

Strategic Chain-of-Thought: Guiding Accurate Reasoning in LLMs through Strategy Elicitation

Anonymous ACL submission

Abstract

The Chain-of-Thought (CoT) paradigm has emerged as a critical approach for enhancing the reasoning capabilities of large language models (LLMs). However, despite their widespread adoption and success, CoT methods often exhibit instability due to their inability to consistently ensure the quality of generated reasoning paths, leading to sub-optimal reasoning performance. To address this challenge, we propose the Strategic Chain-of-Thought (SCoT), a novel methodology designed to refine LLM performance by integrating strategic knowledge prior to generating intermediate reasoning steps. SCoT employs a two-stage approach within a single prompt: first eliciting an effective problem-solving strategy, which is then used to guide the generation of high-quality CoT paths and final answers. Our experiments across eight challenging reasoning datasets demonstrate significant improvements, including a 21.05% increase on the GSM8K dataset and a 24.13% increase on the Tracking_Objects dataset, respectively, using the Llama3-8B model. Additionally, we extend the SCoT framework to develop a few-shot method with automatically matched demonstrations, yielding even stronger results. These findings underscore the efficacy of SCoT, highlighting its potential to substantially enhance LLM performance in complex reasoning tasks.

1 Introduction

The rapid development of large language models (LLMs) has highlighted their remarkable effectiveness in reasoning tasks (Huang and Chang, 2022; Chang et al., 2024), particularly when integrated with various prompting techniques (Sivarajkumar et al., 2023). Among these techniques, Chain-of-Thought (CoT) has become a fundamental component of contemporary LLMs and is now widely adopted in the field of natural language processing.

Despite the demonstrated effectiveness of the CoT approach in various applications, it faces sig-

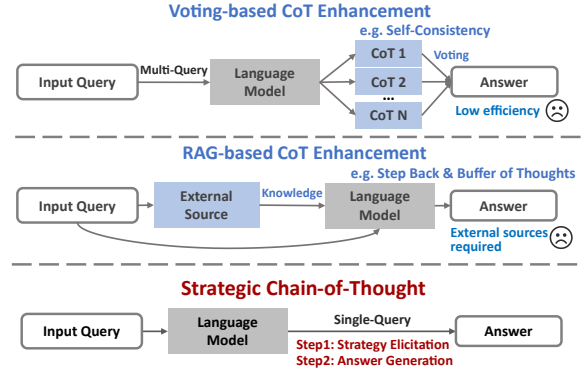


Figure 1: Comparison of some popular methods with SCoT: As a single-query method, SCoT is efficient and does not rely on external knowledge sources, distinguishing it from other approaches.

nificant challenges in complex reasoning tasks. These challenges primarily arise from the variability in the quality of the reasoning paths generated by the CoT method (Wang et al., 2022), which are not consistently of high quality. Consequently, even when LLMs produce a CoT path that aligns with a valid reasoning process, there remains a risk that the final outcome is erroneous.

This phenomenon parallels findings in cognitive science, where different problem-solving strategies, though correct, vary in error likelihood. Sweller’s Cognitive Load Theory (Sweller, 1988) suggests that different strategies impose varying cognitive loads, influencing error probability.

This variability in CoT generation strategies can reduce the reliability of CoT approaches, especially in critical applications requiring precise reasoning. Thus, further refinement is needed to enhance CoT performance in complex reasoning tasks, drawing on insights from both AI and cognitive science.

Several methods address this challenge by improving CoT path quality. Voting-based approaches improve reasoning accuracy by generating diverse paths and selecting the most reliable answer (Wang et al., 2022; Zhang et al., 2023). Retrieval-

Augmented Generation (RAG) uses multi-step prompting to access external knowledge (Lewis et al., 2021; Yang et al., 2024b; Zheng et al., 2023). Additionally, Suzgun and Kalai(2024) integrates prompt enhancement algorithms that dynamically select the optimal one during operation. Prompt-based methods guide models through predefined reasoning patterns by incorporating cue words, such as generating a plan before providing a solution (Wang et al., 2023a).

These approaches help reduce variability in path quality but often come with high costs. For instance, Self-Consistency (Wang et al., 2022) may require up to 40 queries, while methods like BoT (Yang et al., 2024b) and Step-Back (Zheng et al., 2023) involve multi-stage queries. Step-Back abstracts the problem, providing more knowledge but not directly identifying key steps for its solution. Additionally, RAG-based approaches rely on high-quality external expert resources. Prompt-based methods like Plan-and-Solve (Wang et al., 2023a), while optimizing CoT paths, still share similar limitations with traditional CoT methods and may result in suboptimal reasoning paths.

To tackle this challenge, we propose a novel approach called Strategic Chain-of-Thought (SCoT). SCoT is designed to improve the quality of CoT path generation for reasoning tasks by incorporating strategic knowledge. The method involves a two-step process within a single prompt. First, it explores and identifies various problem-solving strategies, eliciting the most effective one as the guiding strategic knowledge. Subsequently, this strategic knowledge directs the model in generating high-quality CoT paths and producing accurate final answers, ensuring a more effective reasoning process. We further extend the SCoT framework by adapting it to a few-shot method. In this approach, strategic knowledge is used to automatically select the most relevant demonstrations. These examples can be employed within both the few-shot and SCoT frameworks to further enhance the model’s reasoning capability. SCoT enhances the model’s reasoning capabilities without the need for multi-query approaches or additional knowledge sources. By eliminating the requirement for multiple queries and external knowledge integration, SCoT reduces computational overhead and operational costs, making it a more practical and resource-efficient solution.

We conducted experiments across eight reasoning datasets spanning five distinct domains: mathe-

matical reasoning, commonsense reasoning, physical reasoning, spatial reasoning, and multi-hop reasoning. The results revealed substantial improvements across various models, including a 21.05% increase in accuracy on the GSM8K dataset and a 24.13% increase on the Tracking_Objects dataset with the Llama3-8b model. These results validate the effectiveness of the SCoT approach.

The contributions of this work are summarized as follows:

- We introduce a two-stage methodology that integrates strategic knowledge, guiding the LLM to generate high-quality CoT paths by first developing a problem-solving strategy and then producing the final answer.
- We propose a method that leverages strategic knowledge to select and match relevant demonstrations, enabling the precise pairing of high-quality CoT examples.
- Our experimental results validate the effectiveness of SCoT, demonstrating promising outcomes in reasoning tasks across multiple domains.

2 Related Work

2.1 Strategic Diversity in Problem Solving

In the realm of problem-solving, there is rarely a one-size-fits-all approach. The complexity of each problem often necessitates a variety of strategies to reach an effective solution. In the fields of education and cognitive science, the phenomenon of using multiple approaches to solve problems is quite common (Sweller, 1988; Rusczyk, 2003). Similarly, researchers have found that LLMs might generate diverse solution paths for a single question, where the problem-solving strategies and answers of these methods might vary significantly (Wang and Zhou, 2024; Wang et al., 2022).

2.2 Enhancement of CoT Path

Current methods for improving model-generated content are diverse and sophisticated.

Some approaches improve reasoning by adding specific phrases to prompt templates. For example, assigning roles at the start of the prompt can elicit more professional responses (Kong et al., 2024), while phrases like “Take a breath” (Yang et al., 2023) or prompting the model to first create a plan or principle before solving the prob-

lem (Zheng et al., 2023; Wang et al., 2023a) can generate higher-quality CoT paths.

Voting-based mechanisms have gained prominence in recent research efforts. Self-Consistency (Wang et al., 2022) enhances reasoning accuracy by generating over 20 CoT paths and then selecting the most consistent answer through voting. Similarly, Step-Back (Zheng et al., 2023) enhances RAG by abstracting the question for better logical structure. Similarly, Yang et al. (2024b) developed a RAG-based method, Buffer of Thoughts, integrating external knowledge into task-specific prompt templates to generate more accurate answers.

Additionally, some methods incorporate external tools for problem-solving. PAL (Gao et al., 2023) utilizes LLMs to parse problems and generated programs as intermediate reasoning steps, delegating the execution of solutions to a runtime environment such as a Python interpreter. Tree of Thoughts (ToT) (Yao et al., 2024) introduced a tree-based reasoning structure to enhance decision-making processes and to improve reasoning capabilities. Suzgun and Kalai (2024) introduced meta-prompting, a technique that integrates existing prompt-based frameworks to enable dynamic selection of the most effective reasoning strategy.

These methods are complex, with some being task-specific and others requiring multi-turn prompting. However, they have proven effective in enhancing LLM reasoning, advancing CoT generation in machine learning.

3 Method

In this section, we introduce the strategic knowledge, the Strategic Chain-of-Thought (SCoT) method, and its few-shot extension.

3.1 Strategic Knowledge

LLMs tend to generate different CoT paths for the same problem, but their quality can vary significantly (Wang and Zhou, 2024; Wang et al., 2022). As shown in Figure 2(a), even methods like Plan-and-Solve, known for higher accuracy, can produce errors when solving problems like the math question “compute the sum of all integers s such that $-26 < s < 24$ ”. An alternative approach, using the arithmetic series sum formula, provides more stable and accurate results. While both methods are valid, the formula-based approach leads to higher-quality, more stable outputs and is consid-

ered strategic knowledge.

Strategic knowledge (strategy) refers to a well-defined method or principle that guides reasoning towards a correct and stable solution. It involves using structured processes that logically lead to the desired outcome, thereby enhancing the stability of CoT generation and improving the overall quality of the results.

Specifically, strategic knowledge should adhere to the following principles:

1. **Correct and Comprehensive Problem-Solving Approach:** The strategy provides a systematic approach that enables the model to generate accurate answers when it follows the reasoning steps.

2. **Straightforward Problem-Solving Steps:** The steps of the method should not be overly complex, while each step should be sufficiently detailed to ensure accuracy and prevent overly brief outputs that could lead to ambiguity.

3.2 Strategic Chain-of-Thought

Building on the concept of strategic knowledge, we propose a prompt-based method to enhance the reasoning quality of LLMs, called Strategic Chain-of-Thought (SCoT).

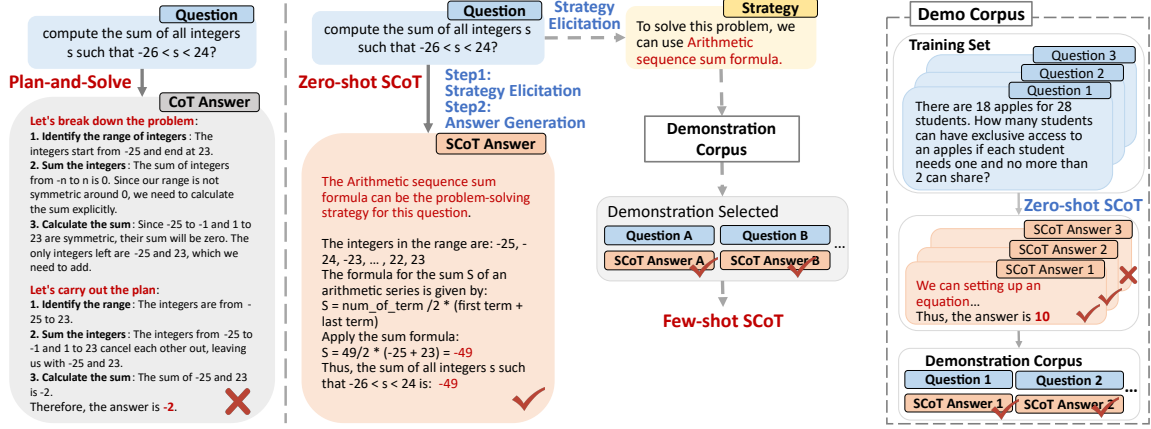
The SCoT method enables the model to first elicit strategic knowledge before generating an answer, rather than producing an answer directly. Specifically, in a single-query setting, SCoT involves two key steps:

1. **Elicitation of Strategic Knowledge:** The model identifies and determines one of the most effective methods for solving the problem, which then serves as the strategic knowledge for the task.

2. **Application of Strategic Knowledge:** The model subsequently applies the identified strategy to solve the problem and derive the final answer.

Figure 3(a) illustrates a prompt template that utilizes the SCoT approach. Our prompt consists of five components: Role, Workflow, Rule, Initialization, and Task Input. The workflow, as shown in Figure 4, comprises three steps integrated into a single prompt. The first two steps are designed to identify and elicit strategic knowledge for solving the problem, while the third step focuses on applying the strategy to generate the answer.

We demonstrate that the details of strategic knowledge identification vary across different domains. In mathematics, strategic knowledge favors generating elegant and efficient solutions, such as using the arithmetic series formula to sum sequences. In physics, it involves selecting the most



(a) Framework of zero-shot and few-shot strategic chain-of-thought. The solid line indicates zero-shot SCoT, while the dashed line indicates few-shot SCoT. Prompt details are omitted for brevity. (b) Construction of the demonstration corpus.

Figure 2: Illustration of zero-shot and few-shot strategic SCoT. Few-shot SCoT builds upon zero-shot SCoT by incorporating strategy-aligned demonstrations. Prompt details are omitted for brevity.

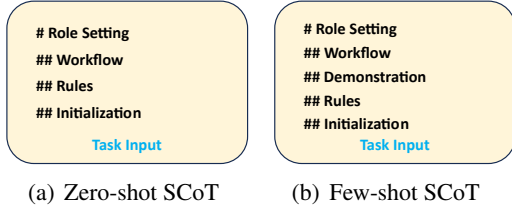


Figure 3: Prompt templates used in Zero-shot SCoT and Few-shot SCoT.

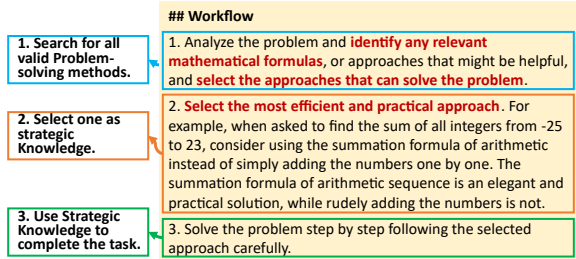


Figure 4: Workflow for a math prompt.

relevant and straightforward formulas or processes, such as applying $F = ma$ to calculate force. For multi-hop reasoning, strategic knowledge focuses on determining the appropriate granularity for problem decomposition and recalling pertinent information. Similarly, in other domains, the model first develops an overarching method or workflow before systematically applying it to solve problems.

3.3 Few-shot Strategic Chain-of-Thought

We refine the SCoT method into a few-shot version by leveraging the strategy to select demonstrations. Our approach is structured into two stages: constructing a strategy-based demonstration corpus

(Figure 2(b)) and performing model inference in a two-query process (Figure 2(a)). This is motivated by the fact that some problems, despite being from different domains, share similar solution strategies. As a result, direct similarity matching based on the problems alone may not yield the most relevant demonstrations.

Stage 1: Strategy-Based Demonstration Corpus Construction.

1. SCoT Answer Generation: We apply zero-shot SCoT to the instances in the training set to generate SCoT answers for each question.

2. Demonstration Corpus Construction: The generated answers are compared to the ground truth, retaining only correct question-SCoT pairs. This step assumes that the strategic knowledge used in these problems is both correct and relevant. The validated pairs are then compiled into a demonstration corpus categorized by strategic knowledge.

Stage 2: Model Inference..

1. Strategy Generation: The LLM generates strategy for each test instance, focusing on understanding the problem rather than giving the answer.

2. Demonstration Matching: The strategy is used to search the demonstration corpus from Stage 1, matching relevant demonstrations to the SCoT answers. This is done by computing embeddings with the m3e-base model and selecting the most similar examples from the corpus.

3. Few-shot Inference: The selected demonstrations are integrated as few-shot examples into the input prompt (Figure 3(b)). This integration guides the model to generate the final prediction based on

the provided examples.

4 Experimental Setup

4.1 Datasets and Tasks

To validate the effectiveness of SCoT, we evaluate models on a range of reasoning datasets spanning mathematical, physical, commonsense, multi-hop, and spatial reasoning. Specifically, we use MathQA (Amini et al., 2019), AQuA (Ling et al., 2017), GSM8K (Cobbe et al., 2021), and MMLU-math (Hendrycks et al., 2021) for mathematical reasoning; ARC (Clark et al., 2018) for physical reasoning; CSQA (Talmor et al., 2019) and StrategyQA (Geva et al., 2021) for commonsense and multi-hop reasoning; and Tracking_Object (BIG-bench authors, 2023) for spatial reasoning.

In the few-shot setting, SCoT is evaluated on MathQA, AQuA, GSM8K, and ARC, which are selected due to their sufficiently large training sets with gold answers to support demonstration construction. We first run zero-shot SCoT and retain only the demonstrations that yield correct answers to form the demonstration corpus. More details on this process are provided in Appendix A.2.

4.2 Models

To verify the effectiveness of the SCoT method, we utilize the Llama3 series (Dubey et al., 2024) (including Llama3-8B, Llama3-70B, Llama3.1-8B, and Llama3.1-70B); the Llama2 series (Touvron et al., 2023) (including Llama2-7B, Llama2-13B, and Llama2-70B); Mistral-7B (Jiang et al., 2023); the Qwen2 series (Yang et al., 2024a) (including Qwen2-7B and Qwen2-72B); and ChatGLM4-9B (Team GLM et al., 2024).

4.3 Baselines

We use zero-shot prompts (Kojima et al., 2022), Self-Consistency (Wang et al., 2022), Step-Back (Zheng et al., 2023), SIMTOM (Wilf et al., 2023), Plan-and-Solve (Wang et al., 2023b), Least-to-Most (Zhou et al., 2022) and Buffer-of-Thought (Yang et al., 2024b) as baselines. Step-Back is evaluated on only six datasets because we were unable to locate examples from the CSQA and Object datasets, and neither of these datasets includes a training set. BoT and Least-to-Most are specifically designed for mathematical problem solving, therefore, we conduct experiments only on the four mathematics datasets.

We use accuracy as the performance metric, calculated as the average of three independent inferences for each model. Experimental parameters are provided in the Appendix A.4.

5 Experimental Results

In this section, we empirically evaluate the effectiveness of SCoT on Llama3-8B and Mistral-7B across all datasets. To further validate its efficacy, we test three representative tasks on all seven models. We also assess the impact of model size, perform ablation studies on SCoT components, and conduct case studies, along with additional discussions to understand the factors influencing SCoT’s effectiveness.

5.1 Results across all Datasets

Results across all datasets using the two models are presented in Table 1 and Table 2. In both vanilla zero-shot and self-consistency settings, SCoT outperforms the CoT approach in most tasks, with significant improvements including a 21.05% accuracy gain on GSM8K and a 24.13% gain on Tracking_Object after incorporating strategic knowledge. However, the Llama3-8B model exhibits a 2.6% performance decline on the ARC dataset, which may be attributed to the greater reliance on factual knowledge relative to others. While some methods, such as SIMTOM, achieve comparable performance (within 2% of SCoT on Llama3-8B), their results fluctuate more across datasets, and their average gains remain lower. In general, SCoT shows an average improvement of 6.92% on all datasets using Llama3-8B, and demonstrates an average improvement of 3.83% using Mistral-7B. Notably, SCoT achieves substantial gains in mathematical and commonsense reasoning tasks, primarily because the gap between strategic and suboptimal reasoning is more significant, highlighting the value of strategy in guiding models toward more accurate and reliable solutions.

Furthermore, we extend the SCoT framework to support few-shot settings by automatically matching demonstrations, resulting in even stronger performance. The SCoT 1-shot[−], as shown in Table 1 and Table 2, refers to CoT prompting with demonstrations selected based on strategic knowledge. Compared to CoT 0-shot¹, SCoT 1-shot[−], which uses strategy-matched demonstrations, shows sig-

¹We do not present the accuracy of CoT 1-shot separately as it was comparable to CoT 0-shot in our experiments.

Table 1: Accuracy (%) using Llama3-8B across all datasets. “SCoT 1-shot[−]” refers to the results obtained using the standard Few-shot CoT template with demonstrations matched by strategy, and “+SC” refers to methods using self-consistency. The highest scores for 0-shot and 1-shot are both bolded.

	Methods	MathQA	AQuA	GSM8K	MMLU	ARC	SQA	CSQA	Object	Gain
0-shot	CoT 0-shot	55.33	49.61	52.11	46.67	80.60	64.60	71.13	44.27	−
	CoT 0-shot+SC	57.00	51.90	48.48	49.52	81.00	66.00	72.06	54.00	+1.83
	Step-Back	56.33	50.39	53.25	47.78	75.80	64.64	−	−	− 0.28
	SIMTOM	56.33	51.07	71.04	49.26	80.00	66.67	74.50	57.00	+4.94
	Plan-and-Solve	50.00	51.57	70.28	49.63	78.00	49.67	71.00	49.40	+0.53
	SCoT 0-shot	56.67	51.85	73.16	50.00	78.02	68.56	74.00	68.40	+6.92
	SCoT 0-shot+SC	60.33	53.94	70.58	52.22	78.00	69.00	75.00	61.60	+6.92
1-shot	Least-to-Most	49.00	35.43	30.18	45.19	−	−	−	−	− 6.23
	Buffer of Thoughts	44.00	38.98	52.24	40.37	−	−	−	−	− 7.28
	SCoT 1-shot [−]	56.33	50.87	74.91	−	73.40	−	−	−	+4.22
	SCoT 1-shot	57.67	55.12	76.57	−	80.60	−	−	−	+7.89

Table 2: Accuracy (%) using Mistral-7B across all datasets.

	Methods	MathQA	AQuA	GSM8K	MMLU	ARC	SQA	CSQA	Object	Gain
0-shot	CoT 0-shot	30.00	29.13	36.26	29.75	67.20	56.22	61.80	21.40	−
	CoT 0-shot+SC	31.42	32.87	34.50	31.88	68.78	53.50	62.69	24.50	+1.04
	Step-Back	31.43	32.87	36.67	31.85	68.00	56.72	−	−	+1.12
	SIMTOM	27.00	29.72	38.67	25.96	68.40	61.67	61.80	24.73	+0.77
	Plan-and-Solve	26.67	30.07	38.22	30.00	66.20	52.67	58.00	23.40	− 0.81
	SCoT 0-shot	30.44	33.60	38.97	32.35	72.20	61.89	68.00	24.75	+3.81
	SCoT 0-shot+SC	31.67	36.22	34.72	32.96	75.40	57.33	66.50	27.60	+3.83
1-shot	Least-to-Most	22.33	29.92	25.31	26.67	−	−	−	−	− 3.72
	Buffer of Thoughts	28.00	24.41	30.17	25.93	−	−	−	−	− 2.88
	SCoT 1-shot [−]	34.33	31.50	45.57	−	67.40	−	−	−	+4.05
	SCoT 1-shot	37.00	35.04	47.38	−	73.20	−	−	−	+6.26

nificant performance improvements across most datasets, highlighting the effectiveness of the matched demonstrations. The SCoT 1-shot, which combines both strategic knowledge and strategy-matched demonstrations, achieves the best results.

5.2 Results across all Models

The experimental results for all models on the three datasets are shown in Table 3. The experiments demonstrate that SCoT can enhance performance across most models. In particular, with the exception of the Llama3.1-8B model, where the addition of SCoT results in a slight decrease in accuracy on the MMLU task, other models exhibit accuracy improvements ranging from 1.11% to 24.13% across the three datasets. Notably, the CoT 0-shot has achieved 100% accuracy with Llama3.1-70B model on Tracking_Object dataset, and SCoT 0-shot maintains this performance.

5.3 Model Scale

We investigate the impact of model size on the effectiveness of SCoT, the results on the Llama2 model series with three different sizes are shown in

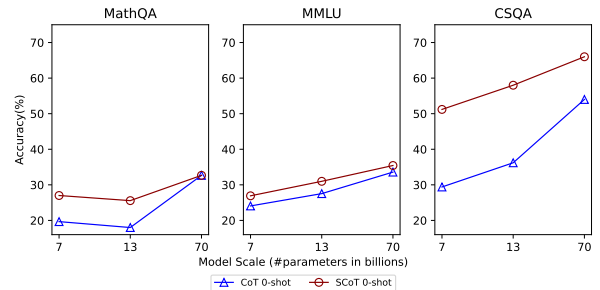


Figure 5: Accuracy (%) across three datasets using different scales of models in Llama2 series.

Figure 5. It demonstrates that SCoT can lead to accuracy improvements across all sizes of the Llama2 models. However, a general trend emerges that performance improvement decreases marginally with model size. Furthermore, manual inspection of the model outputs reveals that larger models are more likely to generate CoT paths containing strategic knowledge in 0-shot settings.

5.4 Ablation Study

We explore the effects of various components within the prompt (such as role, workflow, struc-

Table 3: Accuracy (%) across seven models on MMLU, SQA and Tracking_Object datasets

Datasets	Methods	Llama3-8B	Mistral-7B	Chatglm4-9B	Qwen2-7B	Qwen2-72B	Llama3.1-8B	Llama3.1-70B
MMLU	CoT 0-shot	46.67	29.75	66.67	71.97	84.20	59.63	85.19
	SCoT 0-shot	50.00 _{+3.33}	32.35 _{+2.59}	68.15 _{+1.48}	71.85	85.93 _{+1.73}	56.42	85.19
SQA	CoT 0-shot	64.60	56.22	61.80	61.00	75.22	73.11	64.67
	SCoT 0-shot	68.56 _{+3.96}	61.89 _{+5.67}	64.67 _{+2.87}	61.00	77.67 _{+2.45}	74.22 _{+1.11}	82.33 _{+1.33}
Object	CoT 0-shot	44.27	21.40	61.80	46.20	93.93	62.60	100.00
	SCoT 0-shot	68.40 _{+24.13}	24.67 _{+3.27}	69.00 _{+7.20}	47.53 _{+1.33}	97.47 _{+3.54}	77.60 _{+15.00}	100.00

Table 4: Ablation study on SCoT prompt components: * denotes a non-markdown format, while no * indicates a markdown format.

Methods	AQuA	ARC
Mistral-7B*	29.13%	67.20%
Mistral-7B + Role*	27.95%	69.80%
Mistral-7B + Role	32.28%	71.20%
Mistral-7B + Workflow*	33.07%	70.40%
Mistral-7B + Workflow	31.89%	70.40%
SCoT 0-shot (Ours)	33.60%	72.20%
SCoT 1-shot (Ours)	35.04%	73.20%
SCoT 3-shot (Ours)	35.43%	73.20%

Table 5: Token length comparison between CoT and SCoT (0-shot) across three datasets

Datasets	Methods	Llama3-8B	Mistral-7B
AQuA	CoT 0-shot	158.335	186.889
	SCoT 0-shot	226.192	209.441
GSM8K	CoT 0-shot	130.532	858.507
	SCoT 0-shot	206.278	611.848
Object	CoT 0-shot	106.931	83.475
	SCoT 0-shot	123.482	128.285

ture, and the quantity of demonstrations) on accuracy. The experimental results are illustrated in Table 4. Building on the CoT 0-shot approach, we observed that adding roles, incorporating workflows, and formatting prompts in Markdown progressively increased accuracy. We also explored the impact of the number of demonstrations on accuracy within the few-shot SCoT framework. Experimental results indicate that as the number of demonstrations increases, the performance of SCoT either slightly improves or remains unchanged.

5.5 Case Study

We conduct a detailed case study focusing on the validity of the strategic knowledge elicited from the model. Figure 6 shows several typical cases.

In the domain of mathematics, we observe that the SCoT output tends to favor solving problems using inequalities rather than directly analyzing the problem to reach an answer. For the instance of frog jumping calculation in the Figure 6, an incorrect solution may miscalculate the final jump’s impact. While generating a strategy ensures accurate calculations by considering all constraints and systematically solving the problem.

In physical tasks, we find that the model’s CoT output could be misled by specific phrases in the task input (*e.g.*, “capacitor”), leading to the selection of an incorrect formula. In contrast, the SCoT approach successfully elicited the correct formula.

Similarly, in multi-hop reasoning tasks, CoT output often focuses on details, resulting in incomplete subsequent logical reasoning, whereas SCoT generates answers by considering the overall context.

5.6 Efficiency Analysis

Due to SCoT’s mechanism of generating a strategy before solving problems in one query, it is more efficient than multi-query methods. However, compared to single-query methods, the output token length might be longer, potentially decreasing efficiency. To investigate this, we measure the token lengths for the AQuA, GSM8K, and Tracking_Object using both CoT 0-shot and SCoT 0-shot methods. The results are shown in Table 5.

Except for Mistral-7B on GSM8K, SCoT outputs are longer than CoT, with lengths ranging from 1.03 to 1.8 times that of CoT and averaging around 1.5 times. This is mainly due to the models’ tendency to repeat answer spans, especially in complex or multi-step tasks—a pattern observed in both CoT and SCoT but more pronounced in SCoT. In general, whenever repetitive answer spans occur, the model is unlikely to produce the correct result. For practical deployment, we recommend truncating outputs or employing early-stopping mechanisms when repetition is detected, which can effectively mitigate unnecessary computational overhead while preserving response quality.

Math Question A	A well is 30 meters deep. A frog at the bottom jumps 3 meters up each day but slips back 2 meters each night. How many days will it take for the frog to climb out of the well?
CoT Answer	To solve this problem, we need to consider the frog's daily progress . On the first day, the frog jumps up 3 meters but slips back 2 meters, resulting in a net gain of 1 meter. From the second day onward, the frog gains 1 meter each day. The frog needs to climb a total of 30 meters to get out of the well. After the first day, it still needs to climb 29 meters. Therefore, the total number of days required is 1+29=30 days .
SCoT Answer	To address this problem elegantly, we can utilize a mathematical model . The frog gains a net height of 1 meter each day. However, on the day it reaches the top of the well, it will no longer slip back down. Let n represent the number of days it takes for the frog to climb out of the well. We can formulate the following inequality : $n - 1 + 3 \geq 30$ This simplifies to: $n \geq 28$ Thus, the total number of days for the frog to climb out of the well is 28 days .
Physical Question B	How much work is required to charge a $10\mu\text{F}$ capacitor to a potential difference of 100V?
CoT Answer	Using the formula for work done by a capacitor, which is $W=k \times V \times C$, and substituting the values, we find that 0.09J of work needs to be done.
SCoT Answer	Using the capacitor charging formula. The formula for capacitor charging is $E = (1/2)CV^2$ To substitute the values, it needs to do 0.05J of work
Multi-hop Question C	Is shrimp scampi definitely free of plastic?
CoT Answer	Shrimp scampi is a dish made with shrimp, butter, garlic, and sometimes lemon, usually served with pasta. The ingredients typically include fresh or frozen shrimp, butter, garlic, and lemon juice. There is no connection between shrimp scampi and plastic, as plastic is not an ingredient in this dish . Therefore, shrimp scampi is True definitely free of plastic.
SCoT Answer	Shrimp scampi usually does not contain plastic, but contamination is possible. It is typically made with ingredients like shrimp, garlic, butter, and parsley, which do not contain plastic, and there are no known instances of shrimp scampi being made with plastic-containing ingredients . Therefore, while shrimp scampi is usually made without plastic. So, the answer is: No .

Figure 6: Comparison of the paths generated by CoT and SCoT in different domains using the Qwen2-72B model.

Table 6: Accuracy (%) using DeepSeek-R1-Distill-Qwen-32B.

Methods	MathQA	MMLU	ARC	Object
CoT 0-shot	55.67	85.56	95.00	99.70
SCoT 0-shot	84.33	93.70	94.40	99.80

6 Discussions

6.1 Automatic SCoT

To verify that the performance gains stem from the SCoT concept rather than human-crafted prompts, we conduct a preliminary experiment to assess whether SCoT prompt templates can be automatically generated. We instruct LLMs to generate prompt templates based solely on the SCoT concept and evaluate them on the AQUA dataset. The accuracy using Mistral-7B model is 31.89%. Although the accuracy of prompts automatically generated based on the SCoT concept is lower than that of manually crafted SCoT prompts, it is still superior to 0-shot CoT performance. This suggests that the automatic generation of SCoT-based prompts are feasible.

6.2 Enhancing Smaller-Scale Models

In this paper, we focus on enhancing the reasoning capabilities of smaller-scale models. We exclude larger, more powerful models from our experiments because they already achieve accuracy rates exceeding 95% on our datasets, even in a zero-shot CoT configuration. This indicates that the capabilities of models extremely large parameters on these tasks are already saturated. In future work, we aim to

test SCoT on more challenging reasoning datasets to further validate its efficacy on stronger models.

6.3 Accuracy of the DeepSeek Distilled Model

We also experimented with the popular DeepSeek-R1 series models (DeepSeek-AI et al., 2025), specifically using CoT 0-shot and SCoT 0-shot on the DeepSeek-R1-Distill-Qwen-32B model across four datasets, as shown in Table 6. Except for a slight drop in performance on the ARC dataset, we observed positive gains on the other three datasets, with the highest improvement of 28.66% on MathQA. These results are consistent with those in Table 1 and Table 2, demonstrating that SCoT is effective on the DeepSeek series models.

7 Conclusion

In this paper, we introduce the Strategic Chain-of-Thought, a method that enables large models to autonomously generate an optimal Chain-of-Thought path. By integrating a structured workflow for eliciting and applying strategic knowledge, SCoT enhances the model’s ability to produce high quality outputs. We further extend SCoT to a few-shot version by matching demonstrations through strategic knowledge from a predefined strategic knowledge-based corpus. Experimental results demonstrate the effectiveness of both 0-shot and few-shot SCoT.

Overall, SCoT offers a promising framework for improving the quality of reasoning path in large models. Future research will focus on evaluating its effectiveness with more complex problems and exploring further applications.

8 Limitation

In this paper, we propose a novel methodology to enhance the performance of large language models by incorporating strategic knowledge prior to generating intermediate reasoning steps. However, there are four limitations in our work.

The first limitation is that we only use one-shot SCoT in our main experiments. This choice was made because incorporating additional demonstrations would significantly increase input length, leading to substantial resource consumption. As a result, we limited our experiments to a single dataset and model without conducting large-scale testing.

The second limitation is the exclusion of several recent methods from our baseline comparisons. For example, approaches like Tree-of-Thought or meta prompting could not be reproduced due to the limitation of the resources.

The third limitation is that our experiments focused solely on reasoning tasks, leaving us uncertain about SCoT’s effectiveness in other domains. Additionally, we have not provided a theoretical proof to explain why SCoT is effective.

The fourth limitation lies in our use of only the most basic approach to automatic prompt generation, without applying any prompt optimization techniques. As a result, the performance of the automatically generated prompts presented in this paper is inferior to that of manually crafted prompts. Our work provides only a preliminary attempt at automatic prompt generation, intended primarily to demonstrate that the SCoT framework remains effective even when using automatically generated prompts. In future work, we plan to explore methods for optimizing these prompts, with the goal of developing a more generalizable and robust automatic prompt generation framework.

We plan to address these limitations in future work by expanding the experimental scope and refining our methodology.

9 Ethical Considerations

All datasets and models used in this paper are open-source, and the licenses for the models have been specified. The prompts used in the experiments are provided, and the entire study can be reproduced using widely available large model API frameworks. This ensures the reproducibility and transparency of the research.

References

- Aida Amini, Saadia Gabriel, Peter Lin, Rik Koncel-Kedziorski, Yejin Choi, and Hannaneh Hajishirzi. 2019. [Mathqa: Towards interpretable math word problem solving with operation-based formalisms](#). *Preprint*, arXiv:1905.13319.
- BIG-bench authors. 2023. [Beyond the imitation game: Quantifying and extrapolating the capabilities of language models](#). *Transactions on Machine Learning Research*.
- Yupeng Chang, Xu Wang, Jindong Wang, Yuan Wu, Linyi Yang, Kaijie Zhu, Hao Chen, Xiaoyuan Yi, Cunxiang Wang, Yidong Wang, et al. 2024. A survey on evaluation of large language models. *ACM Transactions on Intelligent Systems and Technology*, 15(3):1–45.
- Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. 2018. [Think you have solved question answering? try arc, the ai2 reasoning challenge](#). *Preprint*, arXiv:1803.05457.
- 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](#). *Preprint*, arXiv:2110.14168.
- DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jiashi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaojun Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiusi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng Zhou, Shuting Pan, S. S. Li, Shuang Zhou, Shaoqing Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun, T. Wang, Wangding Zeng, Wanbiao Zhao, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li,

664	Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin,	<i>Language Technologies (Volume 1: Long Papers)</i> ,	721
665	Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxi-	pages 4099–4113, Mexico City, Mexico. Association	722
666	ang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang,	for Computational Linguistics.	723
667	Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang		
668	Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng	Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying	724
669	Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi,	Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E.	725
670	Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang,	Gonzalez, Hao Zhang, and Ion Stoica. 2023. Effi-	726
671	Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo,	cient memory management for large language model	727
672	Yuan Ou, Yudian Wang, Yue Gong, Yuheng Zou, Yu-	serving with pagedattention. In <i>Proceedings of the</i>	728
673	jia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You,	<i>ACM SIGOPS 29th Symposium on Operating Systems</i>	729
674	Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu,	<i>Principles</i> .	730
675	Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu,		
676	Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan,	Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio	731
677	Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean	Petroni, Vladimir Karpukhin, Naman Goyal, Hein-	732
678	Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao,	rich Küttler, Mike Lewis, Wen tau Yih, Tim Rock-	733
679	Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zi-	täschel, Sebastian Riedel, and Douwe Kiela. 2021.	734
680	jia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song,	<i>Retrieval-augmented generation for knowledge-</i>	735
681	Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu	<i>intensive nlp tasks</i> . <i>Preprint</i> , arXiv:2005.11401.	736
682	Zhang, and Zhen Zhang. 2025. <i>Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning</i> . <i>Preprint</i> , arXiv:2501.12948.		
683			
684			
685	Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey,	Wang Ling, Dani Yogatama, Chris Dyer, and Phil Blun-	737
686	Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman,	som. 2017. <i>Program induction by rationale genera-</i>	738
687	and et al. 2024. <i>The llama 3 herd of models</i> . <i>Preprint</i> ,	<i>tion : Learning to solve and explain algebraic word</i>	739
688	arXiv:2407.21783.	<i>problems</i> . <i>Preprint</i> , arXiv:1705.04146.	740
689			
690	Luyu Gao, Aman Madaan, Shuyan Zhou, Uri Alon,	Richard Rusczyk. 2003. <i>The Art of Problem Solving</i> .	741
691	Pengfei Liu, Yiming Yang, Jamie Callan, and Gram	Mathematical Association of America, Washington,	742
692	Neubig. 2023. Pal: Program-aided language	D.C.	743
693	models. In <i>International Conference on Machine</i>		
694	<i>Learning</i> , pages 10764–10799. PMLR.	Sonish Sivarajkumar, Mark Kelley, Alyssa Samolyk-	744
695		Mazzanti, Shyam Visweswaran, and Yanshan Wang.	745
696	Mor Geva, Daniel Khashabi, Elad Segal, Tushar Khot,	2023. <i>An empirical evaluation of prompting</i>	746
697	Dan Roth, and Jonathan Berant. 2021. <i>Did aris-</i>	<i>strategies for large language models in zero-shot</i>	747
698	<i>totle use a laptop? a question answering bench-</i>	<i>clinical natural language processing</i> . <i>Preprint</i> ,	748
699	<i>mark with implicit reasoning strategies</i> . <i>Preprint</i> ,	arXiv:2309.08008.	749
700	arXiv:2101.02235.		
701		Mirac Suzgun and Adam Tauman Kalai. 2024.	750
702	Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou,	Meta-prompting: Enhancing language models	751
703	Mantas Mazeika, Dawn Song, and Jacob Steinhardt.	with task-agnostic scaffolding. <i>arXiv preprint</i>	752
704	2021. <i>Measuring massive multitask language under-</i>	<i>arXiv:2401.12954</i> .	753
705	<i>standing</i> . <i>Preprint</i> , arXiv:2009.03300.		
706		John Sweller. 1988. <i>Cognitive load during problem</i>	754
707	Jie Huang and Kevin Chen-Chuan Chang. 2022. To-	<i>solving: Effects on learning</i> . <i>Cognitive Science</i> ,	755
708	wards reasoning in large language models: A survey.	12(2):257–285.	756
709	<i>arXiv preprint arXiv:2212.10403</i> .		
710		Alon Talmor, Jonathan Herzig, Nicholas Lourie, and	757
711	Albert Q. Jiang, Alexandre Sablayrolles, Arthur Men-	Jonathan Berant. 2019. <i>CommonsenseQA: A ques-</i>	758
712	sch, Chris Bamford, Devendra Singh Chaplot, Diego	<i>tion answering challenge targeting commonsense</i>	759
713	de las Casas, and et al. 2023. <i>Mistral 7b</i> . <i>Preprint</i> ,	<i>knowledge</i> . In <i>Proceedings of the 2019 Conference</i>	760
714	arXiv:2310.06825.	<i>of the North American Chapter of the Association for</i>	761
715		<i>Computational Linguistics: Human Language Tech-</i>	762
716	Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yu-	<i>nologies, Volume 1 (Long and Short Papers)</i> , pages	763
717	taka Matsuo, and Yusuke Iwasawa. 2022. Large lan-	4149–4158, Minneapolis, Minnesota. Association for	764
718	guage models are zero-shot reasoners. <i>Advances in</i>	<i>Computational Linguistics</i> .	765
719	<i>neural information processing systems</i> , 35:22199–		
720	22213.	Team GLM, Aohan Zeng, Bin Xu, Bowen Wang, Chen-	766
721		hui Zhang, Da Yin, and et al. 2024. <i>Chatglm: A</i>	767
722	Aobo Kong, Shiwan Zhao, Hao Chen, Qicheng Li, Yong	<i>family of large language models from glm-130b to</i>	768
723	Qin, Ruiqi Sun, Xin Zhou, Enzhi Wang, and Xiao-	<i>glm-4 all tools</i> . <i>Preprint</i> , arXiv:2406.12793.	769
724	hang Dong. 2024. <i>Better zero-shot reasoning with</i>		
725	<i>role-play prompting</i> . In <i>Proceedings of the 2024</i>	Hugo Touvron, Louis Martin, Kevin Stone, Peter Al-	770
726	<i>Conference of the North American Chapter of the</i>	bert, Amjad Almahairi, Yasmine Babaei, and et al.	771
727	<i>Association for Computational Linguistics: Human</i>	2023. <i>Llama 2: Open foundation and fine-tuned chat</i>	772
728		<i>models</i> . <i>Preprint</i> , arXiv:2307.09288.	773

774	Lei Wang, Wanyu Xu, Yihuai Lan, Zhiqiang Hu,	Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei,	829
775	Yunshi Lan, Roy Ka-Wei Lee, and Ee-Peng Lim.	Nathan Scales, Xuezhi Wang, Dale Schuurmans,	830
776	2023a. Plan-and-solve prompting: Improving zero-	Claire Cui, Olivier Bousquet, Quoc Le, et al. 2022.	831
777	shot chain-of-thought reasoning by large language	Least-to-most prompting enables complex reason-	832
778	models . In <i>Proceedings of the 61st Annual Meet-</i>	ing in large language models. <i>arXiv preprint</i>	833
779	<i>ing of the Association for Computational Linguistics</i>	<i>arXiv:2205.10625</i> .	834
780	(Volume 1: Long Papers), pages 2609–2634, Toronto,		
781	Canada. Association for Computational Linguistics.		
782	Lei Wang, Wanyu Xu, Yihuai Lan, Zhiqiang Hu,		
783	Yunshi Lan, Roy Ka-Wei Lee, and Ee-Peng Lim.		
784	2023b. Plan-and-solve prompting: Improving zero-		
785	shot chain-of-thought reasoning by large language		
786	models. <i>arXiv preprint arXiv:2305.04091</i> .		
787	Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le,		
788	Ed Chi, Sharan Narang, Aakanksha Chowdhery, and		
789	Denny Zhou. 2022. Self-consistency improves chain		
790	of thought reasoning in language models. <i>arXiv</i>		
791	<i>preprint arXiv:2203.11171</i> .		
792	Xuezhi Wang and Denny Zhou. 2024. Chain-of-		
793	thought reasoning without prompting. <i>arXiv preprint</i>		
794	<i>arXiv:2402.10200</i> .		
795	Alex Wilf, Sihyun Shawn Lee, Paul Pu Liang, and Louis-		
796	Philippe Morency. 2023. Think twice: Perspective-		
797	taking improves large language models’ theory-of-		
798	mind capabilities. <i>arXiv preprint arXiv:2311.10227</i> .		
799	An Yang, Baosong Yang, Binyuan Hui, Bo Zheng,		
800	Bowen Yu, Chang Zhou, and et al. 2024a. Qwen2		
801	technical report. <i>arXiv preprint arXiv:2407.10671</i> .		
802	Chengrun Yang, Xuezhi Wang, Yifeng Lu, Hanxiao		
803	Liu, Quoc V. Le, Denny Zhou, and Xinyun Chen.		
804	2023. Large language models as optimizers . <i>ArXiv</i> ,		
805	abs/2309.03409 .		
806	Ling Yang, Zhaochen Yu, Bin Cui, and Mengdi		
807	Wang. 2025. Reasonflux: Hierarchical llm reason-		
808	ing via scaling thought templates. <i>arXiv preprint</i>		
809	<i>arXiv:2502.06772</i> .		
810	Ling Yang, Zhaochen Yu, Tianjun Zhang, Shiyi Cao,		
811	Minkai Xu, Wentao Zhang, Joseph E Gonzalez,		
812	and Bin Cui. 2024b. Buffer of thoughts: Thought-		
813	augmented reasoning with large language models.		
814	<i>arXiv preprint arXiv:2406.04271</i> .		
815	Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran,		
816	Tom Griffiths, Yuan Cao, and Karthik Narasimhan.		
817	2024. Tree of thoughts: Deliberate problem solving		
818	with large language models. <i>Advances in Neural</i>		
819	<i>Information Processing Systems</i> , 36.		
820	Zhuosheng Zhang, Aston Zhang, Mu Li, Hai Zhao,		
821	George Karypis, and Alex Smola. 2023. Multi-		
822	modal chain-of-thought reasoning in language mod-		
823	els. <i>arXiv preprint arXiv:2302.00923</i> .		
824	Huaixiu Steven Zheng, Swaroop Mishra, Xinyun Chen,		
825	Heng-Tze Cheng, Ed H Chi, Quoc V Le, and Denny		
826	Zhou. 2023. Take a step back: Evoking reasoning via		
827	abstraction in large language models. <i>arXiv preprint</i>		
828	<i>arXiv:2310.06117</i> .		

A Details of Experiments

A.1 Models Details

This experiment involves eleven models, nine of which are public (Llama3-8B, Llama2-7B, Mistral-7B, Llama3.1-8B, Qwen2-7B, ChatGLM4-9B, Llama3-70B, Llama3.1-70B, Llama2-70B, Qwen2-72B and Qwen2.5-72B). ChatGLM4-9B is chat-oriented and other models are instruction-tuned. The sources and licenses for all public models are detailed in Table 7.

A.2 Datasets Details

This experiment involves eight datasets: MathQA, AQuA, GSM8K, MMLU, ARC, StrategyQA, CommonsenseQA, and Tracking_Object.

1. Mathematics and Physical Reasoning: We assess the models using datasets such as MathQA, AQuA, GSM8K, and MMLU-high-school-math for mathematical reasoning tasks. These datasets feature a range of mathematical problems with varying levels of difficulty, demanding strong mathematical reasoning abilities. Additionally, we evaluated the models on ARC_Challenge for physical reasoning, *i.e.*, a popular dataset that presents significant challenges in this domain.

2. Commonsense and Multi-hop Reasoning: We evaluate the models on CommonsenseQA (CSQA) for commonsense reasoning tasks and StrategyQA (SQA) for multi-hop reasoning tasks. These datasets are well-regarded in their respective domains and offer a substantial level of difficulty.

3. Spatial Reasoning: We also evaluate the models using the Tracking_Object (Object) dataset, which represents a less common but highly intriguing type of reasoning task.

Due to the large size of the MathQA, StrategyQA, ARC, CommonsenseQA, and Tracking_Object datasets. We randomly selected a subset from each to serve as the test dataset, while for the other datasets, we used the original full dataset. All datasets used in this study are publicly available, with their sources, licenses and size detailed in Table 8.

MathQA, AQuA, MMLU, ARC, StrategyQA, CommonsenseQA, and Tracking_Object consist of multiple-choice questions. To determine correctness, we compare the predicted choice with the gold choice. For GSM8K, the answers are numerical text spans; we assess correctness by checking if the predicted answer exactly matches the gold answer.

As mentioned in Section 3, we exclusively conducted experiments on the MathQA, AQuA, GSM8K, and ARC datasets, as these were the only four datasets in our study that contained training sets. Moreover, the demonstration corpus for the few-shot version of SCoT required these training sets to provide demonstrations that differed from those in the test set and aligned appropriately with it. Consequently, we limited our experiments to these four datasets.

Initially, we executed the zero-shot version of SCoT using the training sets and subsequently evaluated the final results as the gold answers are provided in the training sets. We assumed that if the model correctly answered a question, the strategic knowledge it generated was also accurate. Therefore, we retained only the correct demonstrations to construct the demonstration corpus for the few-shot version of SCoT. This systematic approach ensured the integrity of our evaluation and enhanced the overall effectiveness of the few-shot learning framework.

A.3 Baseline Details

Step-Back Prompting: We reproduced Step-Back using the demonstration examples and prompt templates provided in the original paper’s appendix.

Least-to-Most: We directly used the demonstrations from the original paper’s appendix to construct few-shot prompts. As the original work focused on math datasets, our experiments were limited to this domain. Results show that Least-to-Most is not particularly effective on more difficult problems.

Buffer-of-Thought: For Buffer-of-Thought, we found that the authors later open-sourced nearly 500 templates in a follow-up work, ReasonFlux (Yang et al., 2025). In our experiments, we utilized the knowledge and demonstration components from the template library provided in ReasonFlux, and reproduced the method based on the descriptions presented in the paper. We only conducted experiments on the mathematic datasets because BoT is focus on the mathematic domain.

SIMTOM For SIMTOM prompting, we followed the approach proposed in Wilf et al. and directly used the prompt examples provided in the original paper as our prompts

A.4 Other Details

We used the standard zero-shot CoT and few-shot CoT templates, with the Step-Back template fol-

Table 7: Models, sources and licenses used in this work

Models	Modelsources	License
Llama2-7B-chat	https://huggingface.co/meta-llama/Llama-2-7b-chat	llama2 license
Llama2-13B	https://huggingface.co/meta-llama/Llama-2-13b-chat	llama2 license
Llama2-70B	https://huggingface.co/meta-llama/Llama-2-70b-chat	llama2 license
Llama3-8B	https://huggingface.co/meta-llama/Meta-Llama-3-8B-Instruct	llama3 license
Llama3.1-8B	https://huggingface.co/meta-llama/Meta-Llama-3.1-8B-Instruct	llama3.1 license
Llama3.1-70B	https://huggingface.co/meta-llama/Meta-Llama-3.1-70B-Instruct	llama3.1 license
Mistral-7B	https://huggingface.co/mistralai/Mistral-7B-Instruct-v0.2	tongyi-qianwen license
Qwen2-7B	https://huggingface.co/Qwen/Qwen2-7B-Instruct	tongyi-qianwen license
Qwen2-72B	https://huggingface.co/Qwen/Qwen2-72B-Instruct	Apache License 2.0
Qwen2.5-72B	https://huggingface.co/Qwen/Qwen2.5-72B-Instruct	Qwen license
ChatGLM4-9b	https://huggingface.co/THUDM/glm-4-9b-chat	glm-4-9b License
DeepSeek-R1-Distill-Qwen-32B	https://huggingface.co/deepseek-ai/DeepSeek-R1-Distill-Qwen-32B	MIT License

Table 8: Datasets, sources and licenses used in this work

Datasets	Sources	Licenses	Size
MathQA	https://huggingface.co/datasets/datafreak/MathQA	Apache License 2.0	300
AQuA	https://github.com/google-deepmind/AQuA	Apache License 2.0	254
GSM8K	https://huggingface.co/datasets/openai/gsm8k	MIT License	1319
MMLU	https://huggingface.co/datasets/cais/mmlu	MIT License	300
ARC	https://huggingface.co/datasets/allenai/ai2_arc	CC-BY-SA-4.0 License	270
StrategyQA	https://huggingface.co/datasets/ChilleD/StrategyQA/viewer/default/test	MIT License	500
CommonsenseQA	https://huggingface.co/datasets/tau/commonsense_qa	MIT License	500
Object Tracking	https://github.com/google/BIG-bench	Apache License 2.0	200

lowing the design from the original paper, and the parameters for self-consistency also based on the original work. Additionally, the selection of the normal one-shot demonstration was done using embedding and cosine similarity, similar to the approach in RAG.

For all experiments except those involving Self-Consistency, we set the temperature to 0 and top_p to 1. For DeepSeek models, the temperature is set to 0.6. Thus, except for tasks with dynamic demonstration selection, the results are expected to be deterministic. For Self-Consistency, following the settings from the original paper (Wang et al., 2022), the temperature is adjusted to 0.5, and top_p is set to 0.5. 20 responses are selected for voting with -SC method.

We utilize vllm (Kwon et al., 2023) as the infer-

ence framework for all deployments. For models under 70 billion parameters (such as Llama3-8B, Llama2-7B, and Mistral-7B), we deploy each on a single 32GB AI graphics card. For models with 70 billion parameters or more (including Llama3-70B, Llama3.1-70B, Llama2-70B, and Qwen2-72B), we utilize two 80GB AI graphics cards per model.

B Results

B.1 All Results

Accuracy is used as the evaluation metric. We conducted three independent inference runs for all experiments and calculated the average results. However, due to the high computational cost, we performed only a single inference for Self-Consistency. The accuracy and standard deviation results are presented in Table 9 and Table 10.

Table 9: Accuracy (%) using Llama2-8B and Mistral-7B across all datasets. SCoT 1-shot[−] refers to the results obtained using the standard few-shot CoT template but with demonstrations matched by strategy.

Models	Methods	MathQA	AQuA	GSM8K	MMLU	ARC	SQA	CSQA	Object
Llama3-8B	CoT 0-shot	56.33 $\pm$ 0.000	49.61 $\pm$ 1.790	52.11 $\pm$ 0.129	46.67 $\pm$ 0.000	80.60 $\pm$ 0.000	64.60 $\pm$ 0.646	71.13 $\pm$ 0.094	44.27 $\pm$ 0.736
	CoT 0-shot+SC	57.00	51.90	48.48	49.52	81.00	66.00	72.06	54.00
	Step-Back	56.33 $\pm$ 0.272	50.39 $\pm$ 0.000	53.23 $\pm$ 0.000	47.78 $\pm$ 0.000	75.80 $\pm$ 0.248	64.64 $\pm$ 0.2722	—	—
	SIMTOM	56.33 $\pm$ 0.000	51.07 $\pm$ 0.000	71.04 $\pm$ 0.000	49.26 $\pm$ 0.000	80.00 $\pm$ 0.000	66.67 $\pm$ 0.000	74.50 $\pm$ 0.000	57.00 $\pm$ 0.000
	Plan-and-Solve	50.00 $\pm$ 0.000	51.57 $\pm$ 0.000	70.28 $\pm$ 0.000	49.63 $\pm$ 0.000	78.00 $\pm$ 0.000	49.67 $\pm$ 0.000	71.00 $\pm$ 0.000	49.40 $\pm$ 0.000
	SCoT 0-shot	56.67 $\pm$ 0.000	51.85 $\pm$ 1.299	73.16 $\pm$ 0.163	50.00 $\pm$ 0.000	78.02 $\pm$ 0.000	68.56 $\pm$ 0.566	74.00 $\pm$ 0.000	68.40 $\pm$ 0.000
	SCoT 0-shot+SC	60.33	53.94	70.58	52.22	78.00	69.00	75.00	61.60
	Least-to-Most	49.00 $\pm$ 0.000	35.43 $\pm$ 0.000	30.18 $\pm$ 0.000	45.19 $\pm$ 0.000	—	—	—	—
	BoT	44.00 $\pm$ 0.000	38.98 $\pm$ 0.000	52.24 $\pm$ 0.000	40.37 $\pm$ 0.000	—	—	—	—
	SCoT 1-shot [−]	56.33 $\pm$ 0.000	50.87 $\pm$ 2.140	74.91 $\pm$ 0.000	—	73.40 $\pm$ 0.000	—	—	—
	SCoT 1-shot	57.67 $\pm$ 0.000	55.12 $\pm$ 0.000	76.57 $\pm$ 0.000	—	80.60 $\pm$ 0.000	—	—	—
Mistral-7B	CoT 0-shot	30.00 $\pm$ 0.000	29.13 $\pm$ 1.245	36.26 $\pm$ 1.854	29.75 $\pm$ 0.924	67.20 $\pm$ 0.356	56.22 $\pm$ 0.314	61.80 $\pm$ 0.000	21.40 $\pm$ 0.000
	CoT 0-shot+SC	31.42	32.87	34.50	31.88	68.78	53.50	62.69	24.50
	Step-Back	31.43 $\pm$ 0.000	32.87 $\pm$ 0.322	36.67 $\pm$ 0.000	31.85 $\pm$ 0.495	68.00 $\pm$ 0.000	56.72 $\pm$ 0.000	—	—
	SIMTOM	27.00 $\pm$ 0.000	29.72 $\pm$ 0.000	38.67 $\pm$ 0.000	25.96 $\pm$ 0.000	68.40 $\pm$ 0.000	61.67 $\pm$ 0.000	61.80 $\pm$ 0.000	24.73 $\pm$ 0.000
	Plan-and-Solve	26.67 $\pm$ 0.000	30.07 $\pm$ 0.000	38.22 $\pm$ 0.000	30.00 $\pm$ 0.000	66.20 $\pm$ 0.000	52.67 $\pm$ 0.000	58.00 $\pm$ 0.000	23.40 $\pm$ 0.000
	SCoT 0-shot	30.44 $\pm$ 0.874	33.60 $\pm$ 1.523	38.97 $\pm$ 0.655	32.35 $\pm$ 1.665	72.20 $\pm$ 0.370	61.89 $\pm$ 0.415	68.00 $\pm$ 0.000	24.75 $\pm$ 0.165
	SCoT 0-shot+SC	31.67	36.22	34.72	32.96	75.40	57.33	66.50	27.60
	Least-to-Most	22.33 $\pm$ 0.000	29.92 $\pm$ 0.000	25.31 $\pm$ 0.000	26.67 $\pm$ 0.000	—	—	—	—
	Buffer of Thoughts	28.00 $\pm$ 0.000	24.41 $\pm$ 0.000	30.17 $\pm$ 0.000	25.93 $\pm$ 0.000	—	—	—	—
	SCoT 1-shot [−]	34.33 $\pm$ 0.000	31.50 $\pm$ 0.964	45.57 $\pm$ 1.087	—	67.40 $\pm$ 0.000	—	—	—
	SCoT 1-shot	37.00 $\pm$ 0.000	35.04 $\pm$ 0.000	47.38 $\pm$ 0.107	—	73.20 $\pm$ 0.000	—	—	—

Table 10: Accuracy (%) across seven models on MMLU, SQA and Tracking_Object datasets

Dataset	Methods	Llama3-8B	Mistral-7B	Chatglm4-9B	Qwen2-7B	Qwen2-70B	Llama3.1-8B	Llama3.1-70B
MMLU	CoT 0-shot	46.67 $\pm$ 0.000	29.75 $\pm$ 0.924	66.67 $\pm$ 0.302	71.97 $\pm$ 0.349	84.20 $\pm$ 0.349	59.63 $\pm$ 0.000	85.19 $\pm$ 0.605
	Math SCoT 0-shot	50.00 $\pm$ 0.000	32.35 $\pm$ 1.665	68.15 $\pm$ 0.907	71.85 $\pm$ 0.302	85.93 $\pm$ 0.302	56.42 $\pm$ 0.175	85.19 $\pm$ 0.000
SQA	CoT 0-shot	64.60 $\pm$ 0.595	56.22 $\pm$ 0.314	61.80 $\pm$ 0.363	61.00 $\pm$ 0.000	75.22 $\pm$ 0.314	73.11 $\pm$ 0.314	64.67 $\pm$ 0.000
	SCoT 0-shot	68.56 $\pm$ 0.566	61.89 $\pm$ 0.415	64.67 $\pm$ 0.408	61.00 $\pm$ 0.157	77.67 $\pm$ 0.272	74.22 $\pm$ 0.157	82.33 $\pm$ 0.157
Object	CoT 0-shot	44.27 $\pm$ 0.736	21.40 $\pm$ 0.000	61.80 $\pm$ 0.000	46.20 $\pm$ 0.000	93.93 $\pm$ 0.525	62.60 $\pm$ 0.411	100.00 $\pm$ 0.000
	SCoT 0-shot	68.40 $\pm$ 0.000	24.67 $\pm$ 0.000	69.00 $\pm$ 0.000	47.53 $\pm$ 0.094	97.47 $\pm$ 0.339	77.60 $\pm$ 0.993	100.00 $\pm$ 0.000

B.2 Case Study

We conducted a detailed case study to assess the validity of the strategic knowledge elicited from the model. Figures 7 and 8 present several representative cases spanning math reasoning, physical reasoning, commonsense reasoning, multi-hop reasoning, and spatial reasoning.

C Comparison with Step-Back prompting

We have presented the results of Step-Back in Table 1 and Table 2. Experimental results show that Step-Back Prompting performs comparably to zero-shot CoT, with some outcomes even falling below baseline.

This method is motivated by the hypothesis that abstracting the original question may help the model focus less on superficial details and more on relevant knowledge. It typically involves generating a principle (in STEM tasks) or a Step-Back

question (in other tasks), then retrieving content based on both the original and abstracted versions. However, the retrieved knowledge or generated principles often do not correspond to strategic knowledge and frequently fail to guide effective reasoning.

In STEM tasks, the abstraction often lacks problem-solving utility. For example, in the question "The price of a product is reduced by 30%. By what percentage should it be increased to return to 100%?", a strategic solution involves setting up an equation. In contrast, Step-Back produces flawed reasoning: "The price is reduced to 70%, so increasing it by 30% restores it to 100%." This overlooks the non-linearity of percentage changes.

In commonsense and multi-hop tasks, the abstracted questions may diverge from the original intent. For instance, the question "Are more people today related to Genghis Khan than Julius Cae-

Math Question A	A well is 30 meters deep. A frog at the bottom jumps 3 meters up each day but slips back 2 meters each night. How many days will it take for the frog to climb out of the well?
CoT Answer	To solve this problem, we need to consider the frog's daily progress . On the first day, the frog jumps up 3 meters but slips back 2 meters, resulting in a net gain of 1 meter. From the second day onward, the frog gains 1 meter each day. The frog needs to climb a total of 30 meters to get out of the well. After the first day, it still needs to climb 29 meters. Therefore, the total number of days required is 1+29=30 days .
SCoT Answer	To address this problem elegantly, we can utilize a mathematical model . The frog gains a net height of 1 meter each day. However, on the day it reaches the top of the well, it will no longer slip back down. Let n represent the number of days it takes for the frog to climb out of the well. We can formulate the following inequality : $n - 1 + 3 \geq 30$ This simplifies to: $n \geq 28$. Thus, the total number of days for the frog to climb out of the well is 28 days .
Physical Question B	How much work is required to charge a 10 μ F capacitor to a potential difference of 100V?
CoT Answer	Using the formula for work done by a capacitor, which is $W=k \times V \times C$, and substituting the values, we find that 0.09J of work needs to be done.
SCoT Answer	Using the capacitor charging formula. The formula for capacitor charging is $E = (1/2)CV^2$ To substitute the values, it needs to do 0.05J of work
Math Question C	Jordan plays video games for 2 hours every day. He also has a part-time job where he earns \$10 an hour. How much money would Jordan earn in one week if he spent his video game time working instead?
CoT Answer	To calculate how much money Jordan would earn in one week if he spent his video game time working instead, we need to follow these steps: Step 1: Jordan plays video games for 2 hours every day. There are 7 days in a week, so he would work for $7 - 2 = 5$ hours each day. Step 2: Jordan earns 10 an hour. If he works 5 hours a day, his daily earnings would be 5 hours * 10/hour = 50. To find his weekly earnings, we multiply his daily earnings by the number of days in a week: 50/day * 7 days/week = 350/week .
SCoT Answer	Calculate the total hours he would have worked in a week by multiplying the number of hours he plays per day by the number of days in a week. Given: 1. Hours Jordan plays video games per day = 2 hours; 2. Days in a week = 7 days Total hours = 2 hours/day * 7 days/week = 14 hours/week Calculate the total amount of money Jordan would have earned in one week if he worked those 14 hours. Total earnings = 14 hours/week * \$10/hour = \$140/week

Figure 7: Comparison of the paths generated by CoT and SCoT in different domains.

sar?" is abstracted as "Who are the ancestors of a large number of people today?", which fails to support direct comparison. Similarly, abstracting "Do the anchors on Rede Globo speak Chinese?" into "What languages are commonly spoken by Rede Globo anchors?" may retrieve irrelevant content (e.g., Portuguese), potentially misleading the model.

Overall, while abstraction occasionally aids reasoning, it often introduces noise or misalignment. In contrast, SCoT explicitly guides the model with strategy, yielding better performance across tasks.

Moreover, Step-Back involves multi-step prompting—two steps for STEM and three for others—resulting in higher latency. It is also less effective with smaller models, which frequently fail to generate valid knowledge in the first step, leading to cascading errors in subsequent steps.

The motivation behind SCoT is to guide models in selecting the most appropriate CoT reasoning path without relying on external inputs, such as user feedback. As outlined in the introduction, existing methods either introduce unnecessary complexity or depend on external resources. In contrast, SCoT offers a more efficient, lightweight approach that generates higher-quality reasoning paths.

It is important to note that the improvements observed in certain models are not always significant. There are three primary reasons for this. First, we did not provide additional external information to the models; instead, we relied solely on the prompt to guide the model in generating strategic knowledge, leveraging its inherent capabilities. This ap-

proach makes it challenging to achieve substantial improvements without sacrificing the model's efficiency in completing tasks.

Second, even when applying the SCoT framework—prompting the model to first generate strategic knowledge before producing an answer—the specifics of the prompt can significantly impact the model's performance. For instance, not all models are compatible with Markdown syntax in prompts. Since we did not optimize the prompts for each model, the current results may not reflect the best possible outcomes.

Finally, the third reason may lie in the inherent limitations of the models themselves. Further analysis of the results revealed that while many of the strategic knowledge outputs are correct and the reasoning processes are sound, errors tend to occur in the final calculations. For example, in one task requiring the calculation of 1.14^4 , most models made calculation errors, leading to incorrect answers. This type of issue is common across all models tested.

D Experimental Prompts

The prompt for standard zero-shot Chain-of-Thought is shown in Figure 9. Prompts for zero-shot Strategic Chain-of-Thought are displayed in Figure 10 (for math reasoning), Figure 11 (for multi-hop reasoning), Figure 13 (for physical reasoning), Figure 12 (for commonsense reasoning) and Figure 14 (for spatial reasoning). Prompts for one-shot Strategic Chain-of-Thought are shown

Object Question A	Alice, Bob, and Claire are playing a game. At the start of the game, they are each holding a ball: Alice has a brown ball, Bob has a red ball, and Claire has a purple ball. As the game progresses, pairs of players trade balls. First, Bob and Claire swap balls. Then, Alice and Claire swap balls. Finally, Alice and Bob swap balls. At the end of the game, Claire has the A) brown ball; B) red ball; C) purple ball
CoT Answer	Let's break down the sequence of swaps: 1. Bob and Claire swap balls: Bob gets the purple ball, and Claire gets the red ball. 2. Alice and Claire swap balls: Alice gets the red ball (which Claire had just gotten), and Claire gets the brown ball. 3. Alice and Bob swap balls: Alice gets the red ball (which she had just given to Claire), and Bob gets the brown ball. So, at the end of the game, the final positions are: Claire has the purple ball .
SCoT Answer	Initial State: Alice has a brown ball; Bob has a red ball; Claire has a purple ball Process Exchanges: 1. Bob and Claire swap balls: Alice still has a brown ball; Bob has a purple ball; Claire has a red ball 2. Alice and Claire swap balls: Alice has a red ball; Bob has a purple ball; Claire has a brown ball 3. Alice and Bob swap balls: Alice has a purple ball; Bob has a red ball; Claire has a brown ball Determine the Answer: The correct answer is that Claire has the brown ball .
Math Question B	Find out which of the following values is the multiple of X, if it is divisible by 9 and 12? A) 36; B) 15; C) 17; D) 5; E) 7
CoT Answer	To find the multiple of a number that is divisible by both 9 and 12, we need to find the common multiple of 9 and 12. The least common multiple (LCM) of 9 and 12 can be found using the prime factorization method or the LCM table. The prime factorization of 9 is 3^2 , and the prime factorization of 12 is $2^2 * 3$. Since 3^2 is a factor of both 3^2 and $2^2 * 3$, the LCM of 9 and 12 is $3^2 * 2^2 * 3 = 3^2 * 2^2 * 3 = 3^2 * 9 = 9^2 = 81$. Now, let's check which option is a perfect square of 9. Option B, 15 , is a perfect square of 9, it is the correct answer..
SCoT Answer	The problem states that X is divisible by 9 and 12. To find the multiple of X that fits the given conditions, we need to find the least common multiple (LCM) of 9 and 12. Choose the most efficient and practical approach: the LCM of 9 and 12 is 36. Therefore, the correct answer is A) 36 .
Multi-hop Question C	Is shrimp scampi definitely free of plastic?
CoT Answer	Shrimp scampi is a dish made with shrimp, butter, garlic, and sometimes lemon, usually served with pasta. The ingredients typically include fresh or frozen shrimp, butter, garlic, and lemon juice. There is no connection between shrimp scampi and plastic, as plastic is not an ingredient in this dish. Therefore, shrimp scampi is True definitely free of plastic.
SCoT Answer	Shrimp scampi usually does not contain plastic, but contamination is possible. It is typically made with ingredients like shrimp, garlic, butter, and parsley, which do not contain plastic, and there are no known instances of shrimp scampi being made with plastic-containing ingredients. Therefore, while shrimp scampi is usually made without plastic. So, the answer is: No.

Figure 8: Comparison of the paths generated by CoT and SCoT in different domains.

in Figure 15. The prompt for one-shot Strategic Chain-of-Thought is provided in Figure 5. Finally, the prompts for automated Strategic Chain-of-Thought are shown in Figure 6. The automated SCoT prompts were generated using the Qwen2.5-72B model by given the idea of SCoT. We did not provide pseudocode directly in the paper; instead, we offered prompts, as utilizing these prompts along with standard API call codes allows for straightforward reproduction of our work.

Zero-shot CoT template	I will provide you with a math problem and 5 options. Please choose the correct option from the five provided and indicate your answer with [Answer]option[Answer], such as [Answer]C[Answer]. Please output the answer at the end in strict accordance with the output format. Problem: [Please Put Your Questions Here] Options: [Please Put Your Options Here] Answer: Let's think step by step.
------------------------	--

Figure 9: An example of prompting for standard zero-shot CoT

Zero-shot SCoT template	# Role A highly skilled mathematician and algorithm expert. # Workflow 1. Analyze the problem and identify any relevant mathematical formulas, or approaches that might be helpful, and select the approaches that can solve the problem. 2. Choose the most efficient and practical approach. For example, when asked to find the sum of all integers from -25 to 23, consider using the summation formula of arithmetic sequence instead of simply adding the numbers one by one. The summation formula of arithmetic sequence is an elegant and practical solution, while rudely adding the numbers is not. 3. Solve the problem step by step following the selected approach carefully. ## Rules 1. Avoid using brute force methods, as they do not reflect the professionalism. 2. Indicate your answer with [Answer]option[Answer], such as [Answer]C[Answer]. Please output the answer at the end in strict accordance with the output format. ## Initialization As <Role>, please follow <Rules> strictly. Your task is to solve the math problem following <Workflow>. I will provide you with a problem and 5 options. Please choose the correct option from the five provided. Problem: [Please Put Your Questions Here] Options: [Please Put Your Options Here] Answer: Let's think step by step.
-------------------------	--

Figure 10: An example of prompting for standard Strategic Chain-of-Thought in math reasoning tasks

Zero-shot SCoT template	# Role An expert of world knowledge with strong logical skills. # Workflow 1. Analyze the problem and break down the complex query into simpler sub-questions. 2. Sequentially finding reliable answers for each sub-question. 3. Integrating these answers to form a comprehensive. Directly answering the main question is rude, but breaking it down, answering the sub-questions, and then integrating the answers is elegant and practical. ## Rules 1. Avoid using brute force methods, as they do not reflect the professionalism. 2. Indicate your answer with [Answer]option[Answer], such as [Answer]C[Answer]. Please output the answer at the end in strict accordance with the output format. ## Initialization As <Role>, please follow <Rules> strictly. Your task is to solve the problem following <Workflow>. I will provide you with a problem and 5 options. Please choose the correct option from the five provided. Problem: [Please Put Your Questions Here] Options: [Please Put Your Options Here] Answer: Let's think step by step.
-------------------------	--

Figure 11: An example of prompting for standard Strategic Chain-of-Thought in multi-hop reasoning tasks

Zero-shot SCoT template	# Role An expert with world knowledge and reasoning abilities. # Workflow 1. Understanding the Question: Identify key concepts and comprehend the question's context. Ensure you grasp the main idea and any analogies being used. Search for any concept, knowledge, or common sense related to the topic. 2. Analyzing the Options: Read each choice carefully, understand its meaning, and relate it to the question's context to determine relevance. 3. Logical Reasoning: Use logical reasoning to eliminate options that are clearly irrelevant or incorrect based on the question's context. Compare the remaining options to identify the one that best aligns with the question's requirements and the context provided. ## Rules 1. Avoid using brute force methods, as they do not reflect the professionalism. 2. Indicate your answer with [Answer]option[Answer], such as [Answer]C[Answer]. Please output the answer at the end in strict accordance with the output format. ## Initialization As <Role>, please follow <Rules> strictly. Your task is to solve the problem following <Workflow>. I will provide you with a problem and 5 options. Please choose the correct option from the five provided. Problem: [Please Put Your Questions Here] Options: [Please Put Your Options Here] Answer: Let's think step by step.
-------------------------	--

Figure 12: An example of prompting for standard Strategic Chain-of-Thought in commonsense reasoning tasks

Zero-shot SCoT template	# Role A careful expert proficient in various world knowledge. # Workflow 1. Careful Question Analysis: - Read the Problem and the Options Carefully: Ensure you understand the background and specific question being asked. - Identify Keywords: Extract key terms or phrases from the Problem and the Options, try recalling their meanings. - Understand the Problem: Ensure you clearly understand what the Problem is asking, including any specific conditions or requirements. Eliminate options that are not relevant to the problem. 2. Identify Relevant Knowledge and approaches: - Recall Related Knowledge or approach: Identify all the relevant concepts, principles, or formulas that might apply to the Problem. - Select Appropriate Knowledge: Choose the knowledge, formulas and approaches that can solve the problem. 3. Choose the Most Efficient and Practical Knowledge and Formulas: When solving the problem, select the most efficient and practical knowledge, formulas or approaches. For example, when the description of a problem is related to potential energy and kinetic energy of an object, after using the formula $PE = mgh$, carefully analyze each option to judge right or wrong, rather than relying on experience or ready-made theorems to select options. 4. Careful Application of Knowledge and Formulas: - Detailed Analysis: When applying formulas and knowledge, pay attention to the specific conditions and variables in the problem. - Logical Reasoning: Carefully analyze each variable in the formula or methodically derive conclusions based on the knowledge point, ensuring the reasoning process is consistent and correct. For example, when using $PE = mgh$, you need to analyze the overall effect of all variables, including m, g, and h, rather than just one variable. ## Rules 1. Avoid using brute force methods, as they do not reflect the professionalism. 2. Indicate your answer with [Answer]option[Answer], such as [Answer]C[Answer]. Please output the answer at the end in strict accordance with the output format. ## Initialization As <Role>, please follow <Rules> strictly. Your task is to solve the problem following <Workflow>. I will provide you with a problem and 5 options. Please choose the correct option from the five provided. Problem: [Please Put Your Questions Here] Options: [Please Put Your Options Here] Answer: Let's think step by step.
-------------------------	---

Figure 13: An example of prompting for standard Strategic Chain-of-Thought in physical reasoning tasks

Zero-shot SCoT template	# Role A very meticulous logical Analyst. # Workflow 1. Initial State: First, list the initial state of the balls each person has according to the problem statement. 2. Process Exchanges: Next, carefully read the problem statement. For each exchange, update the current state of the balls and document the result of each exchange. 3. Determine the Answer: Once all exchanges are completed, identify which friend's ball color is being inquired about in the problem statement and select the correct answer. ## Rules 1. Avoid using brute force methods, as they do not reflect the professionalism. 2. Indicate your answer with [Answer]option[Answer], such as [Answer]C[Answer]. Please output the answer at the end in strict accordance with the output format. ## Initialization As <Role>, please follow <Rules> strictly. Your task is to solve the problem following <Workflow>. I will provide you with a problem and 5 options. Please choose the correct option from the five provided. Problem: [Please Put Your Questions Here] Options: [Please Put Your Options Here] Answer: Let's think step by step.
-------------------------	--

Figure 14: An example of prompting for standard Strategic Chain-of-Thought in spatial reasoning tasks

One-shot SCoT template	# Role A highly skilled mathematician and algorithm expert. # Workflow 1. Analyze the problem and identify any relevant mathematical formulas, or approaches that might be helpful, and select the approaches that can solve the problem. 2. Choose the most efficient and practical approach. For example, when asked to find the sum of all integers from -25 to 23, consider using the summation formula of arithmetic sequence instead of simply adding the numbers one by one. The summation formula of arithmetic sequence is an elegant and practical solution, while rudely adding the numbers is not. 3. Solve the problem step by step following the selected approach carefully. ## Demonstrations Problem: [Please Put Your Demonstration Problem Here] Options: [Please Put Your Demonstration Options Here] Answer: [Please Put Your Demonstration Answer Here] ## Rules 1. Avoid using brute force methods, as they do not reflect the professionalism. 2. Indicate your answer with [Answer]option[Answer], such as [Answer]C[Answer]. Please output the answer at the end in strict accordance with the output format. ## Initialization As <Role>, please follow <Rules> strictly. Your task is to solve the math problem following <Workflow>, <Demonstration> is some examples. I will provide you with a problem and 5 options. Please choose the correct option from the five provided. Problem: [Please Put Your Questions Here] Options: [Please Put Your Options Here] Answer: Let's think step by step.
------------------------	---

Figure 15: An example of prompting for one-shot Strategic Chain-of-Thought