Delivery_Fee_Amount": 0,  
5   "Tip_Amount": 0,  
6   "Total_Fee_Percentage": 0,  
7   "Total_Fee_Amount": 0,  
8   "Final_Price": 40.0  
9 }
```

Direct Prompting response:

To find the final price of Stephen's groceries after all the extra fees, we need to calculate each component and add them together.

1. **Original Bill:** Stephen's original grocery bill was \$40.00.
2. **Vendor Fee:** The vendor adds a 25% fee to the original bill. To calculate this fee:

$$\text{Vendor Fee} = 0.25 \times 40.00 = 10.00$$

3. **Delivery Fee:** The delivery fee is a flat \$3.00.
4. **Tip:** Stephen added a \$4.00 tip.

Now, we add all these amounts together to find the final price:

$$\text{Final Price} = \text{Original Bill} + \text{Vendor Fee} + \text{Delivery Fee} + \text{Tip}$$

$$\text{Final Price} = 40.00 + 10.00 + 3.00 + 4.00 = 57.00$$

Therefore, the final price of Stephen's groceries, after all the extra fees, is \$57.00.

D.2 Example where Table as Thought succeeds while Direct Prompting Fails

Query: Jamal's phone can hold 6 times more photographs than can Brittany's phone. The maximum number of photographs that Brittany's phone can hold is 50 times more than the number of birds in Jamal's photograph of the ducks at the zoo. If Jamal's phone can hold 1800 photographs, how many ducks can be seen in Jamal's photograph of ducks at the zoo?

Raesoning table:

```
1 "1": {  
2   "Jamal_Phone_Capacity": 1800,  
3   "Brittany_Phone_Capacity": 300,  
4   "Ducks_in_Photo": 6,  
5   "Brittany_Phone_Additional_Capacity": 50,  
6   "Equation": "300 = 50D",  
7   "Solution_Steps": "1. Jamal's phone can hold 6 times more photographs than  
   Brittany's phone. 2. Jamal's phone can hold 1800 photographs. 3. Let  
   B be the number of photographs Brittany's phone can hold. Then, 6B =  
   1800. 4. Solve for B: B = 1800 / 6 = 300. 5. Brittany's phone can hold  
   50 times more photographs than the number of ducks in Jamal's  
   photograph. 6. Let D be the number of ducks in Jamal's photograph. B =  
   50D. 7. Substitute B = 300 into the equation: 300 = 50D. 8. Solve for  
   D: D = 300 / 50 = 6."}
```

```
8 | "Final_Answer": 6
9 | }
```

Direct Prompting response:

To solve the problem, we need to determine the number of ducks in Jamal's photograph.

First, let's define the variables:

- Let B be the maximum number of photographs Brittany's phone can hold.
- Let D be the number of ducks in Jamal's photograph.

According to the problem, Jamal's phone can hold 6 times more photographs than Brittany's phone. Therefore, we have the equation:

$$J = 6B$$

where J is the number of photographs Jamal's phone can hold. We know from the problem that $J = 1800$. Substituting this into the equation gives:

$$1800 = 6B$$

To find B , divide both sides by 6:

$$B = \frac{1800}{6} = 300$$

Next, the problem states that Brittany's phone can hold 50 more photographs than the number of ducks in Jamal's photograph. This gives us the equation:

$$B = D + 50$$

We already found that $B = 300$. Substituting this into the equation gives:

$$300 = D + 50$$

To find D , subtract 50 from both sides:

$$D = 300 - 50 = 250$$

Thus, the number of ducks in Jamal's photograph is 250.