

SIPDO: Closed-Loop Prompt Optimization via Synthetic Data Feedback

Anonymous ACL submission

Abstract

Prompt quality plays a critical role in the performance of large language models (LLMs), motivating a growing body of work on prompt optimization. Most existing methods optimize prompts over a fixed dataset, assuming static input distributions and offering limited support for iterative improvement. We introduce **SIPDO** (Self-Improving Prompts through Data-Augmented Optimization), a closed-loop framework for prompt learning that integrates synthetic data generation into the optimization process. SIPDO couples a synthetic data generator with a prompt optimizer, where the generator produces new examples that reveal current prompt weaknesses and the optimizer incrementally refines the prompt in response. This feedback-driven loop enables systematic improvement of prompt performance without assuming access to external supervision or new tasks. Experiments across question answering and reasoning benchmarks show that SIPDO outperforms standard prompt tuning methods, highlighting the value of integrating data synthesis into prompt learning workflows.

1 Introduction

Large language models (LLMs) have demonstrated strong performance across a wide range of natural language tasks, including classification, question answering, and reasoning. However, their output quality is highly sensitive to prompt design—small changes in phrasing, structure, or formatting can lead to significant variations in performance (He et al., 2024; Spiess et al., 2025). This sensitivity has made prompt optimization a core challenge in adapting LLMs to downstream applications.

Prior work in prompt optimization has explored manual tuning, discrete search, and gradient-based methods to improve model responses (Wang et al., 2023; Shin et al., 2020; Cui et al., 2024; Kwon et al., 2024; Zhang et al., 2024). While effective in some settings, these methods typically assume a fixed

input distribution and treat prompt optimization as a one-time procedure. As a result, they may produce prompts that perform well on average but offer limited support for iterative improvement or adaptation to evolving failure modes.

In contrast, data augmentation is a long-standing technique in supervised learning for improving model robustness and generalization by exposing models to diverse training conditions (Mikołajczyk and Grochowski, 2018). In the context of prompt learning, the ability of LLMs to generate high-quality synthetic data opens the door to new forms of self-improvement. However, most existing prompt optimization pipelines do not leverage synthetic data in an adaptive or feedback-driven manner (Singh et al., 2023; Gilardi et al., 2023; Tang et al., 2023; Gao et al., 2023).

We propose **SIPDO** (Self-Improving Prompts through Data-Augmented Optimization), a closed-loop prompt optimization framework that integrates synthetic data generation directly into the learning process. SIPDO couples two components: a synthetic data generator that produces inputs designed to challenge the current prompt, and a prompt optimizer that uses these examples to refine the prompt iteratively. This feedback loop enables prompts to improve over time by addressing their own failures, without requiring access to new tasks or external supervision.

Contributions. This paper makes the following contributions:

- We introduce a feedback-driven framework SIPDO that integrates synthetic data generation into prompt optimization, providing a novel pathway for improving prompt robustness.
- We develop a method to construct synthetic examples that dynamically stress-test prompts, revealing failure modes and guiding refinement.
- We empirically demonstrate that augmenting prompt optimization with synthetic data improves performance across reasoning and QA

benchmarks, surpassing existing prompt tuning methods.

2 Related Work

Automatic Prompt Engineering: Automatically discovering optimal prompts has become a key challenge in the era of large language models (LLMs). Automatic Prompt Engineering (APE) employs optimization-based, generative, and template-driven approaches. Optimization techniques include gradient-based search (Shin et al., 2020), reinforcement learning (Ouyang et al., 2022; Kwon et al., 2024), and evolutionary algorithms (Cui et al., 2024). Generative methods use models like GPT and Gemini to generate candidate prompts, with StablePrompt (Kwon et al., 2024) optimizing prompts via reinforcement learning. Additionally, PromptAgent (Wang et al., 2023) breaks down prompt creation into sub-goals, while template-driven approaches, like fill-in-the-blank formats, ensure clarity (Chen et al., 2024).

Recent work has expanded on automatic prompt optimization techniques. AutoPDL (Spiess et al., 2025) automates the discovery of optimal configurations for agents which successive halving to explore the space of agentic and non-agentic prompting patterns. The sequential optimal learning approach for automated prompt engineering (Wang et al., 2025) uses Bayesian regression and Knowledge-Gradient policies to efficiently identify effective prompt features. Progressively Automatic Prompt Optimization (Qu et al., 2025) introduces an evolution-based algorithm to optimize prompts for visual classification tasks.

We propose a hybrid framework integrating LLM-driven rewriting with natural language feedback (Pryzant et al., 2023), alongside self-reflection (Shinn et al., 2024) and planning (Wang et al., 2023), enhancing prompt adaptability and precision.

Data Synthesis: Using large language models (LLMs) for data synthesis is a relatively new and rapidly evolving approach. Recent advancements have shown that LLMs possess the capability to generate text with fluency and quality comparable to human output (Li et al., 2023; Mukherjee et al., 2023; Eldan and Li, 2023). For instance, prior work (Gao et al., 2022) has explored leveraging pre-trained language models (PLMs) to generate task-specific text data that can be used to train and evaluate. Recent work Magpie (Xu et al., 2024) lever-

ages the auto-regressive nature of aligned LLMs to generate high-quality instruction data. Additionally, Synthetic Text Generation for Training Large Language Models via Gradient Matching (Nguyen et al., 2025) proposes a novel approach to generate synthetic text that matches the gradients of human data. However, these studies have not fully incorporated advanced methodologies such as chain-of-thought (CoT) reasoning, in-context learning, or data synthesis driven by prompts that integrate task descriptions and label information.

In this study, we systematically experimented with a range of techniques, including in-context learning and prompt-driven data synthesis, combining task descriptions and label information. Our results indicate that integrating these approaches produces high-quality synthetic data. To further enhance the robustness and applicability of the synthetic data, we introduced a difficulty tier, making the generated data more challenging. These findings highlight the potential of combining advanced LLM capabilities with tailored prompting strategies to improve data synthesis quality and reliability for prompt optimization.

3 Method

SIPDO presents a two-agent system for optimizing prompts using data augmentation techniques. The workflow has two cooperating agents: (i) Data Generator creates synthetic data with increasing difficulty levels to expose weaknesses in the prompt, and (ii) Auto Prompt Optimizer iteratively analyzes errors and rewrites the prompt to maximize task performance. The two agents advance in lock-step, and together they grow a prompt that remains compact across increasingly challenging samples. An overview of SIPDO is shown in Fig 1.

Notation. We define the true data distribution as S , which governs input-label pairs $(x, y) \in \mathcal{X} \times \mathcal{Y}$. Let N denote the size of an i.i.d. dataset drawn from S , denoted as $\{(x_i, y_i)\}_{i=1}^N \sim S$. We consider a fixed large language model equipped with a prompt $p \in \mathcal{P}$, and define its output function as $f(p, x) \in \mathcal{Y}$. Prediction accuracy is measured using a bounded surrogate loss $L(f(p, x), y)$, where $L \in [0, 1]$.

To help improve the quality of prompts, we introduce a synthetic data generator defined by a distribution $q_\psi(\tilde{x}, \tilde{y})$, parameterized by $\psi \in \Psi$, which produces synthetic samples forming a dataset $D = \{(\tilde{x}_i, \tilde{y}_i)\}_{i=1}^M$, where M is the number of gener-

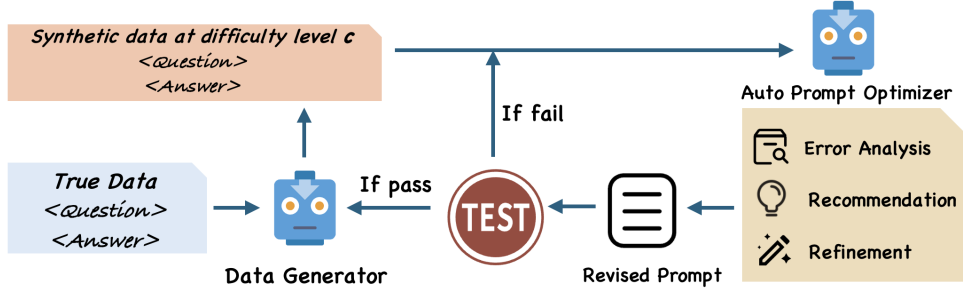


Figure 1: Starting from true data distribution S , the Data Generator(left) produces a synthetic question-answer pair at difficulty level c . The Auto Prompt Optimizer(right) evaluates the current prompt on this synthetic data via three sub-modules-error analysis, recommendation, and refinement-and outputs a revised prompt. The revised prompt is tested on present failures and all previously solved examples. If the prompt still makes errors, then return to the Auto Prompt Optimizer for further refinement; if passes, move on to the next sample(with higher c). The cycle repeats until no error remains or the budget is reached, yielding a self-improved prompt.

ated examples. To ensure that the synthetic labels remain realistic, we estimate the population label prior with $p^*(y)$, and use this to regularize the generator.

3.1 Data Generator

The Data Generator supplies fresh, well-targeted examples that expose the weakness by creating a new synthetic-pair whose difficulty is designed beyond prompt’s current reach.

Sampling rule. The data generator first draws a target label $\tilde{y} \sim p^*(y)$. By sampling a latent variable $z \sim g_\phi(z|S)$ that captures the structure of few-shot S , the decoder q_ψ produces $\tilde{x} = q_\psi(z, \tilde{y}, c)$ where c is a controlled difficulty tier.

Learning objective. The parameters ψ are learned by minimizing a hybrid objective that balances the KL penalty and the bounded surrogate loss:

$$\min_{\psi} R(\psi) + \lambda \mathbb{E}_{(\tilde{x}, \tilde{y}) \sim q_\psi} [L(f(p, \tilde{x}), \tilde{y})], \quad (1)$$

Note that, we penalize deviations from the true label distribution using the Kullback–Leibler divergence term $R(\psi) = \text{KL}(q_\psi(y) \parallel p^*(y))$, scaled by a factor $\lambda^{-1} R(\psi)$ during training.

Progressive complexity. To address tasks of varying complexity, we introduces a progressive complexity parameter c where $c \in \{1, \dots, n\}$ so that prompts could be tested on gradually more challenging examples. This allows the prompts to progressively improve and generalize effectively across task of increasing difficulty. Since q_ψ is conditioned on c , a single latent template (z, y) can

therefore yield n difficulty-aligned variants

$$\{\tilde{x}^{(1)}, \dots, \tilde{x}^{(n)}\} = \{q_\psi(z, y, 1), \dots, q_\psi(z, y, n)\}$$

For curriculum generation, an ordered sequence $c_1 < \dots < c_L$ is sampled and feeds the output of the previous level back into the generator,

$$\tilde{x}^{(1)} = q_\psi(z, y, c_1),$$

$$\tilde{x}^{(2)} = q_\psi(h_\phi(\tilde{x}^{(1)}), y, c_2),$$

$$\vdots$$

$$\tilde{x}^{(L)} = q_\psi(h_\phi(\tilde{x}^{(L-1)}), y, c_L).$$

where h_ϕ is a summarizer that distills the previous sample into a new latent cue, allowing semantic depth to accumulate across levels. The sequence $c_1 < c_2 < \dots < c_L$ guarantees monotone growth of problem complexity, providing a rich gradient of difficulty for the prompt to learn from.

3.2 Auto Prompt Optimizer

After each new synthetic instance is calibrated, the Auto Prompt Optimizer probes the current prompt, identifies the weaknesses, and repairs them before the next instance is drawn. This stage avoids hard-to-diagnose failures and builds a prompt that is both robust and suitable for specific tasks.

Accuracy score. At iteration $t \in \{1, \dots, M\}$, the optimizer improves the current prompt $p^{(t)}$ using the feedback collected from synthetic log $D_t = \{(\tilde{x}_j, \tilde{y}_j)\}_{j=1}^t \subseteq \mathcal{X} \times \mathcal{Y}$. For any prompt p and set $\mathcal{A} \subseteq \mathcal{X} \times \mathcal{Y}$, we define

$$s_{\mathcal{A}}(p) = \frac{1}{|\mathcal{A}|} \sum_{(\tilde{x}, \tilde{y}) \in \mathcal{A}} \mathbb{I}[f(p, \tilde{x}) = \tilde{y}], \quad (1)$$

$\mathbb{I}[\cdot]$ is the indicator function that evaluates to 1 if the prompt’s output matches the target label, and 0 otherwise.

Step 1: Error analysis. We first evaluate $p^{(t)}$ on the whole set and collect the current error slice

$$\mathcal{E}^{(t)} = \{(\tilde{x}, \tilde{y}) \in D \mid f(p^{(t)}, \tilde{x}) \neq \tilde{y}\}.$$

If $\mathcal{E}^{(t)} = \emptyset$, the prompt already "covers" all unseen cases, therefore, we terminate and return $p^* = p^{(t)}$; otherwise, we proceed to the next step.

Step 2: Recommendation. A reflection module $\mathcal{R}_\varphi$ inspects $\mathcal{E}^{(t)}$ and produces a textual-patch suggesting how to modify the prompt:

$$\Delta^{(t)} = \mathcal{R}_\varphi(p^{(t)}, \mathcal{E}^{(t)}).$$

This summarizes why the prompt failed and how it can be amended (e.g., add boundary cases, clarify/revise instructions, drop distracting details).

Step 3: Targeted refinement. A prompt editor $\mathcal{U}_\theta$ applies the patch $\Delta^{(t)}$ to a revised prompt $\tilde{p}^{(t)}$ in order to fix the current error

$$\tilde{p}^{(t)} = \mathcal{U}_\theta(\Delta^{(t)}, p^{(t)}, \mathcal{E}^{(t)}).$$

Local confirmation. We then test revised prompt $\tilde{p}^{(t)}$ only on the current errors: if $s_{\mathcal{E}^{(t)}}(\tilde{p}^{(t)}) < 1$, some errors still remain. In this case, we make the revised prompt as new baseline prompt by setting $p^{(t)} \leftarrow \tilde{p}^{(t)}$, updating $\mathcal{E}^{(t)}$, and repeating Step 2 to generate more sufficient patch $\Delta^{(t)}$; otherwise, proceed to global confirmation.

Global confirmation. Solving the local error slice is not enough—we must ensure that revised prompt "covers" all seen cases. Therefore, we evaluate $\tilde{p}^{(t)}$ on the entire synthetic history collected seen so far by $s_{D_t}(\tilde{p}^{(t)})$. During evaluation, if $\mathcal{E}^{(t)} \neq \emptyset$ at any previous data, we treat them as new error set and sent them back to step 2 with new $\mathcal{E}^{(t)}$ to fix the current error. If $\mathcal{E}^{(t)} = \emptyset$, we accept the revision, set $p^{(t+1)} = \tilde{p}^{(t)}$, draw the next synthetic example, and restart from Step 1 until $t = M$.

Convergence guarantee. Because $s_D(p^{(t)})$ is non-decreasing and bounded above by 1, the process stops at most M successful corrections or the user-chosen cap $T_{\max}$. The final output

$$p^* = \arg \max_{0 \leq t \leq T} s_D(p^{(t)})$$

achieves perfect coverage ($s_{D_T}(p^*) = 1$) whenever it is attainable within the budget.

This revised loop mirrors practical prompt-debugging: it first addresses specific failure case, then confirms that the updated prompt continues to perform correctly on all previously solved examples. By iteratively applying this feedback-driven process, the system systematically refines prompts to improve clarity, adaptability, and overall performance, making the framework highly generalizable across diverse tasks and domains.

3.3 Theoretical Guarantee

Since one of our goals in SIPDO is to demonstrate that, data augmentation, a popular branch of performance improvement in deep learning, can also be used in prompt optimization context, we aim to offer similar performance guarantees as done in previous data augmentation literature (Wang et al., 2022; Chen et al., 2020; Dao et al., 2019).

Assumptions. We first offer the assumptions that we need for the theoretical guarantees.

A1 (Label-preservation) For all $\psi \in \Psi$ and for any (x, y) , the generator’s conditional satisfies $\Pr_{q_\psi}[\tilde{y} = y \mid \tilde{x} \stackrel{q}{\leftarrow} (x, y)] = 1$.

We require the generator never flips the ground-truth label of the base example it is derived from (it may, however, hallucinate novel inputs as long as their labels match the intended classes).

A2 (Approximate maximizer). Let

$$\psi^* = \arg \max_{\psi \in \Psi} \mathbb{E}_{q_\psi} L(f(p, \tilde{x}), \tilde{y}) - \lambda^{-1} R(\psi),$$

The inner-loop training of the generator attains a value at most ε below this supremum.

A perfect maximizer would be ideal but is infeasible; we only need the learned generator to be good enough—within ε of optimal. The residual ε directly appears in the bound.

A3 (Uniform convergence). (Wang et al., 2022) For every prompt p , the empirical loss deviates from its population counterpart by at most $q(|\mathcal{P}|, n, \delta)$ with probability $1 - \delta$, where a standard form of $q(|\mathcal{P}|, n, \delta)$ is $\tilde{O}\left(\sqrt{\frac{\log |\mathcal{P}| + \log(1/\delta)}{N}}\right)$.

PAC(probably approximately correct) guarantee: empirical performance generalizes provided n is large enough.

A4 (Alignment of risks). For any prompt p and generator ψ ,

$$\mathbb{E}_{q_\psi} L(f(p, \tilde{x}), \tilde{y}) \leq \mathbb{E}_S L(f(p, x), y) + \lambda^{-1} R(\psi).$$

The KL penalty controls how far the generator may wander: if it manufactures rare-label outliers, $R(\psi)$ increases and the bound tightens. We can verify that $q_\psi(y)$ is always absolutely-continuous w.r.t. $p^*(y)$; KL is then finite and the inequality follows from the classical Donsker–Varadhan variational formula.

A5 (Surrogate link). The 0–1 loss is upper-bounded by the surrogate loss:

$$\mathbb{1}\{f(p, x) \neq y\} \leq L(f(p, x), y).$$

This is needed in order to translate guarantees on the differentiable training loss to the classification error (e.g. cross-entropy, hinge, logistic).

Theorem 3.1 Regularised Worst-case Data Generation Under Assumptions B1, B2, B4–B6, for any fixed prompt $p \in \mathcal{P}$, with probability at least $1 - \delta$ over the draw of the training set, we have

$$\underbrace{\sup_{\psi \in \Psi} \mathbb{E}_{q_\psi} \mathbb{1}\{f(p, \tilde{x}) \neq \tilde{y}\}}_{\text{population worst-case error}} \leq \underbrace{\frac{1}{n} \sum_{i=1}^n L(f(p, x_i), y_i)}_{\text{empirical risk}} + \underbrace{\lambda^{-1} R(\psi^*)}_{\text{KL penalty of hardest generator}} + \varepsilon + q(|\mathcal{P}|, n, \delta). \quad (2)$$

Practical implication The inequality states that if the empirical loss of the prompt is low, and no generator can inflate that loss without paying a high KL tax, then even a hypothetically all-powerful adversary (generator) cannot cause the prompt to misclassify more than the RHS. Selecting a larger λ tightens the KL tax, thus lowering the worst-case error but potentially harming accuracy—precisely the robustness–performance trade-off observed empirically in Section 4.

4 Experiments

We test SIPDO on four main datasets to measure its resilience across different domains and reasoning tasks. We include all 4689 instances from six BIG-Bench tasks (Tables of Penguins, Geometric Shapes,

Epistemic Reasoning, Object Counting, Temporal Sequences, and Causal Judgment (Srivastava et al., 2022)). To assess logical reasoning, we sample 600 examples from the depth-5 subset of ProofWriter with a balanced label distribution (Tafjord et al., 2021), use 204 test examples from FOLIO that require first-order inference over short passages (Han et al., 2024), and select the 500 most challenging 5-hop scenarios from the fictional-character version of PrOntoQA (Saparov and He, 2022).

4.1 EXPERIMENTAL SETUP

4.1.1 Baselines

We compare SIPDO with four existing prompt optimization strategies:

Chain of Thought (CoT) (Suzgun et al., 2022) improves LLM reasoning by explicitly guiding models through step-by-step decomposition of complex tasks.

Automatic Prompt Engineer (APE) (Wang et al., 2023) automatically generates and refines prompts using a Monte Carlo search based on model feedback, improving instruction quality with minimal human intervention.

PromptAgent (Zhou et al., 2022b) formulates prompt optimization as a strategic planning task, using Monte Carlo Tree Search (MCTS) to explore the prompt space. It refines prompts based on model errors and feedback, aiming to generate prompts with expert-level quality through systematic trial-and-error refinement.

Neuro-Symbolic (Pan et al., 2023) combines neural networks with symbolic rule-based reasoning by transforming LLM outputs into structured symbolic representations. It enables models to handle complex logical inference tasks by integrating natural language understanding with robust symbolic processing, offering a more versatile and human-like approach.

REVOLVE (Zhang et al., 2024) tracks the way an LLM’s responses evolve over successive iterations and updates the prompt with this trajectory-aware signal, giving the optimizer a second-order-style view of progress. By looking beyond the most recent feedback, it avoids local optima and usually reaches a higher-quality prompt in fewer steps.

4.2 Results and analyze

4.2.1 Synthetic Data Generation

We fix the difficulty budget at $c = 10$ for all benchmarks except Penguins and Geometry from BIG-

Bench, where we double it to $c = 20$ to accommodate the intrinsically harder structure of these tasks. For every dataset, we also tie the number of training iterations to the difficulty level by setting $t = c$; the optimiser therefore sees one harder tier at each pass. A sample of synthetic Causal-Judgement data produced by SIPDO at difficulty level 8 (on a 1-10 scale) is shown in the text box below with additional examples in Appendix D. These data ensure a highly consistent format with a well-organized structure and clear hierarchical distinctions. Each output adheres strictly to pre-defined standards, guaranteeing data stability and integrity while significantly reducing the cost of subsequent cleaning and adjustments. Additionally, the structured format enhances readability and parsability, streamlining analysis and processing while improving data reusability and reliability.

Generated Question: In a small town, there is a bakery that makes the best pastries. Every morning, the bakery opens at 7:00 am, and a line of customers forms outside. The bakery owner has a rule that only one person can enter at a time to maintain order. One day, two customers, Alice and Bob, arrive at the same time. Alice follows the rule and waits outside, but Bob ignores the rule and enters the bakery while another customer is still inside. The bakery becomes overcrowded, and a shelf of pastries falls over, ruining the day's batch. Did Bob cause the pastries to be ruined?

Generated Answer: Yes

By generating complex and challenging scenarios gradually, the data stresses the current prompt and express fresh failure modes, improving its ability to handle a wide range of difficulties effectively. This ensures that the prompt becomes more adaptable and robust as it reaches the hardest tiers. These results strongly validate the effectiveness of structured output in standardizing data generation and improving overall data quality, while fostering continuous advancements in prompt optimization.

4.2.2 Implementation

Geometry data generation and prompt improvement. Constructing complex and irregular shapes, including irregular shapes and polygons, cannot be effectively achieved using few-shot methods, so we introduce three measures specifically for SVG (Scalable Vector Graphics) path generation in geometry task: (1) precision normalization: every coordinate produced by the generator is rounded to two decimal places, eliminating the floating-point drift that causes downstream parsers

to miscount line (L) and arc (A) instructions. (2) template-guided retrieval: before generation, a retriever selects a canonical path template whose instruction pattern matches the target shape (e.g., "4 L" for a rectangle, "1 A" for a sector). The generator then perturbs only the vertex coordinates, guaranteeing syntactic correctness while still exposing the prompt to unseen geometries. (3) reverse-generation check: because the label (shape name) is known a priori, we parse the produced SVG with a deterministic rule-based decoder that counts L and A commands; if the inferred label disagrees with the intended one, the sample is rejected and regenerated. The generated SVG path example is shown below. For further prompt refinement, we provide prompt templates in Appendix C.

```
Generated SVG Path: <path d="M 23.45,45.78
L 78.32,45.78 L 78.32,78.56 L 23.45,78.56 L
23.45,45.78"/>
```

```
Generated Answer:
"target_scores": {
  "rectangle": 1,
  "sector": 0,
  "triangle": 0,
  "circle": 0,
  "heptagon": 0,
  "hexagon": 0,
  "kite": 0,
  "line": 0,
  "octagon": 0,
  "pentagon": 0 }
```

4.2.3 Comparison with prompting baselines

We tested SIPDO by different LLMs such as GPT-4o, GPT-4o-mini, Gemini-1.5-flash, and Gemini-1.5-pro on different datasets, including BIG-Bench, FOLIO, PrOntoQA, and ProofWriter.

BIG-Bench. We evaluated SIPDO on six BIG-Bench tasks. As shown in Table 1, GPT-4o and GPT-4o-mini demonstrate particularly strong performance in Temporal Reasoning, Object Counting, and Causal Judgment. While Geometry exhibits comparable accuracy across GPT-4o and GPT-4o-mini, SIPDO achieves the highest overall accuracy for GPT-4o-mini, Gemini-1.5-flash, and Gemini-1.5-pro. For GPT-4o, SIPDO performs at a near-best level, trailing PromptAgent by only 0.01, yet still demonstrating that LLMs benefit from synthetic data generation for reasoning improvements, whereas other methods primarily rely on existing datasets. These results further highlight the advantages of LLM-driven data augmentation in enhancing logical, numerical, and causal reasoning capabilities.

Table 1: Results on BIG-Bench tasks across multiple LLMs. SIPDO consistently outperforms standard prompting baselines (CoT, APE, PromptAgent) across most tasks and models, demonstrating generalization and effectiveness of the optimization.

Model	Method	Accuracy (%)						Avg. (Comparative Acc.)
		Penguins	Geometry	Epistemic	Obj. Count	Temporal	Causal	
GPT-4o	CoT	79.8	79.1	79.3	85.2	98.0	67.8	81.5(↓ 7.6)
	APE	84.8	65.3	84.8	86.0	99.2	74.0	82.4(↓ 6.7)
	PromptAgent	96.1	83.0	91.6	88.2	98.4	77.8	89.2 (↑ 0.1)
	SIPDO	96.4	82.2	86.3	91.1	99.3	79.0	89.1
GPT-4o-mini	CoT	75.8	68.6	85.2	81.5	94.9	63.6	78.3(↓ 9.0)
	APE	83.7	44.5	81.6	86.3	97.2	75.6	78.2(↓ 9.1)
	PromptAgent	89.8	72.0	86.0	84.3	94.6	84.6	85.2(↓ 2.1)
	SIPDO	92.1	73.2	85.1	87.5	98.0	88.0	87.3
Gemini-1.5-flash	CoT	70.4	68.3	85.5	90.1	94.0	66.8	79.2(↓ 3.7)
	APE	37.6	49.4	88.8	84.7	99.4	69.4	71.6(↓ 11.3)
	PromptAgent	67.4	70.3	81.6	86.3	94.2	67.9	78.0(↓ 4.9)
	SIPDO	77.3	68.9	87.0	92.3	98.4	73.2	82.9
Gemini-1.5-pro	CoT	81.8	59.1	82.6	92.8	98.9	61.5	79.5(↓ 3.9)
	APE	40.2	56.6	88.7	78.6	86.0	65.7	69.3(↓ 14.1)
	PromptAgent	73.6	58.3	83.8	72.6	98.4	74.2	76.8(↓ 6.6)
	SIPDO	79.3	64.3	89.3	91.3	98.0	78.3	83.4

FOLIO, PrOntoQA, and ProofWriter. We evaluated SIPDO, REVOLVE, CoT, Neuro-Symbolic, and a plain baseline prompt on FOLIO, PrOntoQA, and ProofWriter with both GPT-4o and GPT-4o-mini, assessing the methods’ ability to perform structured logical reasoning. As shown in Table 2, SIPDO achieves the highest average accuracy score and outperforms all approaches on FOLIO and PrOntoQA. In PrOntoQA, SIPDO surpasses all methods, demonstrating its capability to generate structured logical proofs. Similarly, for FOLIO, SIPDO outperforms Neuro-Symbolic, CoT, and REVOLVE, further validating its effectiveness in formal logic inference.

While neuro-symbolic reasoning remains the best performer on ProofWriter, SIPDO achieves highly competitive results on ProofWriter, trailing by only 0.004 on GPT-4o-mini and 0.02 on GPT-4o, underscoring its strong adaptability to structured reasoning tasks. Crucially, unlike neuro-symbolic approaches that rely on predefined rule-based datasets, SIPDO is trained entirely on generated synthetic data, demonstrating the effectiveness of LLM-driven data augmentation for enhancing logical inference across diverse reasoning benchmarks.

4.2.4 Prompt generalization

The whole process of the prompt improvement on the Penguin data is shown in Appendix C. Follow the three modules-error analysis, recommendation, and refinement-and outputs a revised prompt. The prompts generated by SIPDO out-

perform other baselines, validating the effectiveness of SIPDO in optimizing prompt design and enhancing overall performance. By employing a LEAST-TO-MOST(Zhou et al., 2022a) approach in prompt generation, the LLM to reason in a structured and incremental manner. After each iteration, the prompts are evaluated using real datasets. If the accuracy falls below a predefined threshold, the prompt is refined based on incorrectly answered questions to enhance its effectiveness. All prompts can be viewed in Appendix A

4.3 Ablation Study

Difficulty Gradient. To assess the contribution of the difficulty gradient, we conduct an ablation study by comparing without difficulty gradient. As Table 3 shows, every BIG-Bench sub-task suffers when the difficulty gradient is absent. On average, GPT-4o loses 20.8% accuracy, while the weaker GPT-4o-mini drops 32.1%, confirming that smaller models depends even more on the difficulty gradient. The steepest declines appear on tasks Object Counting (-69.3 % and -119.3 %) and Geometric Shapes (-20.7 % and -54.1 %). Even comparatively simple tasks—Temporal Sequences and Epistemic Reasoning—still lose up to 6 %. These findings underline that a progressive difficulty gradient is essential: it systematically uncovers a prompt’s blind spots and lets the optimizer repair them before moving on to harder examples. The generated prompts tend to be easier and shorter without difficulty gradient placed, often failing to capture complex reasoning patterns(details in Appendix

Table 2: Results on ProofWriter, FOLIO, and PrOntoQA by Neuro-Symbolic, CoT, REVOLVE, SIPDO, and Baseline Prompting methods across GPT-4o and GPT-4o-mini.

Tasks	GPT-4o					GPT-4o-mini				
	Baseline	Neuro-S	CoT	REVOLVE	SIPDO	Baseline	Neuro-S	CoT	REVOLVE	SIPDO
ProofWriter	58.5	81.6	72.3	54.0	79.6	52.6	79.7	61.8	48.6	79.3
FOLIO	71.2	79.2	72.6	65.7	83.2	51.2	73.2	69.3	62.8	81.1
PrOntoQA	80.4	85.2	95.6	85.4	96.3	74.6	79.3	89.3	83.4	91.3
Average	70.0	82.0	80.2	68.4	86.4	59.5	77.4	73.5	64.9	83.9

Table 3: Accuracy (%) after removing the difficulty gradient. Numbers in parentheses show the absolute drop ($\downarrow$) relative to the performance with difficulty gradient placed.

Model	PENGUINS	GEOMETRY	EPISTEMIC	OBJ.CNT.	TEMPORAL	CAUSAL	Avg.
GPT-4o	73.2 ($\downarrow$ 31.7%)	68.1 ($\downarrow$ 20.7%)	81.9 ($\downarrow$ 5.4%)	53.8 ($\downarrow$ 69.3%)	97.0 ($\downarrow$ 2.4%)	68.4 ($\downarrow$ 15.5%)	73.7 ($\downarrow$ 20.8%)
GPT-4o-mini	69.6 ($\downarrow$ 32.3%)	47.5 ($\downarrow$ 54.1%)	80.0 ($\downarrow$ 6.4%)	39.9 ($\downarrow$ 119.3%)	92.1 ($\downarrow$ 6.4%)	67.4 ($\downarrow$ 30.6%)	66.1 ($\downarrow$ 32.1%)

B).

One-Shot Extremes. We experimented with replacing the step-wise difficulty gradient by a one-shot extremes sampler that tells the generator to create the most unusual examples. On our synthetic suites this shortcut delivered no measurable gain. The “extreme” samples were either solved instantly or only slight perturbations of original cases, leaving the optimizer with no fresh errors to exploit. We suspect the idea will pay off in real-world corpora—financial statements, medical notes, legal filings—where genuine edge cases abound and can expose blind spots that our synthetic tasks do not capture.

Reflection Check. We also inserted a reflection check that re-parses each synthetic pair and discards it if the generator’s answer does not match a rule-based output. Because our generator is already conditioned on the ground-truth label (Assumption A1) and draws heavily on templated examples, most outputs are self-consistent. We, therefore, disable the check in the main pipeline for efficiency, but note that it would be prudent to re-enable it when moving to noisier domains.

5 Conclusion

In summary, we introduce SIPDO, a method that transforms data augmentation into a real-time feedback signal to enhance prompt optimization. By integrating a data generator that synthesizes progressively more difficult examples with an auto-prompt optimizer that refines prompts, SIPDO sys-

tematically identifies and resolves prompts’ weaknesses. This iterative feedback-driven loop improves prompt robustness and performance across diverse reasoning tasks and benchmarks. Empirical evaluations show that SIPDO achieves significant accuracy gains, outperforming several prompt optimization baselines. Additionally, we provide theoretical backing by demonstrating bounded worst-case error once training reaches stability, further supporting SIPDO’s reliability in practice. By enabling LLM systems to autonomously recognize and correct their own shortcomings, SIPDO presents a scalable and efficient strategy toward adaptive, self-improving models capable of generalizing reliably across unseen domains and increasingly complex challenges.

Limitations

Although our multi-agent framework shows clear gains on standard reasoning benchmarks, two practical gaps remain. First, all experiments were run on clean public datasets; we have not yet tested SIPDO on messier, domain-specific collections such as financial filings or clinical notes, where labeling rules and text quality vary widely. Second, the current loop relies on repeated calls, which can be slow and costly at scale. Evaluating on real-world data and exploring lighter generator/critic models are therefore important directions for future work.

References

- Shuxiao Chen, Edgar Dobriban, and Jane H Lee. 2020. [A group-theoretic framework for data augmentation](#). *Preprint*, arXiv:1907.10905.
- Yongchao Chen, Jacob Arkin, Yilun Hao, Yang Zhang, Nicholas Roy, and Chuchu Fan. 2024. [Prompt optimization in multi-step tasks \(promst\): Integrating human feedback and heuristic-based sampling](#). *Preprint*, arXiv:2402.08702.
- Wendi Cui, Jiaxin Zhang, Zhuohang Li, Hao Sun, Damien Lopez, Kamalika Das, Bradley Malin, and Sricharan Kumar. 2024. Phasevo: Towards unified in-context prompt optimization for large language models. *arXiv preprint arXiv:2402.11347*.
- Tri Dao, Albert Gu, Alexander J. Ratner, Virginia Smith, Christopher De Sa, and Christopher Ré. 2019. [A kernel theory of modern data augmentation](#). *Preprint*, arXiv:1803.06084.
- Ronen Eldan and Yuanzhi Li. 2023. [Tinystories: How small can language models be and still speak coherent english?](#) *Preprint*, arXiv:2305.07759.
- Jiahui Gao, Renjie Pi, Yong Lin, Hang Xu, Jiacheng Ye, Zhiyong Wu, Weizhong Zhang, Xiaodan Liang, Zhenguo Li, and Lingpeng Kong. 2022. Self-guided noise-free data generation for efficient zero-shot learning. *arXiv preprint arXiv:2205.12679*.
- Jiahui Gao, Renjie Pi, Yong Lin, Hang Xu, Jiacheng Ye, Zhiyong Wu, Weizhong Zhang, Xiaodan Liang, Zhenguo Li, and Lingpeng Kong. 2023. [Self-guided noise-free data generation for efficient zero-shot learning](#). *Preprint*, arXiv:2205.12679.
- Fabrizio Gilardi, Meysam Alizadeh, and Maël Kubli. 2023. Chatgpt outperforms crowd workers for text-annotation tasks. *Proceedings of the National Academy of Sciences*, 120(30):e2305016120.
- Simeng Han, Hailey Schoelkopf, Yilun Zhao, Zhenting Qi, Martin Riddell, Wenfei Zhou, James Coady, David Peng, Yujie Qiao, Luke Benson, Lucy Sun, Alex Wardle-Solano, Hannah Szabo, Ekaterina Zubova, Matthew Burtell, Jonathan Fan, Yixin Liu, Brian Wong, Malcolm Sailor, and 16 others. 2024. [Folio: Natural language reasoning with first-order logic](#). *Preprint*, arXiv:2209.00840.
- Jia He, Mukund Runqta, David Koleczek, Arshdeep Sekhon, Franklin X Wang, and Sadid Hasan. 2024. [Does prompt formatting have any impact on llm performance?](#) *Preprint*, arXiv:2411.10541.
- Minchan Kwon, Gaeun Kim, Jongsuk Kim, Haeil Lee, and Junmo Kim. 2024. Stableprompt: Automatic prompt tuning using reinforcement learning for large language models. *arXiv preprint arXiv:2410.07652*.
- Yifei Li, Zeqi Lin, Shizhuo Zhang, Qiang Fu, Bei Chen, Jian-Guang Lou, and Weizhu Chen. 2023. Making language models better reasoners with step-aware verifier. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5315–5333.
- Agnieszka Mikołajczyk and Michał Grochowski. 2018. Data augmentation for improving deep learning in image classification problem. In *2018 international interdisciplinary PhD workshop (IIPhDW)*, pages 117–122. IEEE.
- Subhabrata Mukherjee, Arindam Mitra, Ganesh Jawahar, Sahaj Agarwal, Hamid Palangi, and Ahmed Awadallah. 2023. Orca: Progressive learning from complex explanation traces of gpt-4. *arXiv preprint arXiv:2306.02707*.
- Dang Nguyen, Zeman Li, Mohammadhossein Batani, Vahab Mirrokni, Meisam Razaviyayn, and Baharan Mirzasoleiman. 2025. [Synthetic text generation for training large language models via gradient matching](#). *Preprint*, arXiv:2502.17607.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, and 1 others. 2022. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744.
- Liangming Pan, Alon Albalak, Xinyi Wang, and William Yang Wang. 2023. Logic-lm: Empowering large language models with symbolic solvers for faithful logical reasoning. *arXiv preprint arXiv:2305.12295*.
- Reid Pryzant, Dan Iter, Jerry Li, Yin Tat Lee, Chenguang Zhu, and Michael Zeng. 2023. [Automatic prompt optimization with "gradient descent" and beam search](#). *Preprint*, arXiv:2305.03495.
- Xiangyan Qu, Gaopeng Gou, Jiamin Zhuang, Jing Yu, Kun Song, Qihao Wang, Yili Li, and Gang Xiong. 2025. [Proapo: Progressively automatic prompt optimization for visual classification](#). *Preprint*, arXiv:2502.19844.
- Abulhair Saparov and He He. 2022. Language models are greedy reasoners: A systematic formal analysis of chain-of-thought. *arXiv preprint arXiv:2210.01240*.
- Taylor Shin, Yasaman Razeghi, Robert L. Logan IV, Eric Wallace, and Sameer Singh. 2020. [Auto-prompt: Eliciting knowledge from language models with automatically generated prompts](#). *Preprint*, arXiv:2010.15980.
- Noah Shinn, Federico Cassano, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. 2024. Reflexion: Language agents with verbal reinforcement learning. *Advances in Neural Information Processing Systems*, 36.
- Avi Singh, John D Co-Reyes, Rishabh Agarwal, Ankesh Anand, Piyush Patil, Xavier Garcia, Peter J Liu, James Harrison, Jaehoon Lee, Kelvin Xu, and 1 others. 2023. Beyond human data: Scaling self-training

for problem-solving with language models. *arXiv preprint arXiv:2312.06585*.

Claudio Spiess, Mandana Vaziri, Louis Mandel, and Martin Hirzel. 2025. [Autopdl: Automatic prompt optimization for llm agents](#). *Preprint*, arXiv:2504.04365.

Aarohi Srivastava, Abhinav Rastogi, Abhishek Rao, Abu Awal Md Shoeb, Abubakar Abid, Adam Fisch, Adam R Brown, Adam Santoro, Aditya Gupta, Adrià Garriga-Alonso, and 1 others. 2022. Beyond the imitation game: Quantifying and extrapolating the capabilities of language models. *arXiv preprint arXiv:2206.04615*.

Mirac Suzgun, Nathan Scales, Nathanael Schärli, Sebastian Gehrmann, Yi Tay, Hyung Won Chung, Aakanksha Chowdhery, Quoc V Le, Ed H Chi, Denny Zhou, and 1 others. 2022. Challenging big-bench tasks and whether chain-of-thought can solve them. *arXiv preprint arXiv:2210.09261*.

Oyvind Tafjord, Bhavana Dalvi Mishra, and Peter Clark. 2021. [Proofwriter: Generating implications, proofs, and abductive statements over natural language](#). *Preprint*, arXiv:2012.13048.

Changli Tang, Wenyi Yu, Guangzhi Sun, Xianzhao Chen, Tian Tan, Wei Li, Lu Lu, Zejun Ma, and Chao Zhang. 2023. Salmonn: Towards generic hearing abilities for large language models. *arXiv preprint arXiv:2310.13289*.

Haohan Wang, Zeyi Huang, Xindi Wu, and Eric Xing. 2022. Toward learning robust and invariant representations with alignment regularization and data augmentation. In *Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pages 1846–1856.

Shuyang Wang, Somayeh Moazeni, and Diego Klabjan. 2025. [A sequential optimal learning approach to automated prompt engineering in large language models](#). *Preprint*, arXiv:2501.03508.

Xinyuan Wang, Chenxi Li, Zhen Wang, Fan Bai, Haotian Luo, Jiayou Zhang, Nebojsa Jojic, Eric P Xing, and Zhiting Hu. 2023. Promptagent: Strategic planning with language models enables expert-level prompt optimization. *arXiv preprint arXiv:2310.16427*.

Zhangchen Xu, Fengqing Jiang, Luyao Niu, Yuntian Deng, Radha Poovendran, Yejin Choi, and Bill Yuchen Lin. 2024. [Magpie: Alignment data synthesis from scratch by prompting aligned llms with nothing](#). *Preprint*, arXiv:2406.08464.

Peiyan Zhang, Haibo Jin, Leyang Hu, Xinnuo Li, Liying Kang, Man Luo, Yangqiu Song, and Haohan Wang. 2024. Revolve: Optimizing ai systems by tracking response evolution in textual optimization. *arXiv preprint arXiv:2412.03092*.

Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei, Nathan Scales, Xuezhi Wang, Dale Schuurmans, Claire Cui, Olivier Bousquet, Quoc Le, and 1 others. 2022a. Least-to-most prompting enables complex reasoning in large language models. *arXiv preprint arXiv:2205.10625*.

Yongchao Zhou, Andrei Ioan Muresanu, Ziwen Han, Keiran Paster, Silviu Pitis, Harris Chan, and Jimmy Ba. 2022b. Large language models are human-level prompt engineers. *arXiv preprint arXiv:2211.01910*.

A Optimized Prompts for different tasks

In this section, we demonstrate optimized prompts by Chain-of-Thought (CoT), Automatic Prompt Engineering (APE), PromptAgent, and our method with Accuracys respectively.

Table 4: Comparison of Optimized Prompts for **Object Counting task**, including CoT, APE, PromptAgent, and our method

Approach	Optimized Prompt	Accuracy
CoT	Your task is to count the total number of objects mentioned in the question. Follow these simple steps to ensure accurate counting: Steps to Follow: 1. Identify Items: Read the question carefully and list all objects mentioned. 2. Count Quantities: For each item, check if a quantity is provided. If no quantity is mentioned, assume it is one. 3. Add Totals: Add up the quantities of all items to calculate the total count. 4. Verify the Total: Double-check to ensure no item is missed or counted twice. Example: - Question: "Count the apples, oranges, and bananas. There are 2 apples, 1 orange, and 3 bananas." - Step 1: Identify items: apples, oranges, bananas. - Step 2: Count quantities: 2 apples, 1 orange, 3 bananas. - Step 3: Add totals: $2 + 1 + 3 = 6$. - Step 4: Verify: All items are accounted for, total is 6. - Output: "The total count is 6." Use this step-by-step method for every question to ensure accurate and clear results.	0.928
APE	Calculate the overall total of all items even those spoken in groups.	0.863
PromptAgent	Carefully examine the provided information. Identify and catalog each mentioned item, ensuring that explicitly stated quantities are accurately recorded. If no quantity is specified for an item, assume it as a single unit. However, for items with defined quantities, count each unit separately and include it in the total. If collective terms or categories are mentioned, break them down into their individual components and associate each with its stated count. When computing the total for such categories, ensure that the sum reflects all individual units rather than just the number of groups or types. Each item should be counted independently, but related items belonging to a common category should be grouped together, with their specific quantities contributing precisely to the final total. Avoid assumptions regarding the classification or nature of items—adhere to standard, widely accepted definitions. Finally, summarize the count by explicitly listing the quantity of each identified item or category, and provide a comprehensive total of individual units rather than just category counts, unless otherwise specified.	0.882

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
Our	Task Requirements: The task involves counting the total number of objects listed in a question. Each distinct object should be considered as part of the total count, regardless of its type or variation. The output should be formatted correctly as specified. Problem Rule Application: Identify all items listed in the question. Count each item exactly once, regardless of type, to determine the total number of objects. Ensure accuracy by verifying that all listed items have been included in the final count. Provide the final result in the required format: The number should be presented in both word form and numerical form, separated by a comma (e.g., "nine, 9"). No extra symbols, characters, or explanations should be included. Judgment Criteria: (Strictly follow these rules) Complete Identification: Extract and recognize all objects in the given list. Do not overlook any item mentioned in the question. Accurate Counting: Each item must be counted exactly once. Ensure no items are omitted or double-counted. Verification Process: Double-check the list to confirm that all objects are included. Cross-verify the final count to avoid errors.	0.923

Table 5: Comparison of Optimized Prompts for **Penguins In A Table task**, including CoT, APE, PromptAgent, and our method

Approach	Optimized Prompt	Accuracy
CoT	You are tasked with answering questions about a table of penguins and their attributes. Use step-by-step reasoning to ensure accuracy in calculations and comparisons. The table is as follows: ““ Name, Age, Height (cm), Weight (kg) Louis, 7, 50, 11 Bernard, 5, 80, 13 Vincent, 9, 60, 11 Gwen, 8, 70, 15 ““ **Reasoning Steps for Each Question:** 1. Identify the target attribute (age, height, or weight) and the type of operation (comparison, ranking, filtering). 2. Extract the relevant rows or columns based on the question’s requirements. 3. Perform the required operation step-by-step using the extracted data. 4. Clearly summarize the answer based on the operation’s result. Example Workflow: - Question: "Who is the tallest penguin?" - Step 1: Identify the target attribute: Height. - Step 2: Extract the height values and corresponding names: [(Louis, 50), (Bernard, 80), (Vincent, 60), (Gwen, 70)]. - Step 3: Find the maximum height: Bernard (80 cm). - Step 4: Output the result: "Bernard is the tallest penguin with a height of 80 cm." Follow this workflow for every question to ensure clarity and correctness.	0.818
APE	Carefully scrutinize the provided table or tables. Understand the query in relation to the information given. Pinpoint the pertinent data and carry out the vital computations or comparisons to determine the right answer from the given choices.	0.848

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
PromptAgent	Answer questions about a table of penguins and their attributes, considering both the penguin table and any additional relevant tables. Please provide step-by-step reasoning for your answers, and ensure to clarify any criteria used for filtering or sorting data. Here is a table where the first line is a header and each subsequent line is a penguin: name, age, height (cm), weight (kg) Louis, 7, 50, 11 Bernard, 5, 80, 13 Vincent, 9, 60, 11 Gwen, 8, 70, 15 For example: the age of Louis is 7, the weight of Gwen is 15 kg, the height of Bernard is 80 cm. What is the name of the last penguin sorted by alphabetic order? Options: (A) Louis (B) Bernard (C) Vincent (D) Gwen (E) James **Instructions**: 1. List the names of the penguins. 2. Sort the names alphabetically and present the sorted list clearly. 3. Identify the last name in the sorted list and indicate the corresponding option letter from the provided options. 4. If the last name does not match any of the options, select the name that is closest to the last name in the original list of penguins. At the end, show the answer option bracketed between <answer> and </answer>.	0.961
Our	Answer questions about a dynamic, comprehensive table of penguins and their attributes that allows penguins and other animals to be added and removed. Perform calculations and comparisons based on the questions asked. Read the question carefully to determine which attribute is being compared (age, height, weight). When comparing an attribute, extract the name and that attribute, and then compare, ignoring the other attributes. Ensure the extracted value is from the correct column corresponding to the requested attribute. When using the table, align the data so that the first number is age, the second is height, and the third is weight. Understand the question correctly, find the key words from it, and then perform calculations or comparisons based on the key words The current table is as follows: Name, Age, Height (cm), Weight (kg) Louis, 7, 50, 11 Bernard, 5, 80, 13 Vincent, 9, 60, 11 Gwen, 8, 70, 15 **Question Rules to Apply** - Identify the rows or columns that meet the specified conditions. - Retrieve the value of the required attribute from the identified rows or columns. When we modify this table (by adding new penguins or removing existing penguins or adding giraffes), we first confirm whether the information we added is a penguin or a giraffe, and then solve the problem of comparing, ranking, and filtering based on attributes between penguins or giraffes, depending on the problem.	0.964

Table 6: Comparison of Optimized Prompts for **Geometric Shapes task**, , including CoT, APE, PromptAgent, and our method

Approach	Optimized Prompt	Accuracy
CoT	Your task is to identify the geometric shape represented by the given SVG path data. Follow these steps to ensure accuracy: Steps to Identify the Shape: 1. Check for 'A' Instructions: - If the path contains 'A', determine: - Circle: 2 or more 'A' instructions. - Sector: 1 'A' instruction. 2. Count 'L' Instructions: - If there are no 'A' instructions, count the 'L' instructions to determine the polygon's shape: - Line: 2 'L'. - Triangle: 3 'L'. - Rectangle: 4 'L'. - Pentagon: 5 'L'. - Hexagon: 6 'L'. - Heptagon: 7 'L'. - Octagon: 8 'L'. - Kite: 4 'L'. 3. Provide the Shape Name: Output only the name of the shape (e.g., "circle", "triangle", "hexagon"). Example: - Input: "M 10 10 L 20 10 L 20 20 L 10 20 Z" - Step 1: No 'A' instructions. - Step 2: Count 'L' instructions: 4 'L'. - Step 3: Shape is a Rectangle. - Output: "rectangle". Use this step-by-step process for all inputs to determine the correct shape.	0.791
APE	Determine the shape each SVG path element is drawing, then pair it with the corresponding letter from the available choices. In this case, C symbolizes hexagon, G is for pentagon, I signifies sector, and B stands for heptagon.	0.650

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
PromptAgent	Analyze the SVG path data to identify the geometric shape it represents. Follow these comprehensive and refined steps to ensure accurate identification: Holistic Path Closure: Determine if the path forms a closed shape by checking if the last point connects back to the starting point. If multiple 'M' commands are present, analyze the segments collectively to identify any closed loops. Treat the entire path as a single entity for a thorough analysis. Segment and Side Analysis: Identify the types of segments used in the path: Line Segments: Count the number of distinct line segments to determine the number of sides. Ensure accurate counting by considering all segments collectively. Arc Segments: For paths using the 'A' command, note that these represent elliptical arcs. Pay attention to the parameters to distinguish between circles and ellipses. In-depth Geometric Properties: - For line segments, analyze the relative lengths of sides and angles between them. Consider properties such as parallel sides, equal side lengths, and right angles to distinguish between different types of polygons. Evaluate the overall shape formed by all segments. - For arc segments, examine the parameters of the 'A' command: - Check if the radii are equal, which indicates a circle. - If the radii differ, consider the shape as an ellipse. Shape Identification and Classification: Use the geometric properties to classify the shape: - For polygons, identify specific types like rectangles, kites, and trapezoids based on their properties. Pay special attention to the number of sides and the relationships between them. Consider the entire path as a single shape to ensure accurate classification. - For arcs, determine if the shape is a circle or an ellipse based on the radii. Options Selection and Interpretation: Choose the most appropriate shape from the given options. Consider multiple interpretations of the path data, especially when multiple 'M' commands are present, to ensure accurate classification. If the path represents multiple shapes, prioritize the most complex or relevant shape. Ambiguity Resolution: In cases where the path data could represent multiple shapes, provide a rationale for selecting the most complex or relevant shape. Consider the context and any additional information that might influence the classification. Visual Verification: If possible, visualize the path to confirm the identified shape. This step can help resolve any remaining ambiguities and ensure the accuracy of the classification. Iterative Refinement: If the initial classification is uncertain, revisit the analysis steps to refine the identification. Consider alternative interpretations and re-evaluate the geometric properties. Contextual Considerations: Take into account any contextual information or additional data that might influence the shape classification, especially in ambiguous cases. Provide your answer by selecting the correct option and enclosing it within <answer> and </answer> tags. Example: - SVG Path: 'path d="M 8.10,55.86 L 1.74,25.57 L 12.08,23.40 L 18.44,53.69 L 8.10,55.86"' - Analysis: The path forms a closed quadrilateral with opposite sides parallel and equal, indicating a rectangle. - Answer: <answer>H</answer> - SVG Path: 'path d="M 16.33,5.98 A 8.87,8.87 275.02 1,0 14.78,23.64 A 8.87,8.87 275.02 1,0 16.33,5.98"/' - Analysis: The path uses elliptical arcs with equal radii, forming a closed loop, indicating a circle. - Answer:	0.830

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
our	Given the following SVG path data: "input" and options, identify the geometric shape it represents and provide ONLY the name of the shape as the 'target'. Task Requirements: 1. Count the instructions in the SVG path 2. Judge the shape of the graphic according to the judgment criteria 3. Provide the exact name of the shape as output. You need to count how many instructions L are in the SVG path: Problem Rule Application: 1. Visualize the path data to understand the overall structure. 2. Find out whether there is instruction A in the instruction. If so, determine whether it is a circle or a sector according to the number of instructions A. If not, determine how many sides it is 3. For polygons, pay attention to the number of edges to identify the shape. The following are the number of instructions corresponding to different shapes: # - triangle: 3 L # - rectangle: 4 L # - hexagon: 6 L # - pentagon: 5 L # - octagon: 8 L # - heptagon: 7 L # - kite: 4 L # - line: 2 L # - circle: 2 or more A # - sector: 1 A Judgment criteria:(Please strictly abide by this rule) No need to pay attention to "M" instructions !! First identify whether there is an instruction "A" in the SVG path. If so, first determine whether it is a circle or a sector. !! If there is no instruction "A", determine the number of sides of the polygon based on the instruction "L". A polygon with n sides requires n "L" instructions.(Please strictly abide by this rule)	0.822

Table 7: Comparison of Optimized Prompts for **Causal Judgment tasks**, including CoT, APE, PromptAgent, and our method

Approach	Optimized Prompt	Accuracy
CoT	Task: Respond to inquiries about causal attribution by identifying the key causes and their contributions to the outcome. Follow the steps below to ensure accurate and clear reasoning: Steps to Analyze Causation: 1. Identify Key Entities: Read the question carefully and highlight the specific entities or factors being discussed. 2. Determine Relevant Causes: Analyze the context to identify immediate and incidental causes contributing to the outcome. - Immediate causes: Directly lead to the outcome. - Incidental causes: Indirectly influence the outcome but may still contribute. 3. Evaluate Interactions: Consider how multiple causes might interact to produce the observed effect (e.g., synergy or independent contributions). 4. Provide the Answer: Clearly state the primary and secondary causes, as well as their roles in creating the outcome. Avoid unsupported assumptions. Use this structured reasoning approach to analyze each inquiry and provide a clear and logical explanation.	0.678

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
APE	For each situation, decide if the result was caused deliberately or not. If the individual or party behind the event was aware of the potential result and chose to go ahead, select 'Yes'. If they didn't intend the result to happen, even if they knew it could possibly occur, select 'No'.	0.756
PromptAgent	When addressing questions about causal attribution, ensure a comprehensive analysis by considering both individual and collective actions that contribute to an outcome. Clearly differentiate between necessary and sufficient causes, and recognize that multiple causes can simultaneously contribute to an outcome. Emphasize the importance of understanding both general and specific intentions, especially when outcomes are unintended. Define "intentional" actions as those where the actor or group had control over maintaining or altering the conditions necessary for the outcome, even if the specific result was not desired. Address scenarios where unintended consequences arise from intentional actions, and provide answers that reflect a nuanced understanding of how different elements interact to produce a result. Use diverse examples to illustrate key concepts like "direct causation," "simultaneity," and "unintended consequences," ensuring a balanced consideration of necessary and sufficient causes. Simplify complex scenarios by breaking them down into clear, manageable components, and provide definitions or examples of key terms to guide your analysis. Additionally, clarify definitions of key terms such as "necessary," "sufficient," "intentional," and "unintended consequences" to ensure precise understanding. Highlight the importance of interactions between multiple causes, especially in complex scenarios, and offer strategies for analyzing scenarios where simultaneity is crucial. Explore the nuances of intentional actions and unintended consequences more deeply, and encourage the use of diverse examples to illustrate different aspects of causation. Pay special attention to the role of individual actions in maintaining necessary conditions and the distinction between collective and individual causation. Emphasize that in collective decision-making, the outcome can be intentional if it aligns with the group's goals, even if individual members disagree. Reinforce the distinction between necessary and sufficient causes, ensuring the model understands that necessary causes alone do not determine the outcome. Clarify that following a protocol does not remove intentionality if the outcome aligns with organizational priorities. Highlight that intentionality can be attributed if the outcome was a foreseeable consequence of the action, regardless of individual opposition.	0.846

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
Our	Task Requirements Determine whether a given event (cause) directly leads to another event (effect). Assess the causal relationship based on logical reasoning, ensuring a clear and definitive answer. The final output must be only "Yes" or "No", strictly adhering to the required format. Problem Rule Application Identify the cause and effect within the question. Assess necessity: Determine if the cause is essential for the effect to occur. Evaluate causation: If the cause did not happen, would the effect still occur? If the effect only happens when the cause is present, then the cause directly leads to the effect. If the effect can still happen independently, then the relationship is not causal. Judgment Criteria Direct Causation: If the cause directly leads to the effect and is a necessary condition, answer "Yes". If the effect would not have occurred without the cause, answer "Yes". Example: "Dropping a glass caused it to shatter." → Yes. No Direct Causation: If the effect can occur without the cause, answer "No". If the cause is only correlated but not necessary, answer "No". Example: "Wearing a red shirt caused the stock market to rise." → No. Verification Process: Check whether the absence of the cause results in the absence of the effect. Ensure logical consistency in the causal assessment.	0.880

Table 8: Comparison of Optimized Prompts for **Epistemic task**, including CoT, APE, PromptAgent, and our method

Approach	Optimized Prompt	Accuracy
CoT	Task: Analyze the logical relationship between a given premise and hypothesis. Your goal is to determine if the premise guarantees the truth of the hypothesis. Choose one of the following answers: 'entailment' or 'non-entailment'. **Steps to Follow:** 1. **Understand the Premise and Hypothesis**: Carefully read the premise and hypothesis to identify the key information in both statements. 2. **Analyze the Logical Relationship**: Determine whether the information in the premise confirms the truth of the hypothesis. - If the premise logically supports and guarantees the hypothesis, choose 'entailment'. - If the premise does not confirm the hypothesis, or if there is uncertainty, choose 'non-entailment'. 3. **Provide the Answer**: Based on your analysis, output the correct answer ('entailment' or 'non-entailment'). Use this step-by-step approach for all premise and hypothesis pairs to ensure accurate reasoning.	0.855
APE	Determine whether the hypothesis is directly implied by the premise or not. If the premise's statement is a direct claim or conviction of the individual mentioned in the hypothesis, choose 'entailment'. However, if the premise is formed on the belief or supposition of someone other than the subject in the hypothesis, opt for 'non-entailment'.	0.888

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
PromptAgent	Determine the relationship between two sentences by evaluating whether the first sentence provides direct or logically implied evidence for the second. Choose from the options 'entailment' or 'non-entailment'. Consider the following: - Entailment: The first sentence directly or through logical implication confirms the truth of the second sentence, even if it involves a chain of beliefs or perceptions, as long as the chain logically supports the hypothesis. - Non-entailment: The first sentence does not confirm the truth of the second sentence, often involving unsupported assumptions, beliefs, or suspicions that do not logically lead to the hypothesis. Guidelines for Analysis: 1. Clarify Belief Chains and Logical Implications: Understand how belief chains work and when they logically support the hypothesis. Pay attention to verbs indicating beliefs, assumptions, or suspicions (e.g., "thinks," "assumes," "suspects") versus those indicating direct evidence (e.g., "learns," "knows," "remembers"). Consider how these verbs interact in belief chains and what they imply about the subject's own beliefs. 2. Evaluate Direct and Implied Evidence: Determine if the premise provides direct or logically implied evidence for the hypothesis, considering how belief chains can logically support the hypothesis. Recognize that indirect beliefs about another person's recognition can imply one's own belief about a situation, especially when the belief chain is logical and straightforward. 3. Consider Perspective and Source of Information: Note any differences in perspective or source of information (e.g., who remembers or assumes something) and how these perspectives contribute to the logical implication of the hypothesis. 4. Conduct a Comprehensive Analysis: Use a step-by-step approach to ensure all relevant details and logical implications are considered in the analysis. Balance the emphasis on direct evidence with the recognition of logical implications from indirect beliefs. Example: Premise: "Charlotte thinks that Richard recognizes that a boy is standing in a pool getting splashed with water." Hypothesis: "Charlotte thinks that a boy is standing in a pool getting splashed with water." Options: (A) entailment (B) non-entailment Analysis: 1. Understanding the Premise: The premise indicates that Charlotte thinks Richard recognizes a specific situation involving a boy in a pool. 2. Understanding the Hypothesis: The hypothesis states that Charlotte thinks a boy is in a pool getting splashed with water. 3. Assessing the Relationship: The premise implies that Charlotte has a belief about the situation (through Richard's recognition), which logically supports the hypothesis. Charlotte's belief about Richard's recognition suggests she also believes in the situation's occurrence. 4. Conclusion: The relationship is one of entailment because Charlotte's belief about Richard's recognition logically implies her belief in the situation. Therefore, the correct answer is: <answer>A</answer> Identify the relation between the following premises and hypotheses, choosing from the options 'entailment' or 'non-entailment'. At the end, show the answer option bracketed between <answer> and </answer>.	0.916

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
our	Task Requirements: Analyze a given premise (primary sentence) and determine whether it fully supports the truth of a hypothesis (subsequent sentence). Classify the relationship as either "Entailment" or "Non-Entailment" based on the logical and factual connections between the two. Provide the classification only as the final output. Problem Rule Application: Entailment: The premise explicitly confirms the hypothesis with clear, direct evidence. No additional information, assumptions, or interpretations are required to validate the hypothesis. Non-Entailment: The premise does not fully or explicitly confirm the hypothesis. If there is ambiguity, uncertainty, or missing logical links, label it as Non-Entailment. Judgment Criteria: (Strictly follow these rules) Language of Uncertainty: Words like "assumes," "believes," "thinks," "feels," "suspects" indicate subjectivity and should not be considered definitive proof. These terms suggest a possibility rather than an explicit factual connection. Specific vs. General Statements: A specific premise (e.g., mentioning a "full face mask") does not necessarily contradict a general hypothesis (e.g., referencing a "mask" in general). However, if the premise is too general to confirm the specific claim, classify as Non-Entailment. Objective Reasoning: Only use the logical and factual ties within the given statements. Do not rely on external knowledge, assumptions, or interpretations unless directly supported by the premise. Decision Process: Determine whether the premise fully supports the hypothesis without needing extra inference → Entailment. If the premise only partially supports or fails to confirm the hypothesis → Non-Entailment.	0.893

Table 9: Comparison of Optimized Prompts for **Temporal task** including CoT, APE, PromptAgent, and our method.

Approach	Optimized Prompt	Accuracy
CoT	Your task is to determine the available time slot for an event, based on a schedule of occupied times. Follow these steps to ensure accuracy: **Steps to Identify Free Time Slots:** 1. **List Occupied Periods**: Organize all occupied time slots in chronological order. 2. **Find Gaps**: Identify gaps between the occupied periods where no activities are scheduled. 3. **Check Constraints**: Ensure that the free time slots fall within operational constraints (e.g., facility closing times). 4. **Select the Slot**: Choose the correct free time slot that satisfies all criteria. **Output Result Format:** - Present the selected free time slot in a clear format, such as "Xpm to Ypm" or "Xam to Yam". Use this step-by-step method to ensure that the identified time slot is accurate and does not overlap with any occupied periods.	0.989

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
APE	Identify the period when the individual was unnoticed and had the possibility to visit the specified place before its closing time.	0.994
PromptAgent	Analyze the timeline of events to determine possible time frames during which certain events could have occurred, even if they were not explicitly observed. Start by constructing a comprehensive timeline, clearly listing all observed and unobserved time slots. Identify gaps where the subject is unobserved, ensuring these gaps fit within any given constraints, such as opening and closing times. Emphasize the importance of constraints by verifying them after identifying potential gaps. Use a step-by-step reasoning approach to systematically evaluate all available information, and include a final review to check for potential errors or overlooked details before finalizing the answer. Define key terms like "unobserved" and "constraints" to ensure clarity in the task requirements. Provide examples to illustrate the reasoning process and expected output format, guiding the model in analyzing timelines and identifying possible time frames for unobserved events. Additionally, incorporate a checklist to ensure all steps are followed, and highlight common pitfalls to avoid during the analysis. Finally, include a summary of the reasoning process to reinforce understanding and ensure the model's conclusions are well-supported. To further enhance the model's performance, include additional examples that cover a wider range of scenarios and constraints, such as overlapping time slots or multiple constraints. Provide explicit guidance on handling complex constraints and ambiguous information. Incorporate interactive feedback mechanisms to help the model learn from mistakes and improve over time. Ensure the prompt is concise and focused, avoiding unnecessary repetition while maintaining clarity and comprehensiveness. Additionally, introduce a section for handling exceptions or unusual cases, offering strategies for dealing with incomplete or conflicting data. This will help the model adapt to a broader range of real-world scenarios and improve its robustness in timeline analysis tasks.	0.984

(Continued on next page)

(Continued from previous page)

Approach	Optimized Prompt	Accuracy
our	**Task Requirements:** Determine the possible time period during which an event could have occurred, based on a detailed schedule of occupied times. Your goal is to identify the correct time slot that fits all the provided criteria without any overlap. **Problem Rule Explanation:** 1. Analyze the schedule to identify all time slots during which the person is occupied. 2. Determine the available time slots by identifying gaps between these occupied periods. 3. Consider any additional constraints, such as closing times, that may limit the available time slots. **Problem Rule Application:** - List all the occupied time slots chronologically. - Identify gaps between these occupied slots where the person is free. - Ensure that the free time slots do not conflict with constraints like closing times. **Result Verification:** - Confirm that the identified time slot is completely free and adheres to any constraints. - Double-check against all occupied periods to ensure there is no overlap. - Avoid selecting time slots that are partially occupied or overlap with any scheduled activities. **Output Result Format:** - Present the correct time slot in a straightforward manner, using the format "Xpm to Ypm" or "Xam to Yam" as appropriate. - Ensure the output is clear and free of any extraneous symbols or text. **Common Mistakes to Avoid:** - Do not include time slots that extend beyond the closing time of the facility. - Avoid selecting time slots that overlap with any scheduled activities. - Ensure the selected time slot is entirely free and does not partially overlap with any occupied period. **General Rules and Analysis:** - Identify all occupied periods and list them chronologically. - Look for gaps between these periods where the person is not scheduled for any activity. - Verify that these gaps fall within any operational constraints, such as closing times. - Ensure the selected time slot is entirely free and does not overlap with any occupied periods. By following these guidelines, you can accurately determine the available time slot for the event in question. Avoid errors by ensuring that the selected time slot is entirely free and does not overlap with any occupied periods.	0.993

Table 10: Optimized Prompts Without Difficulty Scaling in Synthetic Data

Tasks	Optimized Prompt	Accuracy																																												
Penguins	You are provided with two tables containing data about penguins and giraffes. Your task is to focus solely on the giraffe data to answer a specific question regarding the tallest giraffe. **Penguin Data:** <table><tr><th>Name</th><th>Age</th><th>Height (cm)</th><th>Weight (kg)</th></tr><tr><td>Louis</td><td>7</td><td>50</td><td>11</td></tr><tr><td>Bernard</td><td>5</td><td>80</td><td>13</td></tr><tr><td>Vincent</td><td>9</td><td>60</td><td>11</td></tr><tr><td>Gwen</td><td>8</td><td>70</td><td>15</td></tr><tr><td>James</td><td>12</td><td>90</td><td>12</td></tr></table> **Giraffe Data:** <table><tr><th>Name</th><th>Age</th><th>Height (cm)</th><th>Weight (kg)</th></tr><tr><td>Jody</td><td>5</td><td>430</td><td>620</td></tr><tr><td>Gladys</td><td>10</td><td>420</td><td>590</td></tr><tr><td>Marian</td><td>2</td><td>310</td><td>410</td></tr><tr><td>Donna</td><td>9</td><td>440</td><td>650</td></tr></table> **Task Requirements:** 1. Identify the tallest giraffe based on the height provided in the Giraffe Data table. 2. Provide the weight of the tallest giraffe in kilograms. **Problem Rule Explanation:** - Review the height values for each giraffe listed in the Giraffe Data table. - Compare these height values to determine which giraffe is the tallest. **Problem Rule Application:** - Examine the height values for the giraffes: - Jody: 430 cm - Gladys: 420 cm - Marian: 310 cm - Donna: 440 cm - Identify that Donna is the tallest giraffe at 440 cm. - Retrieve the corresponding weight of Donna, which is 650 kg. **Result Verification:** - Ensure that you have considered all entries in the Giraffe Data table. - Confirm that the weight you provide corresponds to the giraffe identified as the tallest. **Output Result Format:** - Provide your answer in the following format: - "Weight of the tallest giraffe: [Weight in kg]" **Example Output:** - "Weight of the tallest giraffe: 650"	Name	Age	Height (cm)	Weight (kg)	Louis	7	50	11	Bernard	5	80	13	Vincent	9	60	11	Gwen	8	70	15	James	12	90	12	Name	Age	Height (cm)	Weight (kg)	Jody	5	430	620	Gladys	10	420	590	Marian	2	310	410	Donna	9	440	650	0.732
Name	Age	Height (cm)	Weight (kg)																																											
Louis	7	50	11																																											
Bernard	5	80	13																																											
Vincent	9	60	11																																											
Gwen	8	70	15																																											
James	12	90	12																																											
Name	Age	Height (cm)	Weight (kg)																																											
Jody	5	430	620																																											
Gladys	10	420	590																																											
Marian	2	310	410																																											
Donna	9	440	650																																											
Geometry	" Given the following input: ""input"", you must provide ONLY the correct value for the 'target'. **Rules:** 1. Do NOT provide any explanations. 2. Do NOT provide any sentences, text, or words other than the 'target' value. 3. The answer must be the exact value contained in the ""target"" and any unauthorized additions are prohibited."	0.681																																												

(Continued on next page)

(Continued from previous page)

Tasks	Optimized Prompt	Accuracy
Object Counting	**Task Requirements:** - Determine the total number of fruits by accurately identifying and counting each type listed in the question. **Problem Rule Explanation:** - The task involves listing and counting each fruit mentioned. - Each fruit should be counted as one unless a specific quantity is provided. **Problem Rule Application:** - Carefully read through the list to identify all items that are fruits. - Count each fruit once unless otherwise specified with a different quantity. - Avoid including any non-fruit items or miscounting due to misinterpretation of the list. **Result Verification:** - Re-examine the list to ensure all fruits have been correctly identified and counted. - Verify that the total count reflects only the fruits listed, with no errors in inclusion or exclusion. **Output Result Format:** - Provide the total number of fruits in both word and numeral forms, such as: ["ten", "10"]. - Ensure the output is clear and free from special symbols or formatting errors."	0.538
Causal Judgment	Analyze the scenario to determine if the described action directly caused the outcome. Provide a definitive 'Yes' or 'No' answer based on a logical assessment of the causal relationship as described in the scenario. **Problem Rule Explanation:** A causal relationship exists when an action directly leads to an outcome without other factors influencing the result. The outcome should not occur without the action. Avoid assumptions and base your analysis solely on the information provided. **Problem Rule Application:** - Identify the key action and the resulting outcome within the scenario. - Determine if the outcome is a direct result of the action, ensuring no additional factors are at play. - Evaluate whether the outcome would still occur without the initial action, focusing on the explicit roles, responsibilities, and conditions mentioned. - Avoid external assumptions and concentrate on the details provided in the scenario. **Result Verification:** - Confirm that the action directly causes the outcome, with no interference from other factors. - Ensure the outcome logically follows from the action, considering the context and rules provided. - Review the scenario for any overlooked details that could affect the causal link, ensuring a comprehensive analysis. **Output Result Format:** - Answer 'Yes' if the action directly causes the outcome, with the outcome being a direct consequence of the action. - Answer 'No' if there is no direct causal relationship or if other factors could have contributed to the outcome.	0.684

(Continued on next page)

(Continued from previous page)

Tasks	Optimized Prompt	Accuracy
Temporal	**Task Requirements:** Determine the available time slots for an unscheduled activity within a given daily schedule, ensuring these slots do not conflict with scheduled events and comply with any facility operating hours. **Problem Rule Explanation:** 1. Review the entire schedule to identify all events and their specific time frames. 2. Identify gaps between these events or after the last scheduled event to find potential time slots for the unscheduled activity. 3. Consider any additional constraints, such as facility operating hours, to ensure the proposed time slot is feasible. **Problem Rule Application:** - List all scheduled events with their respective time frames. - Identify gaps between these events or available time after the last scheduled event. - Ensure that the identified time slots comply with any additional constraints, like facility operating hours. **Result Verification:** - Confirm that the identified time slots do not overlap with any scheduled events. - Verify that the time slots fall within the facility's operating hours. **Output Result Format:** Present the time range in a clear and concise format, such as "Xpm to Ypm" or "Xam to Yam", ensuring clarity and precision. **Example Application:** Given the schedule: - Breakfast: 8am to 9am - Business meeting: 9am to 11am - Art gallery: 11am to 1pm - Lunch: 1pm to 2pm - Cinema: 3pm to 5pm - Dinner party: 6pm to 8pm - Gym closes at 10pm Determine the available time for the gym: - Identify the gaps: 5pm to 6pm and 8pm to 10pm. - Ensure these time slots do not overlap with scheduled events and are within the gym's operating hours. - The correct answer is "5pm to 6pm" and "8pm to 10pm", as they fit within the gym's operating hours and do not overlap with any scheduled events.	

C Prompt Self-improvement

C.1 Error Analysis: Prompt Template

1. **Misinterpretation of the Question:** The model may have misinterpreted the question, focusing on the structure of the data rather than the specific request for the height of the tallest penguin. This could lead to confusion and an irrelevant answer.

2. **Inattention to Numerical Data:** The model might have overlooked the numerical values provided in the table, failing to recognize that it needed to compare the heights of the penguins to determine the tallest one.

3. **Irrelevant Output Generation:** The answer "rectangle" does not relate to the context of the question. This suggests that the model may have generated a response based on unrelated patterns or associations rather than the specific data presented.

4. **Lack of Contextual Understanding:** The model may not have fully grasped the context of the data table, leading to a failure in recognizing that the question was asking for a specific value derived from the table.

5. **Failure to Process Tabular Data:** The model might struggle with processing tabular data effectively, which can lead to incorrect conclusions or irrelevant outputs when asked to analyze such formats.

C.2 Improvement Recommendation: Prompt Template

- Clarify the question to emphasize the need for numerical comparison.
- Provide explicit instructions to focus on extracting specific values from the data.
- Ensure the model is trained to recognize and process tabular data more effectively.
- Avoid ambiguity in the question to prevent misinterpretation of the request.

C.3 Improvement Recommendations: Revised Prompt

Here is a table of penguins with their respective attributes:

name	age	height (cm)	weight (kg)
------	-----	-------------	-------------

Louis	7	50	11
-------	---	----	----

Bernard	5	80	13
---------	---	----	----

Vincent	9	60	11
---------	---	----	----

Gwen	8	70	15
------	---	----	----

****Task Requirements:****

Analyze the table to identify the tallest penguin based on the 'Height (cm)' column.

****Problem Rule Explanation:****

Compare the numerical values in the 'Height (cm)' column to determine the maximum height.

****Problem Rule Application:****

Identify the maximum value from the 'Height (cm)' column.

****Result Verification:****

Ensure that the identified height corresponds to the tallest penguin.

Table 11: Examples of Generated Data for BIG-Bench tasks

Tasks	Generated Data
Causal Judgment	In a small town, there is a bakery that makes the best pastries. Every morning, the bakery opens at 7:00 am, and a line of customers forms outside. The bakery owner has a rule that only one person can enter at a time to maintain order. One day, two customers, Alice and Bob, arrive at the same time. Alice follows the rule and waits outside, but Bob ignores the rule and enters the bakery while another customer is still inside. The bakery becomes overcrowded, and a shelf of pastries falls over, ruining the day's batch. Did Bob cause the pastries to be ruined?
Geometry	Generated question: This SVG path element This SVG path element <code><path d="M 50.00,30.00 L 66.18,35.09 L 72.45,50.00 L 66.18,64.91 L 50.00,70.00 L 33.82,64.91 L 27.55,50.00 L 33.82,35.09 L 50.00,30.00"/></code>
Object Counting	Generated question: I have two violins, a drum, a piano, a flute, and a trumpet. Additionally, I have a cat, a rabbit, a dog, a chicken, and a goat. How many musical instruments do I have?
Epistemic	Premise: Olivia suspects that Ethan recognizes that a group of musicians gather in a park, tuning their instruments as the sun sets behind the city skyline. Hypothesis: Ethan recognizes that a group of musicians gather in a park, tuning their instruments as the sun sets behind the city skyline.
Temporal	Today, Alex attended several events. Between what times could he have gone to the gym? We know that: Alex had breakfast at 8am. He attended a meeting from 9am to 11am. He was seen at the art gallery from 11am to 1pm. He had lunch with friends from 1pm to 2pm. He was at the cinema from 2pm to 4pm. He visited his grandmother from 4pm to 6pm. The gym closes at 10pm. Between what times could Alex have gone to the gym?
Penguins	Generated question: Here is a table where the first line is a header and each subsequent line is a penguin: name, age, height (cm), weight (kg) Louis, 7, 50, 11 Bernard, 5, 80, 13 Vincent, 9, 60, 11 Gwen, 8, 70, 15 For example: the age of Louis is 7, the weight of Gwen is 15 kg, the height of Bernard is 80 cm. What is the age of Vincent?

E Detailed proof of Theorem 3.3

1. Surrogate domination. Because the surrogate loss upper-bounds the 0–1 loss point-wise,

$$\sup_{\psi \in \Psi} \mathbb{E}_{q_\psi}[\mathbb{1}\{f(p, \tilde{x}) \neq \tilde{y}\}] \leq \sup_{\psi \in \Psi} \mathbb{E}_{q_\psi}[L(f(p, \tilde{x}), \tilde{y})] = \sup_{\psi \in \Psi} J(\psi),$$

where we abbreviated $J(\psi) := \mathbb{E}_{q_\psi} L(f(p, \tilde{x}), \tilde{y})$.

2. Reduce to the near-optimal generator. Let ψ^* be any generator that ε -maximises the regularised objective,

$$\psi^* = \arg \max_{\psi \in \Psi} \left\{ J(\psi) - \lambda^{-1} R(\psi) \right\} \quad \text{s.t.} \quad J(\psi^*) - \lambda^{-1} R(\psi^*) \geq \sup_{\psi \in \Psi} (J(\psi) - \lambda^{-1} R(\psi)) - \varepsilon.$$

Rearranging, $J(\psi) \leq J(\psi^*) + \lambda^{-1}(R(\psi) - R(\psi^*)) + \varepsilon$ for every ψ , hence

$$\sup_{\psi \in \Psi} J(\psi) \leq J(\psi^*) + \varepsilon.$$

3. Bound the hard generator via KL. Applying the risk-alignment inequality to ψ^* ,

$$J(\psi^*) \leq \mathbb{E}_{(x,y) \sim P} L(f(p, x), y) + \lambda^{-1} R(\psi^*).$$

4. Sample–population substitution. With probability at least $1 - \delta$ over the draw of the training set,

$$\mathbb{E}_{(x,y) \sim P} L(f(p, x), y) \leq \frac{1}{n} \sum_{i=1}^n L(f(p, x_i), y_i) + q(|\mathcal{P}|, n, \delta).$$

Combine. Chaining 1-4 we obtain

$$\sup_{\psi \in \Psi} \mathbb{E}_{q_\psi} \mathbb{1}\{f(p, \tilde{x}) \neq \tilde{y}\} \leq \frac{1}{n} \sum_{i=1}^n L(f(p, x_i), y_i) + \lambda^{-1} R(\psi^*) + \varepsilon + q(|\mathcal{P}|, n, \delta),$$

which is exactly the bound claimed in Theorem 3.3. $\square$