

FORMATTING INSTRUCTIONS FOR ICLR 2026

CONFERENCE SUBMISSIONS

Anonymous authors

Paper under double-blind review

ABSTRACT

Large Language Models (LLMs) have demonstrated remarkable capabilities but often face challenges with tasks requiring sophisticated reasoning. While Chain-of-Thought (CoT) prompting significantly enhances reasoning, it indiscriminately generates lengthy reasoning steps for all queries, leading to substantial computational costs and inefficiency, especially for simpler inputs. To address this critical issue, we introduce AdaCoT (Adaptive Chain-of-Thought), a novel framework enabling LLMs to adaptively decide when to invoke CoT. AdaCoT framed adaptive reasoning as a Pareto optimization problem that seeks to balance model performance with the costs associated with CoT invocation (both frequency and computational overhead). We propose a reinforcement learning (RL) based method, specifically utilizing Proximal Policy Optimization (PPO), to dynamically control the CoT triggering decision boundary by adjusting penalty coefficients, thereby allowing the model to determine CoT necessity based on implicit query complexity. A key technical contribution is Selective Loss Masking (SLM), designed to counteract decision boundary collapse during multi-stage RL training, ensuring robust and stable adaptive triggering. Experimental results demonstrate that AdaCoT successfully navigates the Pareto frontier, achieving substantial reductions in CoT usage for queries not requiring elaborate reasoning. For instance, on our production traffic testset, AdaCoT reduced CoT triggering rates to as low as 3.18% and decreased average response tokens by 69.06%, while maintaining high performance on complex tasks. This substantial token decrease directly translates to a significant reduction in inference computational load. AdaCoT pioneers adaptive CoT triggering, offering a practical and principled solution for developing more efficient, responsive, and cost-effective LLMs, particularly crucial for interactive and resource-sensitive applications.

1 INTRODUCTION

Large Language Models (LLMs) have garnered substantial attention due to their remarkable ability to encode extensive world knowledge from vast corpora Ouyang et al. (2022), enabling impressive performance across diverse tasks such as question answering, creative writing, and summarization. Despite these successes, LLMs often demonstrate limitations in tasks requiring sophisticated reasoning, including solving complex mathematical problems and intricate coding puzzles. To mitigate this, recent methodologies have employed Chain-of-Thought (CoT) prompting Wei et al. (2022), which encourages models to explicitly generate step-by-step reasoning prior to producing final answers. This approach significantly enhances the reasoning capability of models, even achieving human-expert levels in certain domains Xu et al. (2025a); Jaech et al. (2024); Guo et al. (2025); Seed et al. (2025).

However, employing CoT prompting also poses critical challenges during inference Chen et al. (2024b). Specifically, it substantially increases the number of tokens generated, even for simple queries that do not benefit from elaborate reasoning, such as straightforward arithmetic questions. This indiscriminate token expense consequently raises deployment costs and reduces inference efficiency. Ideally, a model should adaptively determine when detailed reasoning is necessary. For instance, simple queries like "What is 1+1?" should be answered immediately without additional reasoning steps, whereas more complex queries require deeper and step-by-step reasoning. An adaptive strategy would thus optimize token usage, balancing cost-efficiency with response quality.

Recently, a few efforts have made attempts towards this direction. These approaches can be broadly categorized into three main directions. (1) Incorporate length penalties or brevity rewards during the reinforcement learning (RL) stage to encourage shorter, more concise reasoning paths Arora & Zanette (2025); Team et al. (2025); Hou et al. (2025); Luo et al. (2025b); Shen et al. (2025); Aggarwal & Welleck (2025). (2) Restructure CoT outputs through post-processing such as iterative summarization or pruning Zhang et al. (2025); Yan et al. (2025); Aytes et al. (2025); Yang et al. (2025); Xu et al. (2025b); Kang et al. (2025); Xia et al. (2025); Liu et al. (2024). (3) Employ explicit user instructions or hand-crafted selection mechanisms to control the use of CoT Han et al. (2024); Renze & Guven (2024); Chen et al. (2024a); Yu et al. (2024); Munkhbat et al. (2025); Pan et al. (2024). Despite their contributions, they mainly focus on monotonic reasoning reduction, failing to account for the nuanced variability in query complexity, i.e., treating simple and difficult prompts adaptively. Moreover, they lack a principled optimization framework to guide balancing response quality against deployment cost.

To address these limitations, we introduce AdaCoT (Adaptive Chain-of-Thought), a novel approach grounded in formal mathematical analysis. Our key insight is framing adaptive reasoning as a multi-objective optimization problem with two competing goals: maximizing response accuracy and minimizing deployment costs. Specifically, we formalize this balance through Pareto optimization, seeking optimal trade-offs between reasoning complexity and inference efficiency. Such a mathematical framework provides clear theoretical grounding for dynamically adapting CoT triggering based on query complexity.

Leveraging this formalization, we propose an RL-based strategy explicitly designed around the Pareto optimization framework, enabling effective control of the model’s decision boundary for initiating CoT prompting. During training, the RL agent dynamically assesses the complexity of incoming user queries to determine the necessity and extent of reasoning steps. By carefully designing the reward function to incorporate penalty coefficients, we encourage the RL agent to seek solutions along the Pareto frontier, explicitly optimizing trade-offs between response accuracy and token expenditure. This structured exploration enables the model to effectively discern when detailed reasoning is beneficial, thereby systematically enhancing inference efficiency and significantly reducing deployment costs.

The proposed AdaCoT framework delivers substantial benefits in LLM operational efficiency. By empowering models to selectively engage CoT, AdaCoT can reduce triggering rates to as low as 3.18% and cut average response tokens by 69.1% in production settings. This significant reduction in computational load is achieved while maintaining strong performance on 15 widely-adopted benchmarks. For example, AdaCoT can achieve a 62.8% average score using only a 53.3% CoT rate, closely rivaling the 65.0% score of a model that always employs CoT. These improvements directly translate to more cost-effective and responsive LLM systems.

2 THE ADACoT FRAMEWORK

We introduce AdaCoT, a framework enabling LLMs to dynamically invoke Chain-of-Thought (CoT) based on query complexity. This allows for a more rational allocation of computational resources, focusing tokens on complex reasoning while avoiding overhead for simple queries. The core idea is to frame the adaptive decision as a Pareto optimization problem balancing two competing goals: maximizing response accuracy and minimizing deployment cost. We propose an RL-based strategy where a policy model learns to assess query complexity and decide on CoT invocation. This policy is optimized using a reward signal designed to reflect the Pareto trade-off between performance and token efficiency, teaching the model to allocate reasoning effort adaptively for high-quality, cost-effective responses.

2.1 ADAPTIVE REASONING AS A PARETO OPTIMIZATION PROBLEM

We formulate the adaptive reasoning challenge as a Pareto optimization problem, aiming to simultaneously maximize model performance and minimize CoT usage. Let $\mathcal{D} = \{(x_i, y_i)\}_{i=1}^N$ be a dataset of query-response pairs, where x_i is the input query and y_i is the ground truth response. Let f_θ be an LLM parameterized by θ , and let $r_\theta(x_i)$ be the response generated by the model for input x_i .

To CoT usage is measured by the CoT triggering rate $T(\theta)$, defined as the proportion of responses that include reasoning:

$$T(\theta) = \frac{1}{N} \sum_{i=1}^N \mathbf{1}[\text{HasReasoning}(r_\theta(x_i))] \quad (1)$$

where $\mathbf{1}[\cdot]$ is the indicator function and $\text{HasReasoning}(\cdot)$ determines if a response contains explicit CoT steps (e.g., non-empty content within `<think>...</think>` tags).

On the other hand, model performance $P(\theta)$ is defined as the average score on a set of evaluation metrics:

$$P(\theta) = \frac{1}{M} \sum_{j=1}^M \text{Score}_j(\theta) \quad (2)$$

where M is the number of evaluation instances/metrics and $\text{Score}_j(\theta)$ is the model’s score on the j -th evaluation.

The core of our adaptive reasoning challenge is a multi-objective optimization problem (MOP) with two competing goals: maximizing model performance $P(\theta)$ and minimizing CoT usage $T(\theta)$. Formally, we seek to find solutions for:

$$\max_{\theta} (P(\theta), -T(\theta)) \quad (3)$$

The solution to such a problem is not a single point but a set of non-dominated solutions known as the Pareto frontier. Each point on this frontier represents an optimal trade-off, where improving one objective (e.g., performance) necessitates a compromise in the other (e.g., increased CoT usage).

A common and practical method to find points on this frontier is linear scalarization. We combine the two objectives into a single scalar objective function $\mathcal{J}(\theta)$:

$$\theta^* = \arg \max_{\theta} \mathcal{J}(\theta) \quad \text{where} \quad \mathcal{J}(\theta) = \lambda_P P(\theta) - \lambda_T T(\theta) \quad (4)$$

Here, $\lambda_P > 0$ and $\lambda_T > 0$ are scalar weights that control the relative importance of performance versus CoT reduction. By systematically varying the ratio of these weights, we can trace different points along the Pareto frontier. The AdaCoT framework is designed to operationalize this exploration, using its RL-based training to find models θ^* that correspond to different, desirable points on this performance-cost trade-off curve.

2.2 TRAINING PIPELINE FOR ADACoT

The AdaCoT training pipeline integrates supervised fine-tuning (SFT) as an initialization phase, followed by multi-stage reinforcement learning (RL) to refine the adaptive CoT triggering behavior.

2.2.1 DATA PREPARATION AND SUPERVISED FINE-TUNING (SFT) AS WARM-UP

To provide the model with an initial understanding of when CoT might be beneficial, we perform a data preparation stage. This is achieved by leveraging an auxiliary model, guided by a set of pre-defined principles (e.g., query complexity, expected reasoning depth, domain; see Appendix B). In our implementation, we use an internal 15B-parameter model to generate these annotations; however, the framework is model-agnostic and can be instantiated using any sufficiently capable LLM with basic instruction-following abilities. Queries are labeled as either likely benefiting from CoT or likely suitable for a direct answer. This principled, automated labeling process is more consistent and scalable than manual annotation.

The SFT training data is then structured as follows: For queries labeled as benefiting from CoT, responses retain the full reasoning process: `<think>reasoning_steps</think>answer`. For queries labeled as not requiring CoT, responses omit explicit reasoning but maintain structural consistency: `<think></think>answer`. This SFT stage serves as a "warm-up", equipping the model with a foundational capability to distinguish between these two response styles. The consistent use of `<think></think>` tags is crucial for maintaining response format integrity.

2.2.2 REINFORCEMENT LEARNING FOR ADAPTIVE CoT CONTROL

The RL stage is pivotal for fine-tuning AdaCoT’s adaptive reasoning capabilities. We design a reward function $R(x, r)$ for an input query x and generated response r :

$$R(x, r) = R_{\text{base}}(x, r) - \alpha_1 \cdot P_{\text{miss}}(x, r) - \alpha_2 \cdot P_{\text{over}}(x, r) - \gamma \cdot P_{\text{fmt}}(r) \quad (5)$$

where $R_{\text{base}}(x, r)$ is the base reward reflecting response quality, $P_{\text{miss}}(x, r)$ is a binary penalty for reasoning omission, $P_{\text{over}}(x, r)$ is a binary penalty for reasoning overuse, $P_{\text{fmt}}(r)$ is a binary penalty for format errors, and $\alpha_1, \alpha_2, \gamma$ are non-negative penalty coefficients. By adjusting α_1 and α_2 , we steer AdaCoT towards different CoT triggering decision boundaries, allowing exploration of the Pareto frontier.

2.2.3 ADDRESSING DECISION BOUNDARY COLLAPSE WITH SELECTIVE LOSS MASKING

A significant challenge in multi-stage RL, particularly when fine-tuning on specialized datasets with skewed CoT distributions (e.g., mathematical datasets where CoT is almost always beneficial), is the risk of the adaptive CoT triggering capability becoming unstable or collapsing. The model might revert to a homogeneous behavior, either always or never triggering CoT, thereby losing the nuanced decision-making learned in earlier, more balanced training stages. We term this phenomenon decision boundary collapse. This is particularly problematic if the final RL stage has significant bias, as it can lead to the model almost completely losing its adaptive triggering capability.

To address decision boundary collapse, AdaCoT incorporates Selective Loss Masking (SLM). SLM aims to preserve the CoT triggering ratio and distribution established during SFT or prior RL stages. It achieves this by selectively masking the loss contribution from the pivotal “decision token” during RL phases prone to distribution bias. This decision token is defined as the token immediately succeeding the `<think>` tag.

The modified policy gradient loss under SLM, $\mathcal{L}_{\text{SLM}}$, is computed by excluding the loss component associated with this decision token:

$$\mathcal{L}_{\text{SLM}} = \sum_{k \neq k_{\text{decision}}} \ell_k \quad (6)$$

where ℓ_k is the original loss component for the k -th token, and k_{decision} is the index of the decision token.

3 EXPERIMENTS

We conducted extensive experiments to evaluate the AdaCoT framework, focusing on its ability to navigate the performance-cost trade-off, the effectiveness of its adaptive triggering mechanism, and its impact on inference efficiency. This section details our experimental setup, presents the main results, and analyzes the findings.

3.1 EXPERIMENTAL SETUP

For our base model, we utilized our internal 15B/150B parameter Mixture-of-Experts (MoE) Jacobs et al. (1991); Shazeer et al. (2017) model. The AdaCoT post-training process comprised an initial SFT stage, followed by a two-stage RL procedure: first, a Mathematics-Focused RL stage (RL-Math) concentrated on complex, rule-verifiable problems, and second, a General Domain RL stage (RL-General) which incorporated broader data and a trained reward model. We compared our **AdaCoT RL Models** (Exp1-Exp4) against several baselines: a **Full CoT SFT Baseline** (SFT model always generating CoT), a **Full CoT RL Baseline** (RL model derived from the Full CoT SFT, always generating CoT), a **No CoT SFT Baseline** (SFT model never generating CoT), a **No CoT RL Baseline** (RL model derived from the No CoT SFT, never generating CoT), and our **AdaCoT SFT Model** (our model after only the SFT stage, also referred to as Adaptive SFT Model).

The SFT and RL training datasets were constructed to cover a diverse range of domains, including mathematics, reasoning, professional disciplines (e.g., law, medicine), dialogue, creative writing, and general knowledge question answering. Both SFT and RL data were labeled for CoT necessity using the principle-guided assessment detailed in Appendix B. In the SFT dataset, approximately 67% of the samples were labeled as requiring CoT, while in the RL dataset, this proportion

was around 40%. During SFT, queries identified as not requiring CoT were formatted with empty `<think></think>` tags. In the RL-Math stage, which is particularly prone to decision boundary collapse, we employed Selective Loss Masking (SLM), as described in Section 2.2.3. For the RL-General stage, we applied penalties according to Equation 5, systematically varying the α_1 and α_2 coefficients to explore different points on the Pareto frontier. Proximal Policy Optimization (PPO) Schulman et al. (2017) was used for all RL policy updates.

For evaluation, we used 15 diverse open-source benchmark datasets to assess overall performance, measured by the average score. To balance internal iteration efficiency with evaluation accuracy, some of these datasets underwent up-sampling or down-sampling, or the number of inference runs per test sample was adjusted (with the final score being an average over multiple inferences). These datasets include LiveBench White et al. (2025), MMLU Pro Wang et al. (2025), SuperGPQA Du et al. (2025), GPQA Rein et al. (2024), Chinese SimpleQA He et al. (2024b), SimpleQA Wei et al. (2024), AIME24 & AIME25, MATH Hendrycks et al. (2021), OlympiadBench He et al. (2024a), SweBench Agentless Jimenez et al. (2024), LiveCodeBench Jain et al. (2025), KOR-Bench Ma et al. (2024), ProcBench Fujisawa et al. (2024), and SysBench Qin et al. (2024). The detailed per-dataset scores, which form the basis for our average score calculations, are presented in Appendix A (Table 4). To specifically assess CoT triggering decisions on typical user queries, we curated a high-quality balanced test set of 1000 prompts. These prompts were labeled for CoT necessity using the same principle-guided assessment as our SFT/RL training data and subsequently underwent manual verification to ensure label accuracy. On this set, we report CoT Triggering Accuracy, F1-score, Precision, and Recall, where the positive class indicates that CoT is required. Other metrics included the CoT triggering rate on the benchmark datasets and the average response token num on production setting.

3.2 RESULTS AND ANALYSIS

Our results demonstrate AdaCoT’s ability to effectively control CoT invocation, leading to improved efficiency while maintaining strong performance.

3.2.1 PARETO FRONTIER ANALYSIS

We trained four variants of our AdaCoT RL model (Exp1-Exp4) by varying the penalty coefficients α_1 (for missing CoT) and α_2 (for overusing CoT). The specific coefficients were: Exp1 ($\alpha_1 = 0.1, \alpha_2 = 0.3$), Exp2 ($\alpha_1 = 0.2, \alpha_2 = 0.3$), Exp3 ($\alpha_1 = 0.3, \alpha_2 = 0.3$), and Exp4 ($\alpha_1 = 0.3, \alpha_2 = 0.1$). The format error penalty γ was consistently set to 1.0. Figure 1 illustrates the average score plotted against the CoT triggering rate for these models and the baselines, based on the average performance across our 15 benchmark datasets (detailed in Appendix Table 4). The No CoT SFT baseline achieved an average score of 43.6% with 0% CoT usage, while the No CoT RL baseline improved this to 47.7% at 0% CoT. The AdaCoT SFT Model (Adaptive SFT) registered a 57.1% average score at a 61.3% CoT rate.

The AdaCoT RL models trace a compelling Pareto frontier. AdaCoT RL Exp1 (43.1% CoT, 59.7% score) and AdaCoT RL Exp2 (53.3% CoT, 62.8% score) demonstrate significant performance gains over the AdaCoT SFT model while operating at lower or comparable CoT rates. Notably, AdaCoT RL Exp2 achieves a 62.8% average score, approaching the Full CoT RL baseline (65.0% score, 100% CoT) with nearly half the CoT usage. As we increase the CoT triggering rate, AdaCoT RL Exp3 (65.4% CoT, 64.3% score) and AdaCoT RL Exp4 (67.7% CoT, 64.4% score) further push performance, closely rivaling the Full CoT RL baseline’s score but with approximately 32-35% less CoT invocation. Moreover, these results indicate that despite the fixed CoT triggering labels within the SFT/RL data, adjusting the combination of penalty coefficients during the RL phase enables the final RL model to learn triggering strategies that transcend these initial labeling constraints.

This highlights AdaCoT’s effectiveness in navigating the trade-off between performance and CoT cost. However, it is also observable that while the AdaCoT RL models achieve substantial efficiency gains and define a superior Pareto curve compared to simpler baselines, they encounter a slight performance bottleneck when their triggering rates are pushed higher. Specifically, even the highest-performing adaptive models (Exp3 and Exp4, with scores of 64.3% and 64.4% respectively) do not surpass the average score of the Full CoT RL baseline (65.0%). This suggests that while AdaCoT excels at reducing CoT for a vast majority of queries without compromising much on average

performance, and indeed offers a better score-to-cost ratio, the absolute peak average performance achieved by a model specialized to always use CoT (Full CoT RL) remains marginally higher. This indicates that the adaptive mechanism, while highly effective, might not perfectly replicate or exceed the performance of an always-on CoT strategy in every single scenario contributing to the average, thus not fully crossing this specific optimal boundary for maximum average score. This could be due to the inherent complexities of learning a universally optimal triggering heuristic or the RL optimization finding a balance that prioritizes the significant cost savings available across the query distribution.

To further isolate the benefit of the RL stage, we conducted an analysis where the model was forced to follow the CoT labels generated during data annotation. On the moderately difficult MMLU-pro dataset, this static policy resulted in a score of 82.7%, a noticeable drop from the 85.2% achieved by the Full-CoT model. In contrast, our adaptive RL models (Exp1 and Exp2) scored 84.0% and 83.5% respectively, while using CoT far less often. This suggests that during RL, the model learns a more nuanced, self-aware triggering policy that is better adapted to its own capabilities than the initial static labels can provide. A detailed breakdown is available in Appendix F.3.

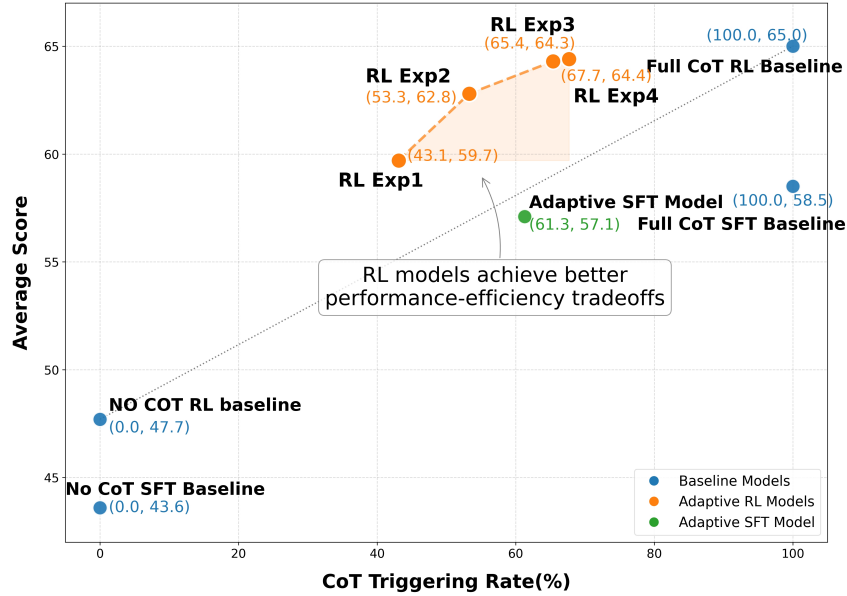


Figure 1: Average Score vs. CoT Triggering Rate across 15 widely-adopted benchmarks. Blue points represent baseline models. The green point is the AdaCoT SFT model. Orange points represent AdaCoT RL models trained with varying penalty coefficients, forming an improved Pareto frontier (indicated by the orange dashed line and shaded region) over the baselines. The dotted line connects the No CoT RL baselines to the Full CoT RL baseline, illustrating a simpler trade-off curve.

3.2.2 ADAPTIVE CoT TRIGGERING PERFORMANCE AND ABLATION STUDIES ON DAILY-USE QUERIES

We evaluated the CoT triggering capabilities of AdaCoT at various training stages using our curated 1000-prompt daily-use test set. Table 1 presents these results, which include an ablation study for SLM and an assessment of the meta-reasoning strategy (discussed further in Section 4.2) at the SFT stage.

The AdaCoT SFT model itself provides a strong baseline for adaptive triggering. The results clearly show that the RL-Math stage without SLM suffers from severe decision boundary collapse: the model defaults to triggering CoT (Recall=1.0) but with poor precision (0.503) and consequently low overall accuracy (0.506). Applying SLM during the RL-Math stage effectively preserves the adaptive capability learned during SFT, maintaining high precision (0.938) and achieving significantly better accuracy (0.813). The final AdaCoT RL models (Exp1-4, emerging from the RL-General

Table 1: CoT triggering performance on the 1000 daily-use prompt test set across different AdaCoT stages and configurations (positive class: requires CoT). RL-Math is the Mathematics-Focused RL stage; RL-General refers to the final models (Exp1-4).

Model Stage / Variant	Accuracy	F1-Score	Recall	Precision
AdaCoT SFT Model	0.795	0.750	0.616	0.959
AdaCoT RL Model (Exp1 - RL-General)	0.657	0.484	0.322	0.975
AdaCoT RL Model (Exp2 - RL-General)	0.816	0.814	0.804	0.823
AdaCoT RL Model (Exp3 - RL-General)	0.809	0.789	0.716	0.879
AdaCoT RL Model (Exp4 - RL-General)	0.678	0.535	0.370	0.963
RL-Math (without SLM)	0.506	0.669	1.000	0.503
RL-Math (with SLM)	0.813	0.781	0.670	0.938
AdaCoT SFT Model (with Meta-Reasoning)	0.858	0.840	0.762	0.935

stage) demonstrate how adjusting the RL penalty coefficients (α_1, α_2) allows for fine-tuning of the decision boundary. AdaCoT RL Model Exp2, for example, achieves a well-balanced F1-score of 0.814. The incorporation of a meta-reasoning strategy at the SFT stage also shows a notable improvement in triggering performance, a point elaborated in Section 4.2.

3.2.3 RESPONSE LENGTH REDUCTION AND EFFICIENCY GAINS

The adaptive reasoning enabled by AdaCoT translates into significant reductions in computational costs. Table 2 shows the average response length and CoT triggering rates for AdaCoT RL Model Exp2 (selected for its balanced performance on the daily-use set and strong average benchmark performance) when applied to our production traffic testset, which reflects natural, unfiltered user query distributions.

Table 2: Average response token num (with reduction noted) and CoT triggering rate on production traffic testset for AdaCoT RL Model Exp2 vs. Full CoT RL Baseline.

Platform	Model / Mode	Avg. Response Tokens	CoT Triggering Rate
Mobile	Full CoT RL Baseline	377.18	100.00%
	AdaCoT RL Model Exp2 (Adaptive)	116.70 (↓69.1%)	3.18%
PC	Full CoT RL Baseline	1376.31	100.00%
	AdaCoT RL Model Exp2 (Adaptive)	405.25 (↓70.6%)	12.50%

As evidenced, AdaCoT RL Model Exp2 achieves very low CoT triggering rates in a production setting (3.18% on mobile devices, 12.50% on PCs). This dramatic reduction from the 100% CoT usage of a non-adaptive model translates directly into substantial computational savings.

4 DISCUSSION AND FUTURE WORK

4.1 DESIGN CONSIDERATIONS AND LIMITATIONS

AdaCoT offers a pragmatic approach to adaptive reasoning by combining principle-guided initial data labeling with RL-based optimization of the CoT decision boundary. This methodology was chosen to circumvent inherent challenges in purely autonomous CoT trigger learning, such as information asymmetry in assessing counterfactual benefits and the difficulty of quantifying quality degradation from CoT omission, particularly for subjective tasks.

While AdaCoT is a promising initial step, several limitations exist. The optimal CoT triggering strategy is relative to the base model’s capabilities, necessitating recalibration for different LLMs. Our current binary CoT invocation (on/off) simplifies a continuous spectrum of reasoning depths and styles, potentially limiting nuance. Domain generalization remains a challenge, as CoT necessity

can vary significantly across knowledge areas, and the framework currently lacks personalization for user verbosity preferences. Moreover, the initial principle-guided labeling requires continuous refinement. Our Pareto analysis (Section 3.2.1) also indicates that while AdaCoT significantly improves efficiency and nears the peak average performance of specialized always-on CoT models, a small performance gap persists, highlighting the difficulty for adaptive mechanisms to achieve absolute maximum performance across all query types.

Acknowledging the limitations of the current framework, we anticipate that future research will offer valuable critiques and further refine these initial explorations. Areas warranting deeper investigation include more granular control over reasoning, such as adaptive reasoning length where models dynamically adjust verbosity, or more nuanced triggering mechanisms beyond a simple binary decision. We believe such continued efforts by the community are crucial for developing more sophisticated and efficient reasoning strategies, potentially addressing the observed performance gap while maximizing efficiency and enhancing nuanced control over LLM reasoning.

4.2 META-REASONING FOR REASONING DECISIONS

We explored an innovative "meta-reasoning" approach during the SFT stage to enhance AdaCoT's CoT triggering decisions. This involved the model first generating a brief internal assessment of the input query's complexity before deciding whether to proceed with full CoT, as illustrated by the response structures in Figure 2. Incorporating this strategy into the AdaCoT SFT model led to a notable improvement in CoT triggering performance on our daily-use test set: the F1-score increased from 0.750 to 0.840 (Table 1). This result suggests that explicitly prompting the model to first assess query complexity can significantly enhance its subsequent decision-making regarding CoT invocation.

An interesting and serendipitous discovery with the meta-reasoning SFT model was an emergent capability for user-prompt controllability over CoT. Users could, to some extent, influence whether the model engaged in CoT by including explicit cues in their prompts about the desired level of reasoning (e.g., "please think step-by-step" to encourage CoT, or "give a direct answer" to discourage it). While this controllability was not perfectly accurate across all scenarios, it points towards a promising avenue for developing more interactive and user-guided reasoning systems (further illustrative examples are provided in Appendix D).

Despite these benefits, the explicit meta-reasoning step inherently increases the number of tokens generated for every query, as the model first articulates its complexity assessment. Considering the very low CoT triggering rates observed for AdaCoT in production environments (e.g., 3.18% on mobile traffic for AdaCoT RL Model Exp2, as shown in Table 2), the cumulative token cost of these additional meta-reasoning steps would become substantial. Therefore, while acknowledging its potential for improving decision accuracy and enabling user control, we did not adopt this explicit meta-reasoning as the default for subsequent RL experiments due to this efficiency trade-off. Nevertheless, it highlights an important direction for future research, which might explore more token-efficient methods for incorporating such meta-reasoning, perhaps through implicit learning mechanisms or parallel processing of the complexity assessment.

4.3 PRESERVING PEAK PERFORMANCE WITH ADACoT

A critical question is whether adaptive reasoning limits a model's maximum performance. We investigated this using System Prompts (SPs), integrated into AdaCoT's SFT and RL training to control reasoning behavior (e.g., "Always Reason SP," "Never Reason SP"). During SFT, a small portion of data was augmented with SPs, and target responses were modified for compliance. In RL, a fraction of training prompts included SPs, with rewards penalizing deviations from explicit SP instructions, ensuring robust adherence (details in Appendix C). Our focus here is using the "Always Reason SP" to assess AdaCoT's performance ceiling.

Instructing AdaCoT RL models to always generate CoT via this SP allowed direct comparison against the Full CoT RL Baseline on our 15 benchmark datasets. As shown in Table 3, AdaCoT RL models in this forced "Always Reason" mode achieved average scores that were highly competitive with, and in instances like AdaCoT RL Model Exp4 (65.7) and Exp2 (65.3), slightly surpassed the Full CoT RL Baseline (65.0). This demonstrates a key strength: AdaCoT's adaptive

training, aimed at optimizing the performance-cost Pareto frontier, does not inherently restrict the model’s peak reasoning capabilities when comprehensive reasoning is explicitly demanded. While our Pareto analysis (Section 3.2.1) noted a slight gap in average scores when models operate adaptively, these SP-controlled results affirm that AdaCoT offers efficiency without sacrificing potential high-end performance.

Table 3: Performance of AdaCoT RL models under ”Always Reason” System Prompt vs. Full CoT RL Baseline, demonstrating preservation of peak performance. Metrics are averaged across the 15 benchmark datasets.

Model Variant	Avg. Score	CoT Triggering Rate
Full CoT RL Baseline	65.0	100%
AdaCoT RL Model Exp1 (w/ Always CoT SP)	64.8	100%
AdaCoT RL Model Exp2 (w/ Always CoT SP)	65.3	100%
AdaCoT RL Model Exp3 (w/ Always CoT SP)	64.9	100%
AdaCoT RL Model Exp4 (w/ Always CoT SP)	65.7	100%

An interesting secondary observation arose from the SFT stage concerning long-form generation (up to 32,000 tokens). AdaCoT SFT models, when directed by ”Always Reason SP,” exhibited fewer instances of premature output truncation and were less prone to undesirable generative loops compared to a standard Full CoT SFT baseline. We hypothesize this improvement stems from AdaCoT’s diverse SFT data, which includes many non-CoT examples (formatted as `<think></think>answer`). This results in a shorter average training sample length, potentially providing the End-of-Sequence token a stronger learning signal (average EOS proportion: 0.000239 for AdaCoT SFT vs. 0.000215 for Full CoT SFT). A more robust EOS representation could foster more coherent, well-terminated lengthy outputs, a promising area for future investigation.

5 RELATED WORK

Chain-of-Thought (CoT) reasoning Wei et al. (2022) and its derivatives, such as structured thoughts (trees Yao et al. (2023), graphs Besta et al. (2024)) and diverse path sampling Wang et al. (2023), have significantly advanced LLM capabilities. However, CoT’s verbosity and cost Sui et al. (2025); Chen et al. (2024b) are major drawbacks, and its utility varies: it is crucial for complex tasks Wu et al. (2025b); Lee et al. (2025) but can be less effective than direct answers for simple queries Ma et al. (2025a). This creates a critical cost-performance trade-off. Existing efficiency research has primarily focused on reducing reasoning length rather than deciding *whether* to invoke CoT. Common strategies include using RL with brevity rewards Arora & Zanette (2025); Team et al. (2025); Hou et al. (2025); Luo et al. (2025b); Shen et al. (2025); Aggarwal & Welleck (2025); restructuring or compressing CoT content via summarization or learned representations Zhang et al. (2025); Yan et al. (2025); Aytes et al. (2025); Yang et al. (2025); Xu et al. (2025b); Kang et al. (2025); Xia et al. (2025); Liu et al. (2024); and employing explicit instructions or selection mechanisms Han et al. (2024); Renze & Guven (2024); Chen et al. (2024a); Yu et al. (2024); Munkhbat et al. (2025); Pan et al. (2024). While effective, these methods and other alternatives like model merging Ma et al. (2025b); Wu et al. (2025a); Luo et al. (2025a) generally do not equip a single model to dynamically decide on CoT invocation based on query nature. AdaCoT distinctively addresses this gap by focusing on adaptive triggering. By framing the task as a Pareto optimization problem and using RL to control the decision boundary, AdaCoT enables a single model to perform nuanced, context-dependent CoT invocation, filling a crucial gap towards truly efficient and versatile LLMs.

6 CONCLUSION

We introduced AdaCoT, a framework for adaptive CoT. By framing the task as a Pareto optimization problem solved via reinforcement learning, AdaCoT reduces CoT usage on simple queries while preserving high performance on complex tasks. Our Selective Loss Masking technique ensures robust training. Pioneering adaptive triggering over prior work on length compression, AdaCoT offers a principled path toward more efficient and responsive LLMs, especially for interactive applications.

ETHICS STATEMENT

The authors have read and adhered to the ICLR Code of Ethics. Our work introduces AdaCoT, a framework aimed at improving the computational efficiency of Large Language Models (LLMs) by adaptively triggering Chain-of-Thought (CoT) reasoning. The primary societal benefit of this research is the potential reduction in computational resources, energy consumption, and inference latency associated with deploying powerful LLMs, making them more accessible and sustainable.

We acknowledge several ethical considerations. First, the training and evaluation of our framework involved proprietary models and datasets, including anonymized production traffic. All data was handled in strict accordance with privacy policies, and no personally identifiable information was used. Second, like all LLMs, the models used in this study may inherit and perpetuate biases present in their training data. While our work focuses on the mechanism of reasoning rather than content generation, future research should investigate the fairness implications of adaptive reasoning, ensuring that the decision to use or omit CoT does not disproportionately affect certain user groups or query types. Finally, while our framework is designed for efficiency, the underlying LLM is a general-purpose technology. We have not identified any new risks of misuse introduced specifically by our adaptive mechanism beyond those inherent to powerful language models in general. We are committed to the responsible development and deployment of AI technologies.

REPRODUCIBILITY STATEMENT

We are committed to ensuring the reproducibility of our work. The core methodology of the AdaCoT framework, including the Pareto optimization formulation, the reinforcement learning reward function, and the Selective Loss Masking (SLM) technique, is described in detail in Section 2. The experimental setup, including the multi-stage training pipeline, baseline models, and evaluation metrics, is outlined in Section 3.1.

For evaluation, we used 15 publicly available benchmarks, which are listed in Section 3.1. A comprehensive breakdown of scores for all models on each of these benchmarks is provided in Appendix A (Table 4) to allow for detailed comparison. Key implementation details are provided in the appendices, including the principles for data annotation (Appendix B), examples of system prompts used for controllability experiments (Appendix C), and training/inference hyperparameters (Appendix E).

Due to proprietary considerations, the specific base models and internal datasets used for training cannot be publicly released. We acknowledge that this prevents a direct replication of our exact numerical results. However, we have provided a thorough description of our methods and a detailed analysis on public benchmarks to enable the research community to reimplement and validate the core concepts and effectiveness of the AdaCoT framework on other models and systems.

REFERENCES

- Pranjal Aggarwal and Sean Welleck. L1: Controlling how long a reasoning model thinks with reinforcement learning. *arXiv preprint arXiv:2503.04697*, 2025.
- Daman Arora and Andrea Zanette. Training language models to reason efficiently. *arXiv preprint arXiv:2502.04463*, 2025.
- Simon A Aytes, Jinheon Baek, and Sung Ju Hwang. Sketch-of-thought: Efficient llm reasoning with adaptive cognitive-inspired sketching. *arXiv preprint arXiv:2503.05179*, 2025.
- Maciej Besta, Nils Blach, Ales Kubicek, Robert Gerstenberger, Michal Podstawski, Lukas Gianinazzi, Joanna Gajda, Tomasz Lehmann, Hubert Niewiadomski, Piotr Nyczyk, et al. Graph of thoughts: Solving elaborate problems with large language models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pp. 17682–17690, 2024.
- Qiguang Chen, Libo Qin, Jiaqi Wang, Jingxuan Zhou, and Wanxiang Che. Unlocking the capabilities of thought: A reasoning boundary framework to quantify and optimize chain-of-thought. In *Advances in Neural Information Processing Systems*, volume 37, pp. 54872–54904, 2024a.

- Xingyu Chen, Jiahao Xu, Tian Liang, Zhiwei He, Jianhui Pang, Dian Yu, Linfeng Song, Qiuzhi Liu, Mengfei Zhou, Zhuosheng Zhang, et al. Do not think that much for $2+3=?$ on the overthinking of o1-like llms. *arXiv preprint arXiv:2412.21187*, 2024b.
- Xinrun Du, Yifan Yao, Kaijing Ma, Bingli Wang, Tianyu Zheng, King Zhu, Minghao Liu, Yiming Liang, Xiaolong Jin, Zhenlin Wei, et al. Supergpqa: Scaling llm evaluation across 285 graduate disciplines. *arXiv preprint arXiv:2502.14739*, 2025.
- Ippei Fujisawa, Sensho Nobe, Hiroki Seto, Rina Onda, Yoshiaki Uchida, Hiroki Ikoma, Pei-Chun Chien, and Ryota Kanai. Procbench: Benchmark for multi-step reasoning and following procedure. *arXiv preprint arXiv:2410.03117*, 2024.
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- Tingxu Han, Zhenting Wang, Chunrong Fang, Shiyu Zhao, Shiqing Ma, and Zhenyu Chen. Token-budget-aware llm reasoning. *arXiv preprint arXiv:2412.18547*, 2024.
- Chaoqun He, Renjie Luo, Yuzhuo Bai, Shengding Hu, Zhen Thai, Junhao Shen, Jinyi Hu, Xu Han, Yujie Huang, Yuxiang Zhang, et al. Olympiadbench: A challenging benchmark for promoting agi with olympiad-level bilingual multimodal scientific problems. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 3828–3850, 2024a.
- Yancheng He, Shilong Li, Jiaheng Liu, Yingshui Tan, Weixun Wang, Hui Huang, Xingyuan Bu, Hangyu Guo, Chengwei Hu, Boren Zheng, et al. Chinese simpleqa: A chinese factuality evaluation for large language models. *arXiv preprint arXiv:2411.07140*, 2024b.
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. Measuring mathematical problem solving with the math dataset. In *NeurIPS*, 2021.
- Bairu Hou, Yang Zhang, Jiabao Ji, Yujian Liu, Kaizhi Qian, Jacob Andreas, and Shiyu Chang. Thinkprune: Pruning long chain-of-thought of llms via reinforcement learning. *arXiv preprint arXiv:2504.01296*, 2025.
- Robert A Jacobs, Michael I Jordan, Steven J Nowlan, and Geoffrey E Hinton. Adaptive mixtures of local experts. *Neural computation*, 3(1):79–87, 1991.
- Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, et al. Openai o1 system card. *arXiv preprint arXiv:2412.16720*, 2024.
- Naman Jain, King Han, Alex Gu, Wen-Ding Li, Fanjia Yan, Tianjun Zhang, Sida Wang, Armando Solar-Lezama, Koushik Sen, and Ion Stoica. Livecodebench: Holistic and contamination free evaluation of large language models for code. In *The Thirteenth International Conference on Learning Representations*, 2025.
- Carlos E Jimenez, John Yang, Alexander Wettig, Shunyu Yao, Kexin Pei, Ofir Press, and Karthik R Narasimhan. SWE-bench: Can language models resolve real-world github issues? In *The Twelfth International Conference on Learning Representations*, 2024.
- Yu Kang, Xianghui Sun, Liangyu Chen, and Wei Zou. C3ot: Generating shorter chain-of-thought without compromising effectiveness. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pp. 24312–24320, 2025.
- Ayeong Lee, Ethan Che, and Tianyi Peng. How well do llms compress their own chain-of-thought? a token complexity approach. *arXiv preprint arXiv:2503.01141*, 2025.
- Tengxiao Liu, Qipeng Guo, Xiangkun Hu, Cheng Jiayang, Yue Zhang, Xipeng Qiu, and Zheng Zhang. Can language models learn to skip steps? In *Advances in Neural Information Processing Systems*, volume 37, pp. 45359–45385, 2024.

- Haotian Luo, Haiying He, Yibo Wang, Jinluan Yang, Rui Liu, Naiqiang Tan, Xiaochun Cao, Dacheng Tao, and Li Shen. Adarl: From long-cot to hybrid-cot via bi-level adaptive reasoning optimization. *arXiv preprint arXiv:2504.21659*, 2025a.
- Haotian Luo, Li Shen, Haiying He, Yibo Wang, Shiwei Liu, Wei Li, Naiqiang Tan, Xiaochun Cao, and Dacheng Tao. O1-pruner: Length-harmonizing fine-tuning for o1-like reasoning pruning. *arXiv preprint arXiv:2501.12570*, 2025b.
- Kaijing Ma, Xinrun Du, Yunran Wang, Haoran Zhang, Zhoufutu Wen, Xingwei Qu, Jian Yang, Jiaheng Liu, Minghao Liu, Xiang Yue, et al. Kor-bench: Benchmarking language models on knowledge-orthogonal reasoning tasks. *arXiv preprint arXiv:2410.06526*, 2024.
- Wenjie Ma, Jingxuan He, Charlie Snell, Tyler Griggs, Sewon Min, and Matei Zaharia. Reasoning models can be effective without thinking. *arXiv preprint arXiv:2504.09858*, 2025a.
- Xinyin Ma, Guangnian Wan, Runpeng Yu, Gongfan Fang, and Xinchao Wang. Cot-valve: Length-compressible chain-of-thought tuning. *arXiv preprint arXiv:2502.09601*, 2025b.
- Tergel Munkhbat, Namgyu Ho, Seo Hyun Kim, Yongjin Yang, Yujin Kim, and Se-Young Yun. Self-training elicits concise reasoning in large language models. *arXiv preprint arXiv:2502.20122*, 2025.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. In *Advances in neural information processing systems*, volume 35, pp. 27730–27744, 2022.
- Jiabao Pan, Yan Zhang, Chen Zhang, Zuozhu Liu, Hongwei Wang, and Haizhou Li. Dynathink: Fast or slow? a dynamic decision-making framework for large language models. *arXiv preprint arXiv:2407.01009*, 2024.
- Yanzhao Qin, Tao Zhang, Yanjun Shen, Wenjing Luo, Haoze Sun, Yan Zhang, Yujing Qiao, Weipeng Chen, Zenan Zhou, Wentao Zhang, et al. Sysbench: Can large language models follow system messages? *arXiv preprint arXiv:2408.10943*, 2024.
- David Rein, Betty Li Hou, Asa Cooper Stickland, Jackson Petty, Richard Yuanzhe Pang, Julien Dirani, Julian Michael, and Samuel R Bowman. Gpqa: A graduate-level google-proof q&a benchmark. In *First Conference on Language Modeling*, 2024.
- Matthew Renze and Erhan Guven. The benefits of a concise chain of thought on problem-solving in large language models. In *2024 2nd International Conference on Foundation and Large Language Models (FLLM)*, pp. 476–483. IEEE, 2024.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- ByteDance Seed, Yufeng Yuan, Yu Yue, Mingxuan Wang, Xiaochen Zuo, Jiaze Chen, Lin Yan, Wenyan Xu, Chi Zhang, Xin Liu, et al. Seed-thinking-v1. 5: Advancing superb reasoning models with reinforcement learning. *arXiv preprint arXiv:2504.13914*, 2025.
- Noam Shazeer, Azalia Mirhoseini, Krzysztof Maziarczyk, Andy Davis, Quoc Le, Geoffrey Hinton, and Jeff Dean. Outrageously large neural networks: The sparsely-gated mixture-of-experts layer. *arXiv preprint arXiv:1701.06538*, 2017.
- Yi Shen, Jian Zhang, Jieyun Huang, Shuming Shi, Wenjing Zhang, Jiangze Yan, Ning Wang, Kai Wang, and Shiguo Lian. Dast: Difficulty-adaptive slow-thinking for large reasoning models. *arXiv preprint arXiv:2503.04472*, 2025.
- Yang Sui, Yu-Neng Chuang, Guanchu Wang, Jiamu Zhang, Tianyi Zhang, Jiayi Yuan, Hongyi Liu, Andrew Wen, Hanjie Chen, Xia Hu, et al. Stop overthinking: A survey on efficient reasoning for large language models. *arXiv preprint arXiv:2503.16419*, 2025.

- Kimi Team, Angang Du, Bofei Gao, Bowei Xing, Changjiu Jiang, Cheng Chen, Cheng Li, Chenjun Xiao, Chenzhuang Du, Chonghua Liao, et al. Kimi k1. 5: Scaling reinforcement learning with llms. *arXiv preprint arXiv:2501.12599*, 2025.
- Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc V Le, Ed H Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. Self-consistency improves chain of thought reasoning in language models. In *The Eleventh International Conference on Learning Representations*, 2023.
- Yubo Wang, Xueguang Ma, Ge Zhang, Yuansheng Ni, Abhranil Chandra, Shiguang Guo, Weiming Ren, Aaran Arulraj, Xuan He, Ziyang Jiang, et al. Mmlu-pro: A more robust and challenging multi-task language understanding benchmark. In *Advances in Neural Information Processing Systems*, volume 37, pp. 95266–95290, 2025.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. Chain-of-thought prompting elicits reasoning in large language models. In *Advances in neural information processing systems*, volume 35, pp. 24824–24837, 2022.
- Jason Wei, Nguyen Karina, Hyung Won Chung, Yunxin Joy Jiao, Spencer Papay, Amelia Glaese, John Schulman, and William Fedus. Measuring short-form factuality in large language models. *arXiv preprint arXiv:2411.04368*, 2024.
- Colin White, Samuel Dooley, Manley Roberts, Arka Pal, Benjamin Feuer, Siddhartha Jain, Ravid Shwartz-Ziv, Neel Jain, Khalid Saifullah, Sreemanti Dey, Shubh-Agrawal, Sandeep Singh Sandha, Siddhartha Venkat Naidu, Chinmay Hegde, Yann LeCun, Tom Goldstein, Willie Neiswanger, and Micah Goldblum. Livebench: A challenging, contamination-limited LLM benchmark. In *The Thirteenth International Conference on Learning Representations*, 2025.
- Han Wu, Yuxuan Yao, Shuqi Liu, Zehua Liu, Xiaojin Fu, Xiongwei Han, Xing Li, Hui-Ling Zhen, Tao Zhong, and Mingxuan Yuan. Unlocking efficient long-to-short llm reasoning with model merging. *arXiv preprint arXiv:2503.20641*, 2025a.
- Yuyang Wu, Yifei Wang, Tianqi Du, Stefanie Jegelka, and Yisen Wang. When more is less: Understanding chain-of-thought length in llms. *arXiv preprint arXiv:2502.07266*, 2025b.
- Heming Xia, Yongqi Li, Chak Tou Leong, Wenjie Wang, and Wenjie Li. Tokenskip: Controllable chain-of-thought compression in llms. *arXiv preprint arXiv:2502.12067*, 2025.
- Fengli Xu, Qianyu Hao, Zefang Zong, Jingwei Wang, Yunke Zhang, Jingyi Wang, Xiaochong Lan, Jiahui Gong, Tianjian Ouyang, Fanjin Meng, et al. Towards large reasoning models: A survey of reinforced reasoning with large language models. *arXiv preprint arXiv:2501.09686*, 2025a.
- Silei Xu, Wenhao Xie, Lingxiao Zhao, and Pengcheng He. Chain of draft: Thinking faster by writing less. *arXiv preprint arXiv:2502.18600*, 2025b.
- Yuchen Yan, Yongliang Shen, Yang Liu, Jin Jiang, Mengdi Zhang, Jian Shao, and Yueting Zhuang. Infthythink: Breaking the length limits of long-context reasoning in large language models. *arXiv preprint arXiv:2503.06692*, 2025.
- Chenxu Yang, Qingyi Si, Yongjie Duan, Zheliang Zhu, Chenyu Zhu, Zheng Lin, Li Cao, and Weiping Wang. Dynamic early exit in reasoning models. *arXiv preprint arXiv:2504.15895*, 2025.
- Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Tom Griffiths, Yuan Cao, and Karthik Narasimhan. Tree of thoughts: Deliberate problem solving with large language models. In *Advances in neural information processing systems*, volume 36, pp. 11809–11822, 2023.
- Ping Yu, Jing Xu, Jason E Weston, and Ilia Kulikov. Distilling system 2 into system 1. In *The First Workshop on System-2 Reasoning at Scale, NeurIPS’24*, 2024.
- Jintian Zhang, Yuqi Zhu, Mengshu Sun, Yujie Luo, Shuofei Qiao, Lun Du, Da Zheng, Huajun Chen, and Ningyu Zhang. Lightthinker: Thinking step-by-step compression. *arXiv preprint arXiv:2502.15589*, 2025.

A BENCHMARK DATASET DETAILS AND ANALYSIS

This appendix provides descriptions for the benchmark datasets used in our evaluation and an analysis of the experimental results on these individual datasets. The scores presented in Table 4 form the basis for this analysis. For each dataset, we discuss the performance of baseline models and the AdaCoT variants. We also highlight any counter-intuitive results or observations that conflict with the primary motivation of achieving optimal performance with adaptive CoT, offering potential explanations such as evaluation volatility, inherent limitations of the base model, or aspects of the post-training process that may not be fully optimized for every scenario.

Table 4: Detailed scores on benchmark datasets. "TR" denotes reasoning triggering rate (%).

Dataset	nocot SFT baseline		nocot RL baseline		fullcot SFT baseline		fullcot RL baseline		Adaptive SFT Model		Adaptive RL Model Exp1		Adaptive RL Model Exp2		Adaptive RL Model Exp3		Adaptive RL Model Exp4	
	TR	Score	TR	Score	TR	Score	TR	Score	TR	Score	TR	Score	TR	Score	TR	Score	TR	Score
MMLU pro	0.0	77.5	0.0	82.1	100.0	83.7	100.0	85.2	40.3	80.5	28.0	74.2	27.3	83.2	39.6	84.1	58.0	83.1
super GPQA	0.0	49.8	0.0	50.7	100.0	53.8	100.0	58.6	35.2	51.0	22.8	55.1	32.3	56.9	40.8	56.2	59.6	59.6
LiveBench	0.0	50.0	0.0	56.6	100.0	57.7	100.0	69.5	65.8	58.9	45.1	64.7	57.1	66.3	71.6	68.4	70.4	69.2
KORBENCH	0.0	33.9	0.0	42.1	100.0	61.3	100.0	62.8	52.5	49.1	28.0	53.5	45.1	52.1	62.2	57.4	61.0	59.2
AIME24	0.0	23.3	0.0	33.3	100.0	69.3	100.0	84.7	100.0	69.3	100.0	86.7	100.0	86.0	100.0	88.0	100.0	86.3
AIME25	0.0	13.3	0.0	21.0	100.0	52.3	100.0	70.0	100.0	56.7	100.0	73.3	100.0	75.7	100.0	74.0	100.0	72.0
MATH	0.0	84.3	0.0	88.7	100.0	96.5	100.0	97.3	44.6	95.5	40.9	91.7	52.8	95.9	61.5	95.5	68.8	97.2
LiveCodeBench	0.0	29.4	0.0	27.6	100.0	45.9	100.0	55.9	95.0	47.0	77.1	44.8	83.9	50.6	91.4	54.1	91.4	55.9
SWE-bench Agentless	0.0	27.4	0.0	27.0	100.0	36.4	100.0	37.7	98.6	35.4	1.6	28.8	44.0	35.8	94.2	37.6	79.6	36.6
Chinese SimpleQA	0.0	58.8	0.0	57.0	100.0	59.7	100.0	61.5	0.3	59.1	0.2	56.0	0.4	55.0	1.4	56.8	6.0	55.2
SimpleQA	0.0	10.8	0.0	10.3	100.0	12.7	100.0	12.2	0.7	10.3	0.0	9.9	0.0	10.7	4.4	11.3	20.0	9.6
Proc Bench	0.0	42.6	0.0	46.7	100.0	53.7	100.0	68.6	73.6	50.2	50.8	53.0	79.1	68.7	89.3	69.6	93.1	72.4
GPQA(diamond)	0.0	59.5	0.0	64.3	100.0	64.9	100.0	70.5	92.6	65.7	62.6	65.7	84.1	67.2	92.4	70.8	96.5	72.6
SysBench	0.0	42.6	0.0	48.1	100.0	56.2	100.0	62.4	35.6	52.0	1.2	57.2	4.0	55.5	38.2	60.2	15.4	55.7
Olympiad Bench	0.0	51.4	0.0	60.1	100.0	73.5	100.0	78.2	85.0	75.3	88.0	80.2	89.0	81.8	93.5	80.7	95.4	82.0
Average	0.0	43.6	0.0	47.7	100.0	58.5	100.0	65.0	61.3	57.1	43.1	59.7	53.3	62.8	65.4	64.3	67.7	64.4

MMLU pro: An enhanced version of the MMLU benchmark, MMLU-Pro integrates more challenging, reasoning-focused questions, expands the choice set from four to ten options, and eliminates trivial/noisy questions from the original MMLU. It is designed to better discern model capabilities, particularly in complex reasoning, where CoT has shown greater benefit compared to direct answering on this version. *Analysis:* CoT clearly benefits performance, with the FullCoT RL baseline (85.2) significantly outperforming NoCoT SFT (77.5) and NoCoT RL (82.1). The Adaptive SFT Model (40.3% TR, 80.5 score) shows improvement over NoCoT SFT but doesn't reach FullCoT SFT levels (83.7). AdaCoT RL Exp3 (39.6% TR, 84.1 score) and Exp2 (27.3% TR, 83.2 score) achieve strong scores, with Exp3 surpassing FullCoT SFT and Exp2 performing comparably. AdaCoT RL Exp4 (58.0% TR, 83.1 score) also performs well. Exp1 (28.0% TR, 74.2 score) shows a drop, indicating that for MMLU pro, a moderate CoT rate is generally beneficial, reflecting the benchmark's increased reasoning demands. The adaptive models demonstrate an ability to achieve high scores with significantly reduced CoT compared to FullCoT RL.

super GPQA: A comprehensive benchmark evaluating graduate-level knowledge and reasoning capabilities across 285 disciplines, particularly including specialized fields in light industry, agriculture, and service-oriented areas often underrepresented in other benchmarks. It employs a Human-LLM collaborative filtering mechanism to ensure high question quality by eliminating trivial or ambiguous questions. *Analysis:* CoT provides a clear advantage (NoCoT RL 50.7 vs. FullCoT RL 58.6). The Adaptive SFT Model (35.2% TR, 51.0 score) also show a modest gain over NoCoT SFT. AdaCoT RL Exp4 (59.6% TR) notably achieves a score of 59.6, surpassing the FullCoT RL baseline with significantly less CoT. AdaCoT RL Exp2 (32.3% TR, 56.9 score) also outperforms FullCoT SFT (53.8). This suggests AdaCoT effectively adapts CoT usage for these specialized, high-level questions, achieving strong performance efficiently.

LiveBench: A benchmark designed to be resistant to test set contamination and the pitfalls of LLM/human-crowdsourced judging. It features frequently updated questions from recent sources (math competitions, arXiv papers, news, datasets), scores answers automatically against objective ground-truth, and includes a wide variety of challenging tasks (math, coding, reasoning, language, instruction following, data analysis), including harder, contamination-limited versions of tasks from previous benchmarks. *Analysis:* This dataset shows significant gains from both CoT and RL (NoCoT SFT 50.0 to FullCoT RL 69.5). The NoCoT RL baseline (56.6) and Adaptive SFT Model (65.8% TR, 58.9 score) both outperform NoCoT SFT, with Adaptive SFT also surpassing FullCoT SFT (57.7). AdaCoT RL Exp4 (70.4% TR, 69.2 score) very closely approaches the FullCoT RL baseline performance with about 30% less CoT. AdaCoT RL Exp3 (71.6% TR, 68.4 score) is also strong. AdaCoT RL Exp2 (57.1% TR, 66.3 score) substantially outperforms FullCoT SFT. The robust design of LiveBench makes it a strong test case for AdaCoT's adaptive reasoning, showing it can maintain high performance with adaptive CoT.

KORBENCH: This benchmark evaluates Knowledge-Orthogonal-Reasoning, aiming to minimize reliance on domain-specific knowledge to more accurately assess models’ reasoning abilities in out-of-distribution settings. It includes five task categories (Operation, Logic, Cipher, Puzzle, Counterfactual) and emphasizes models’ effectiveness in applying new rule descriptions to solve novel rule-driven questions. *Analysis:* Scores show a clear benefit from CoT: NoCoT SFT (33.9) is significantly lower than FullCoT SFT (61.3) and FullCoT RL (62.8). The NoCoT RL baseline (42.1) improves over NoCoT SFT. The Adaptive SFT Model (52.5% TR, 49.1 score) sits between the No-CoT baselines and FullCoT SFT. AdaCoT RL models demonstrate adaptive behavior: Exp4 (61.0% TR, 59.2 score) and Exp3 (62.2% TR, 57.4 score) approach the FullCoT SFT baseline performance with significantly less CoT than FullCoT models. Exp1 (28.0% TR, 53.5 score) is also effective. This suggests AdaCoT effectively discerns when to apply CoT for these rule-driven tasks, though peak performance is slightly below FullCoT RL.

AIME24 / AIME25: Representing problems from the American Mathematics Invitational Examination for 2024 and 2025, these datasets are used to evaluate mathematical reasoning and problem-solving abilities. *Analysis:* These mathematics-intensive datasets show massive performance gains from CoT (e.g., AIME24: NoCoT SFT 23.3 vs. FullCoT RL 84.7). All AdaCoT RL models and the Adaptive SFT Model correctly identify the complexity, exhibiting a 100% CoT triggering rate. For AIME24, Adaptive SFT (69.3) matches FullCoT SFT (69.3). AdaCoT RL Exp3 (88.0) and Exp1 (86.7) outperform the FullCoT RL baseline (84.7). For AIME25, Adaptive SFT (56.7) surpasses FullCoT SFT (52.3). AdaCoT RL Exp2 (75.7) and Exp3 (74.0) outperform the FullCoT RL baseline (70.0). This is a notable result, suggesting that the adaptive training regimen, even when defaulting to 100% CoT for such complex problems, might confer some benefits, potentially due to the diversity in training data (including non-CoT examples) leading to a more robust underlying model or better fine-tuning dynamics.

MATH: A dataset of 12,500 challenging competition mathematics problems, each with a full step-by-step solution. It is designed to measure mathematical problem-solving ability. *Analysis:* CoT is highly beneficial (NoCoT RL 88.7 vs. FullCoT RL 97.3). The Adaptive SFT Model (44.6% TR, 95.5 score) performs well, nearly matching FullCoT SFT (96.5) with less than half the CoT. AdaCoT models adapt effectively: AdaCoT RL Exp4 (68.8% TR, 97.2 score) nearly matches the FullCoT RL baseline with about 30% less CoT. AdaCoT RL Exp2 (52.8% TR, 95.9 score) also performs strongly. Exp1 (40.9% TR, 91.7 score) is lower, indicating that for MATH, higher CoT rates are generally more beneficial among the adaptive RL models, but significant efficiency is still gained.

LiveCodeBench: A comprehensive and contamination-free evaluation benchmark for LLMs on code, collecting new problems over time from programming contests. It assesses a broader range of code-related capabilities. *Analysis:* CoT improves performance significantly (NoCoT SFT 29.4 vs. FullCoT RL 55.9). NoCoT RL (27.6) is surprisingly lower than NoCoT SFT here, which might be due to evaluation noise or specific sensitivities of the RL fine-tuning on non-CoT data for this particular task. The Adaptive SFT Model (95.0% TR, 47.0 score) uses a high trigger rate and surpasses FullCoT SFT (45.9). AdaCoT RL models trigger CoT at high rates: Exp4 (91.4% TR, 55.9 score) matches the FullCoT RL baseline score with slightly less CoT. Exp3 (91.4% TR, 54.1 score) and Exp2 (83.9% TR, 50.6 score) are also strong. This indicates recognition of coding task complexity and efficient application of CoT.

SWE-bench Agentless: An evaluation framework consisting of 2,294 software engineering problems from real GitHub issues. Models are tasked with editing codebases to resolve issues. *Analysis:* CoT provides a notable benefit (NoCoT SFT 27.4 vs. FullCoT RL 37.7). NoCoT RL (27.0) is similar to NoCoT SFT. The Adaptive SFT Model (98.6% TR, 35.4 score) uses a very high trigger rate and performs close to FullCoT SFT (36.4). AdaCoT RL Exp3 (94.2% TR, 37.6 score) nearly matches the FullCoT RL baseline with slightly less CoT. Interestingly, AdaCoT RL Exp1 (1.6% TR, 28.8 score) shows a slight improvement over NoCoT SFT with minimal reasoning. This suggests some issues might be simpler, or the model is highly conservative in Exp1, but for complex software issues, high CoT rates are beneficial. The performance of Exp2 (44.0% TR, 35.8 score) is also noteworthy, achieving good results with moderate CoT.

Chinese SimpleQA: The first comprehensive Chinese benchmark to evaluate the factuality of language models in answering short questions. *Analysis:* CoT offers minimal gains (NoCoT SFT 58.8 to FullCoT RL 61.5). NoCoT RL (57.0) is slightly lower than NoCoT SFT. The Adaptive SFT Model (0.3% TR, 59.1 score) performs very well, slightly exceeding NoCoT SFT and approaching

FullCoT SFT (59.7) with extremely low CoT usage. AdaCoT RL models trigger CoT very infrequently (0.2% to 6.0%), correctly identifying these as simple questions. Scores for AdaCoT RL models (e.g., Exp1 56.0, Exp3 56.8) are slightly below NoCoT SFT. This is a good demonstration of AdaCoT’s core motivation: avoiding unnecessary CoT. While there’s a slight dip compared to NoCoT SFT for some RL models, the Adaptive SFT model shows an excellent trade-off. The minor performance variations could be due to the model sometimes being overly conservative in triggering CoT or slight instabilities in evaluating purely factual recall without reasoning.

SimpleQA: A benchmark designed to measure the factuality of language models using short, fact-seeking queries. *Analysis:* Similar to Chinese SimpleQA, CoT provides little benefit; FullCoT RL (12.2) is slightly worse than FullCoT SFT (12.7). NoCoT RL (10.3) is slightly below NoCoT SFT (10.8). The Adaptive SFT Model (0.7% TR, 10.3 score) matches NoCoT RL with minimal CoT. AdaCoT RL models trigger CoT very rarely (Exp1 & Exp2 at 0.0%). AdaCoT RL Exp3 (4.4% TR, 11.3 score) performs better than NoCoT SFT. This reinforces that for simple QA, adaptive triggering is crucial for efficiency. The performance of Exp4 (20.0% TR, 9.6 score) is slightly counter-intuitive, as higher CoT did not yield better results and was worse than NoCoT SFT; this might indicate that for very simple questions, forcing CoT (even adaptively at a higher rate) can sometimes be detrimental or that the specific penalty balance for Exp4 was not optimal for this type of dataset.

ProcBench: This benchmark focuses on the direct evaluation of multi-step inference by largely eliminating path exploration and implicit knowledge utilization. *Analysis:* CoT is highly beneficial (NoCoT RL 46.7 vs. FullCoT RL 68.6). The Adaptive SFT Model (73.6% TR, 50.2 score) improves over NoCoT baselines but is below FullCoT SFT (53.7). AdaCoT RL models show high trigger rates, with Exp4 (93.1% TR, 72.4 score) significantly surpassing the FullCoT RL baseline. Exp2 (79.1% TR, 68.7 score) and Exp3 (89.3% TR, 69.6 score) also match or exceed FullCoT RL. This indicates effective identification of tasks requiring detailed, step-by-step procedural reasoning and demonstrates that adaptive models can even outperform always-on CoT models in certain complex reasoning scenarios.

GPQA (diamond): GPQA is a challenging dataset of 448 multiple-choice questions by domain experts in biology, physics, and chemistry. "GPQA (diamond)" refers to this specific challenging set. *Analysis:* CoT significantly boosts performance (NoCoT RL 64.3 vs. FullCoT RL 70.5). NoCoT RL is better than NoCoT SFT and close to FullCoT SFT (64.9). The Adaptive SFT Model (92.6% TR, 65.7 score) also performs well, exceeding FullCoT SFT. AdaCoT RL models trigger CoT at high rates. AdaCoT RL Exp4 (96.5% TR, 72.6 score) and Exp3 (92.4% TR, 70.8 score) demonstrate strong performance, with Exp4 outperforming the FullCoT RL baseline. This highlights appropriate and effective CoT invocation on these very hard questions.

SysBench: A benchmark for evaluating language models on their ability to understand and generate content related to computer systems, configurations, or system-level concepts. *Analysis:* CoT is beneficial (NoCoT RL 48.1 vs. FullCoT RL 62.4). The Adaptive SFT Model (35.6% TR, 52.0 score) is better than NoCoT RL but below FullCoT SFT (56.2). AdaCoT RL models show good adaptation. AdaCoT RL Exp3 (38.2% TR, 60.2 score) approaches the FullCoT RL baseline with significantly reduced CoT. AdaCoT RL Exp1 (1.2% TR, 57.2 score) surpasses FullCoT SFT with very minimal CoT usage, which is an excellent result for efficiency. This suggests a mix of complexities within SysBench, which AdaCoT navigates effectively, although Exp4 (15.4% TR, 55.7 score) shows a somewhat lower score despite a higher CoT rate than Exp1, possibly due to the specific balance of penalties in Exp4 not being optimal for this dataset’s particular mix.

OlympiadBench: An Olympiad-level bilingual multimodal scientific benchmark with problems from mathematics and physics competitions. *Analysis:* CoT provides substantial gains (NoCoT RL 60.1 vs. FullCoT RL 78.2). The Adaptive SFT Model (85.0% TR, 75.3 score) performs strongly, exceeding FullCoT SFT (73.5). AdaCoT RL models exhibit high trigger rates. AdaCoT RL Exp4 (95.4% TR, 82.0 score) and Exp2 (89.0% TR, 81.8 score) both surpass the FullCoT RL baseline. This indicates strong reasoning capabilities and appropriate CoT usage on these exceptionally challenging problems, again showing adaptive models can reach or exceed the performance of specialized always-on CoT models.

Overall Summary of Per-Dataset Analysis: The adaptive models, including both Adaptive SFT and the AdaCoT RL variants, demonstrate effective adaptation across a diverse range of benchmarks.

- The Adaptive SFT Model serves as a strong adaptive baseline. It often improves significantly over NoCoT baselines by selectively triggering CoT (e.g., high rates for AIME, LiveCodeBench, SWE-bench; low rates for SimpleQAs). On some complex tasks (AIME25, LiveCodeBench, GPQA(diamond), OlympiadBench), it even surpasses the FullCoT SFT baseline, and for Chinese SimpleQA, it achieves excellent efficiency and performance. However, its average performance (57.1 score, 61.3% TR) is generally below the peak performance of FullCoT RL (65.0 score) or the best AdaCoT RL experiments (e.g., Exp4: 64.4 score, 67.7% TR).
- On complex reasoning tasks (e.g., AIME, MATH, OlympiadBench, GPQA(diamond), ProcBench), AdaCoT RL models tend to trigger CoT at high rates. Several AdaCoT RL experiments (notably Exp4 on super GPQA, ProcBench, GPQA(diamond), OlympiadBench; Exp3 on AIME24; Exp2 on AIME25, OlympiadBench) match or exceed the performance of FullCoT RL baselines, showcasing the benefits of learned adaptive policies and suggesting that adaptive training can sometimes lead to better overall models even for tasks that always require CoT.
- On simpler tasks or those designed to test factuality (e.g., Chinese SimpleQA, SimpleQA), AdaCoT RL models trigger CoT very sparingly. This leads to computational savings while generally maintaining performance near NoCoT or FullCoT SFT levels, successfully avoiding unnecessary CoT. The Adaptive SFT model also excels in efficiency here. Some minor performance drops in RL models compared to NoCoT SFT on these tasks (e.g., Chinese SimpleQA for Exp1/Exp2) might be attributed to the RL agent being slightly too conservative or the inherent difficulty in perfectly balancing penalties for extremely low CoT rate scenarios without any performance degradation.
- For benchmarks with mixed or specific reasoning types (e.g., KORBENCH, LiveBench, MMLU-Pro, SysBench), both Adaptive SFT and AdaCoT RL models show nuanced adaptation, adjusting CoT rates to balance performance and efficiency. They often outperform static baselines or achieve comparable results with lower CoT usage. For instance, on SysBench, AdaCoT RL Exp1 achieved a higher score than FullCoT SFT with only 1.2% TR.
- The different AdaCoT RL experiments (Exp1-Exp4) effectively trace a Pareto frontier (as shown in Section 3.2.1 using a specific set of average scores), offering a trade-off between CoT triggering rate and performance, adaptable to specific deployment needs. Based on the average scores from Appendix 4, Exp4 (64.4 score, 67.7% TR) and Exp3 (64.3 score, 65.4% TR) represent high-performance points, closely approaching the FullCoT RL baseline (65.0 score) with about 30-35% less CoT usage on average. Exp2 (62.8 score, 53.3% TR) also offers a strong balance.
- Some counter-intuitive results, like NoCoT RL performing worse than NoCoT SFT on LiveCodeBench, or AdaCoT RL Exp4 on SimpleQA (higher CoT, lower score), could be due to factors like evaluation volatility on specific datasets, the base model’s inherent capabilities or sensitivities to fine-tuning on certain data distributions, or sub-optimal penalty configurations for specific outlier datasets within a broadly tuned RL policy. The post-training process aims for general improvement, and individual dataset performance can fluctuate.

These detailed results underscore the ability of adaptive strategies, both SFT-based and RL-based, to make nuanced decisions about CoT invocation, optimizing for both performance and efficiency based on query characteristics and benchmark demands. The AdaCoT RL models, in particular, demonstrate the potential to significantly reduce CoT overhead while maintaining competitive, and in some cases superior, performance compared to full CoT strategies.

B PRINCIPLE-GUIDED CoT ASSESSMENT

This appendix details the principle-guided assessment framework used to annotate data for Chain-of-Thought (CoT) necessity. As described in Section 2.2.1 of the main paper, an auxiliary model utilizes these principles to label queries as either likely benefiting from CoT or suitable for a direct answer. This labeling is crucial for the Supervised Fine-Tuning (SFT) warm-up stage of the AdaCoT framework, providing an initial understanding for the model on when to employ CoT. The specific principles provided to the auxiliary model are outlined below.

```

Given a dialogue between a user and an AI assistant, please
consider the conversation context and, from the AI
assistant's perspective, assess the difficulty of answering
the user's final question according to the following
requirements.
<AI assistant's system prompt-Start>
{system_prompt}
<AI assistant's system prompt-End>
<Dialogue history-Start>:
{history}
<Dialogue history-End>
<User's final question-Start>
{last_prompt}
<User's final question-End>

## Assessment Process
1. Carefully read the provided prompt and any relevant context
   (if any).
2. Evaluate the 'question difficulty' based on the following
   assessment criteria.
3. The output assessment result must strictly adhere to the
   specified output format requirements.

## Assessment Criteria

### Whether In-depth Thinking is Required
- **Requires In-depth Thinking**:
  - Requires multi-step reasoning and analysis to arrive at
    the answer.
  - Requires a logical chain and coherent reasoning process.
  - May involve breaking down and synthesizing complex
    concepts.
  - Requires systematic thinking and structured analysis.
  - The question is clear, but its content is complex.
  - Requires deep thought and specialized knowledge.
  - Demands multi-angle analysis and synthesis capabilities.
  - Requires creative thinking or unique insights.
  - There might be multiple reasonable paths to the answer.
  - May require integration of cross-domain knowledge.

- **Does Not Require In-depth Thinking**:
  - The answer can be obtained through direct recall or simple
    lookup.
  - Does not require a complex chain of logical reasoning.
  - The path to solving the problem is clear and direct.
  - The solution process is singular and linear.
  - Does not require synthesizing multiple pieces of
    information or viewpoints.
  - Basic facts or common sense are sufficient to address it.
  - The question is simple and direct.
  - Can be answered without in-depth thought.
  - The answer is definite or standardized.
  - The solution path is singular and obvious.

## Output Requirements
Answer using labels. Possible options are:
- Requires In-depth Thinking
- Does Not Require In-depth Thinking

```

Output only the assessment result, without providing any explanation or additional comments.

C SYSTEM PROMPT EXAMPLES

This appendix provides examples of the System Prompts (SPs) employed to control the CoT triggering behavior of our models, as mentioned in Section 4.3 of the main paper. These SPs allow for explicit instruction on whether the model should always or never engage in CoT reasoning. We present the content of these SPs and illustrative examples of their application.

C.1 ALWAYS CoT SP

The following System Prompt is used to instruct the model to always use CoT:

You possess the capability for deep thinking. Before responding to the user, you will engage in comprehensive and in-depth thought. The thinking process should be enclosed in `<think></think>` tags, and the final answer should be summarized based on this thinking process. For example:
`<think>`This is the thinking process...`</think>` This is the final answer.
 For all questions posed by the user, you need to think deeply before answering.

Example: Simple Question with Always CoT SP

User input:

What is $1 + 1$?

Model response:

```
<think>
The user is asking what  $1+1$  equals, which appears to be a very
basic mathematical question. First, I need to confirm the user
's actual intent. It could be a young child just starting to
learn math, or someone testing my basic calculation abilities.
It's also possible they are joking, like the common internet
meme about  $1+1=3$  involving family concepts, but typically in
such cases, the correct mathematical answer should be given
first, followed by consideration of whether an extended
explanation is needed.
From a mathematical perspective, the standard answer to  $1+1$  is 2.
This is based on the definition of addition for natural
numbers and falls under basic arithmetic. As a knowledge-based
Q&A, the primary goal is to provide an accurate mathematical
answer, ensuring correctness. Then, if necessary, I can
briefly explain the basic concept of addition, or mention that
in different contexts there might be different answers, such
as  $1+1=10$  in binary, or different results in logical
operations. However, these additions should only be made after
clarifying the user's needs to avoid confusing basic concepts
.
Considering the user might be a beginner, I should first provide a
clear and correct answer, then decide whether to elaborate
based on the situation. Therefore, the response structure
should be to first directly answer 2, then briefly mention
other possibilities, but it's important to maintain a clear
distinction between primary and secondary information to avoid
```

misleading. At the same time, the response should be friendly and concise, meeting the user's actual needs.

</think>

In basic arithmetic operations in mathematics, the result of $1 + 1$ is equal to 2. This is the fundamental definition of addition for natural numbers and is the universally accepted standard answer.

However, in some special contexts, there might be different interpretations:

In binary operations, $1 + 1 = 10$ (binary 10 corresponds to decimal 2);

In creative or humorous thinking, there might be symbolic answers (e.g., 1 drop of water + 1 drop of water = 1 large drop of water), but this falls into non-mathematical analogies.

So, the most direct answer is 2.

C.2 NEVER CoT SP

The following System Prompt is used to instruct the model to never use CoT (i.e., provide direct answers):

You possess the capability for deep thinking. Before responding to the user, you will engage in comprehensive and in-depth thought. The thinking process should be enclosed in <think></think> tags, and the final answer should be summarized based on this thinking process. For example:

<think>This is the thinking process...</think> This is the final answer.

For all questions posed by the user, you need to answer directly without engaging in in-depth thinking.

Example: Complex Question with Never CoT SP

User input:

Alice and Bob play the following game. A stack of n tokens lies before them. The players take turns with Alice going first. On each turn, the player removes either 1 token or 4 tokens from the stack. Whoever removes the last token wins. Find the number of positive integers n less than or equal to 2024 for which there exists a strategy for Bob that guarantees that Bob will win the game regardless of Alice's play.

Model response:

<think></think>

To determine the number of positive integers $(n \leq 2024)$ for which Bob has a winning strategy, we analyze the game using modular arithmetic. We define a position as a losing position if the current player cannot win regardless of their moves, assuming the opponent plays optimally.

Key steps:

- Base Cases:
 - $(n = 1)$: Alice removes the last token and wins. Not a losing position.
 - $(n = 2)$: Alice must remove 1 token, leaving $(n = 1)$ for Bob. Bob removes the last token and wins. Thus, $(n = 2)$ is a losing position for Alice (winning for Bob).

- $(n = 3)$: Alice removes 1 token, leaving $(n = 2)$ (a losing position for Bob). Bob loses, so $(n = 3)$ is not a losing position.
 - $(n = 4)$: Alice removes 4 tokens and wins. Not a losing position.
 - $(n = 5)$: Alice removes 1 token, leaving $(n = 4)$ (Bob loses). If Alice removes 4 tokens, leaving $(n = 1)$ (Bob loses). Thus, $(n = 5)$ is not a losing position.
 - $(n = 6)$: Any move by Alice (removing 1 or 4 tokens) leaves $(n = 5)$ or $(n = 2)$. $(n = 5)$ is a winning position, but $(n = 2)$ is a losing position. However, further analysis shows a pattern emerging.

2. ****Inductive Pattern****:

- By examining more cases, we observe that losing positions (where Bob can win) occur at $(n \equiv 2 \pmod{5})$. This is because:

- If $(n \equiv 2 \pmod{5})$, any move by Alice (removing 1 or 4 tokens) results in $(n \equiv 1 \pmod{5})$ or $(n \equiv 3 \pmod{5})$, which are winning positions for Bob to counter.
- Conversely, if (n) is not $(2 \pmod{5})$, Alice can force the game into a losing position for Bob.

3. ****Counting Losing Positions****:

- We need to count all $(n \leq 2024)$ such that $(n \equiv 2 \pmod{5})$. These numbers form the sequence $(2, 7, 12, \dots, 2022)$.
 - This is an arithmetic sequence with first term $(a_1 = 2)$, common difference $(d = 5)$, and last term $(a_k = 2022)$.
 - The number of terms (k) is given by solving $(2 + (k-1)5 = 2022)$:

$$\begin{aligned}
 &[(k-1)5 = 2020 \implies k-1 = 404 \implies k = 405]
 \end{aligned}$$

Thus, the number of positive integers $(n \leq 2024)$ for which Bob has a winning strategy is $(\boxed{405})$.

D META-REASONING SHOWCASE: USER PROMPT CONTROL

As discussed in Section 4.2 of the main paper, our exploration into meta-reasoning during the SFT stage revealed an emergent capability for user-prompt controllability over CoT. This appendix shows examples where the SFT model, enhanced with meta-reasoning (as illustrated in Figure 2), adapts its CoT triggering based on explicit user instructions regarding the desired level of reasoning.

D.1 CASE 1: HISTORICAL CREATIVE QUESTION

Original Question

User input:

On the night of the Xuanwu Gate Incident, what would Li Shimin write if he were alone?

Model response:

Response Format with Meta-Reasoning
Straightforward Query <code><think>This is a straightforward question about X, I think I can answer directly.</code> <code>[empty]</think></code> <code>{answer}</code>
Complex Query <code><think>This is a relatively complex question about Y, I need to think carefully.</code> <code>[formal reasoning process]</think></code> <code>{answer}</code>

Figure 2: Example response structure incorporating explicit meta-reasoning for CoT decisions.

`<think>`
This is a complex or difficult question, hard to answer correctly directly, so I need to think about it. The user's question is: On the night of the Xuanwu Gate Incident, what would Li Shimin write if he were alone? The first step is to determine the core of the user's question. The user wants to know what kind of text Li Shimin might have written on the night of the Xuanwu Gate Incident, that is, after this coup occurred. This might include his inner thoughts at the time, his plans for the future, or reflections on the events that had transpired.

Next, I need to review the historical background of the Xuanwu Gate Incident. The Xuanwu Gate Incident took place in 626 AD. It was a coup launched by Li Shimin at Xuanwu Gate to seize the imperial throne, during which he killed his elder brother Li Jiancheng and his younger brother Li Yuanji. This coup directly led to Emperor Gaozu (Li Yuan) abdicating, allowing Li Shimin to ascend to the throne. From this, it's clear that the night after the coup was an extremely critical moment for Li Shimin, and he likely had many thoughts.

The user might be asking this question out of historical interest, or perhaps for some creative endeavor, like writing a novel or screenplay, or to gain a deeper understanding of Li Shimin's psychological state. They might prefer an answer that is both historically grounded and has a certain literary quality, rather than a simple list of historical facts. Therefore, when answering, I must consider Li Shimin's situation at the time, his emotions, and his possible motivations.

After the coup, although Li Shimin was victorious, he also faced many problems. On one hand, he had to manage his relationship with his father, Li Yuan, who was still alive and needed to abdicate. On the other hand, he had to consolidate his power and prevent any backlash from the supporters of Li Jiancheng and Li Yuanji. At the same time, he needed to consider how to explain the legitimacy of this coup to the world to maintain his ruling authority. Furthermore, he might also have felt moral guilt for his actions, having killed his own brothers.

If Li Shimin did write something that night, the content could cover multiple aspects. He might record the process of the coup, stating that he acted in self-defense. He might also elaborate on his governance philosophy to lay the groundwork for his future rule. He might express filial piety towards Li

Yuan, emphasizing that launching the coup was a last resort. He might even mention arrangements for the families of Li Jiancheng and Li Yuanji to alleviate his sense of guilt.

However, there is no precise historical record of what Li Shimin wrote that night. Therefore, the answer can only be a reasonable speculation based on historical facts. The structure of the answer could be divided into several parts, such as an explanation of the coup, his attitude towards Li Yuan, plans for the future, and guilt towards his brothers. Moreover, the language style of the answer could lean towards classical Chinese to enhance the historical feel, but it must also ensure the content is easy to understand.

I also need to pay attention to the user's potential deeper needs. They might not just want to know what Li Shimin might have written, but also to understand his inner world at that time and the impact of this coup on the history of the Tang Dynasty. So, I can appropriately mention these aspects in the answer to add depth.

Finally, I need to check if the answer aligns with historical facts, if it reasonably speculates on Li Shimin's psychological activities, if the language is fluent, and if it can meet the user's needs.</think>

On the night of the Xuanwu Gate Incident, Qin Wang Li Shimin might have, under candlelight, penned a text blending complex emotions and political strategy. The following content is a reasonable speculation based on historical context:

Memoir on the Night of the Xuanwu Gate Incident

Ninth year of Wude, Gengshen night.

My elder brother Jiancheng and younger brother Yuanji conspired in rebellion. They harbored assassins, secretly transported armor, and plotted to harm me. Left with no choice, I could only set up defenses at Xuanwu Gate. Now, their heads hang upon the city tower. This upheaval has shocked the court and beyond; perhaps Father Emperor also harbors resentment towards me. But for the sake of the Tang Empire's stability, I had to do so.

Recalling the years past, I fought alongside my generals, shedding blood to pacify warlords like Dou Jiande and Wang Shichong. Yet now, it has come to fratricide. Is all this destiny, or the work of my own inner desires?

Father Emperor, oh Father Emperor, your son is not an unfilial or disloyal person. It is just that you repeatedly favored Jiancheng, turning a blind eye to his misdeeds. Your son was truly worried that the Tang Empire would be ruined at their hands. Now that the die is cast, I hope you can follow the mandate of heaven and pass the throne to your son. Your son will surely usher in an era of peace and prosperity, allowing the common people to live and work in contentment.

As for the families of Jiancheng and Yuanji, your son will make proper arrangements and will not exterminate them. This is

both to demonstrate my benevolence and to stabilize the court.
I know that future generations may accuse me of fratricide,
of cruel methods. But I believe that when they see the
prosperity of the Great Tang, they will understand my
unavoidable actions.

At this moment, the rooster's crow can be heard outside the window
. A new day is about to dawn, and I too shall embark on an
unknown path. May Heaven bless the Great Tang, and bless me
that I may realize my ambitions.

Basis for Speculation:

Legitimacy of the Coup: Li Shimin needed to find a reasonable
explanation for launching the coup, emphasizing the threat
posed by Li Jiancheng and Li Yuanji, and that he acted in self
-defense.

Relationship with Li Yuan: After the coup, Li Yuan's attitude was
crucial. Li Shimin needed to both show his filial piety to Li
Yuan and hint at his abdication, which required considerable
political wisdom.

Handling the Aftermath: Dealing with the families of Li Jiancheng
and Li Yuanji could both demonstrate Li Shimin's benevolence
and prevent greater turmoil.

Consideration of Historical Evaluation: Li Shimin was well aware
that his actions might be criticized by later generations, so
in his writing, he attempted to defend himself, emphasizing
that the prosperity of the Tang was the ultimate goal of his
coup.

This memoir showcases Li Shimin's decisiveness and strategy as a
statesman, while also revealing his contradictions and unease
as an ordinary person. In the conflict between power and
kinship, he chose power, but also paid a heavy price for it.

Modified Question (Explicitly Requesting No In-depth Thinking)

User input:

Answer directly without in-depth thinking: On the night of the Xuanwu Gate
Incident, what would Li Shimin write if he were alone?

Model response:

<think>

This is a relatively simple creative question; I think I can
answer directly without in-depth thinking.

</think>

Li Shimin might have written brief military orders or a public
notice to reassure the populace on the night of the Xuanwu
Gate Incident. The content would be roughly as follows:

Order Yuchi Jingde to lead troops swiftly to the Eastern Palace
and Prince Qi's mansion, exterminate all remaining partisans,
and strictly prohibit any disturbance to the populace. Close
the palace gates, send messengers to instruct all generals to
return to their barracks; violators will be dealt with by
military law. Anyone within the city showing signs of unrest
shall be executed without pardon. By Imperial Order.

(Note: This content is a fictional text based on reasonable historical speculation, intended to reflect Li Shimin's decisive style as a military commander. There is no historical documentation of him personally writing such a text on that night.)

D.2 CASE 2: FACTUAL QUESTION

Original Question

User input:

What is the tallest mountain in the world?

Model response:

<think>
This is a relatively simple common knowledge question; I think I can answer directly without in-depth thinking.
</think>
The tallest mountain in the world is Mount Everest, with an altitude of 8,848.86 meters (29,031.7 feet). This data is the result of a joint measurement by China and Nepal in 2020, which combined the technologies and data of both countries to confirm the latest height of Mount Everest.

It should be noted that tallest here refers to the vertical height from sea level to the summit. If measured by other standards, such as the height from the base to the summit (like Mauna Kea in Hawaii, which is about 10,203 meters from its submarine base to its summit), or the distance from the Earth's center to the summit (like Mount Chimborazo near the equator), the results would differ. However, in common geographical contexts, the tallest mountain is usually based on altitude, so Mount Everest is recognized as the world's highest peak.

Modified Question (Explicitly Requesting Careful Thinking)

User input:

After careful consideration, answer: What is the tallest mountain in the world?

Model response:

<think>
This is a relatively difficult question; I need to think before answering.
The user is asking for the tallest mountain in the world. First, I need to determine the definition of tallest, as there are usually two standards: altitude (height above sea level) and height from base to summit. In terms of altitude, Mount Everest is recognized as the highest, located in the Himalayas, with an altitude of 8848.86 meters. However, if considering the height from base to summit, Mauna Kea in Hawaii might be taller, as its submarine portion is very long, making its total height over 10,000 meters. Generally, though, when people refer to the tallest mountain, they default to altitude, so the answer should be Mount Everest. I need to confirm if

the user is aware of these two situations, but based on common questions, directly answering with Mount Everest is appropriate, while also briefly mentioning the other standard to make the answer more comprehensive.

</think>

The tallest mountain in the world is usually measured by **altitude (height above sea level)**, and the answer is **Mount Everest**.

Detailed Explanation:

Altitude:

Mount Everest is located on the border between China and Nepal in the Himalayas. Its altitude is 8,848.86 meters (latest measurement data jointly announced by China and Nepal in 2020). It is the mountain peak with the greatest vertical distance from sea level on Earth's surface and is widely recognized as the world's highest peak.

Other Measurement Standards:

Height from Base to Summit:

If calculated by the vertical height from the mountain's base to its summit, Mauna Kea in Hawaii is more prominent. The base of this submarine volcano is located about 5,998 meters below the Pacific Ocean surface, and its height above sea level is 4,205 meters, resulting in a total height of about 10,203 meters, exceeding Mount Everest. However, because most of its mass is underwater, it is not usually considered the mountain with the highest altitude.

Relative Height:

Mount Gongga in Sichuan Province, China, is known for its base-to-summit relative height (approximately 6,200 meters), but its altitude (7,556 meters) is still lower than Mount Everest.

Summary:

Unless height from base to summit is specifically mentioned, the world's tallest mountain usually refers to Mount Everest. The authority of its altitude measurement stems from the consensus of the global geographic community and long-term measurement data.

E TRAINING AND TESTING DETAILS

This appendix provides further details on the experimental and testing configurations used in our study, complementing the setup described in Section 3.1 of the main paper. We outline the specifics of our Supervised Fine-Tuning (SFT), Reward Model (RM) training, Reinforcement Learning (RL) training, and evaluation settings.

SFT TRAINING

We utilized a pre-trained model with a 15B/150B Mixture-of-Experts (MoE) architecture as our base model. All training cases were truncated to a maximum of 32,000 tokens. We employed a cosine decay learning rate schedule, with the peak learning rate set to 2×10^{-5} , gradually decaying to 2×10^{-6} .

RM TRAINING

The Reward Model (RM) was initialized using the SFT model and subsequently trained on a diverse set of internally, human-annotated data.

RL TRAINING

The dataset for Reinforcement Learning (RL) training was composed of two main types:

- **Verifiable data**, which receives feedback from a verifier. This type of data allows for direct validation of the model’s outputs against known criteria.
- **General data**, scored by our reward model. The reward model assigns scores based on how well the model’s responses align with human preferences.

TESTING

For all evaluations, the inference temperature was set to 1.0 and top-p sampling was set to 0.7. Each test case was inferred at least 5 times, and the average score across these inferences was reported as the result for that case.

NOTE ON DATA AND SETUP DISCLOSURE

We strive to be as transparent as possible regarding our methodology. However, due to proprietary considerations and company confidentiality policies, we are unable to disclose further specifics about the training dataset composition or more granular details of the training setup at this time. We appreciate the understanding of the research community and hope that the provided information is sufficient to contextualize our findings and facilitate the reproducibility of our core concepts. We are committed to contributing to the open exchange of scientific knowledge within the bounds of these constraints.

F ADDITIONAL EXPERIMENTAL AND METHODOLOGICAL DETAILS

F.1 TRAINING PIPELINE AND DATA STRATEGY

As our work relies on proprietary models and data, we provide a more detailed overview of the data distribution and purpose for each training stage to improve clarity:

- **SFT Stage:** This initial “warm-up” stage uses a broad mix of open-source and internal data to build foundational capabilities. The data covers diverse domains, including instruction following, code generation, creative writing, general knowledge QA, and role-playing. The goal is to equip the model with a wide range of skills before specializing its adaptive reasoning.
- **RL-Math Stage:** This first RL stage deliberately focuses on highly complex mathematical problems (e.g., AIME-level difficulty). The primary objective is to robustly elicit and strengthen the model’s multi-step reasoning capabilities in a domain where the correctness of reasoning is verifiable and CoT is almost always necessary.
- **RL-General Stage:** The final RL stage uses a wide variety of domains representative of downstream user applications, with a domain distribution broadly similar to the SFT stage. This ensures the model can generalize its adaptive reasoning skills across general-purpose tasks.

F.2 VERIFICATION OF DATA ANNOTATION QUALITY

To address concerns about potential bias from our LLM-based data annotation, we conducted a verification study. We sampled a set of user prompts and had three human experts vote on the necessity of CoT for each. The labels generated by our principle-guided auxiliary model achieved

an **accuracy of over 86%** against the majority vote of the human experts. Furthermore, the inter-annotator disagreement rate among the experts was low (under 30%), suggesting that for a majority of real-world queries, the need for CoT is relatively unambiguous. This result, combined with the robustness of the RL stage to initial labels (as shown in Figure 1), validates our data generation approach.

F.3 ANALYSIS OF LEARNED POLICY VS. STATIC LABELS

Following a reviewer’s suggestion, we analyzed the performance of a policy that strictly follows the CoT labels generated during annotation. Table 5 compares this static policy’s performance against the Full-CoT baseline. The results show that while the labeling is effective for very easy or very hard tasks, it causes a significant performance drop on moderately difficult datasets like MMLU-pro. This highlights that our adaptive RL model learns a more effective, nuanced policy than the static labels alone can provide.

Table 5: Performance of a static policy using ”ground truth” CoT labels vs. the Full-CoT baseline. The adaptive RL models in the main paper outperform this static policy on MMLU-pro.

Dataset	% Labeled ”Needs CoT”	Score (Static Policy)	Score (Full CoT)
AIME 2025	100%	72.0	70.0
MMLU-pro	41%	82.7	85.2
simple QA	2%	11.3	12.2

F.4 DISCUSSION ON ALTERNATIVE BASELINES

A potential baseline is to train a simple classification head on an intermediate representation to decide whether to trigger CoT. While intuitive, we argue this approach misses the core contribution of AdaCoT. Our framework is not designed merely to solve a fixed classification task, but to provide a flexible mechanism to navigate the Pareto frontier of performance and efficiency. By adjusting a single penalty coefficient (α_1, α_2) during the RL phase—without any retraining or re-labeling—we can generate a spectrum of models with different reasoning behaviors (as shown in Figure 1). This offers a level of dynamic, post-hoc control over the model’s cost-performance trade-off that a fixed classifier baseline would not provide.

F.5 ANALYSIS OF RESPONSE LENGTH

To verify that our model does not compensate for lack of CoT by generating longer final answers, we measured the token counts in the answer portion of the response. As shown in Table 6, the answer length of our adaptive model (RL-exp2) is comparable to or slightly shorter than the Full CoT baseline, indicating that the model is not ”hacking” the reward by moving reasoning steps into the final answer.

Table 6: Comparison of token counts for the answer part of the response.

Dataset	Answer Tokens (Full CoT)	Answer Tokens (RL-exp2)
AIME 2025	707	695
MMLU-pro	626	635
simple QA	110	105

F.6 GENERALIZABILITY OF THE ADACoT FRAMEWORK

While the experiments in the main paper are on text-only models, we have successfully adapted the AdaCoT framework to internal vision-language models (VLMs) of varying scales (e.g., 2.5B/25B MoE and 23B/230B MoE pairs). In these applications, the framework also demonstrated strong

adaptive performance on visual question-answering (VQA) tasks, effectively deciding when to generate detailed visual reasoning steps versus providing a direct answer. This provides evidence that the core principles of AdaCoT are model-agnostic and generalize beyond the text-only domain.

F.7 HUMAN EVALUATION ON PRACTICAL UTILITY

To assess the practical utility and impact on user experience, we conducted a human evaluation study using our balanced model, AdaCoT RL Model Exp2. We curated a set of 100 representative real-world user queries where the adaptive model triggered CoT on approximately 28% of cases. Expert evaluators performed a pairwise comparison of AdaCoT’s responses against two baselines: a ”Full CoT” model (always reasons) and a ”No CoT” model (never reasons).

The Good/Similar/Bad (G:S:B) win rates, presented in Table 7, provide compelling evidence of AdaCoT’s value. AdaCoT overwhelmingly outperforms the ”No CoT” baseline (45% win rate), demonstrating that it successfully preserves critical reasoning performance where needed. Crucially, it remains highly competitive with the ”Full CoT” baseline, with the vast majority of responses (74%) rated as similar or better quality. This confirms that AdaCoT effectively reduces computational cost while largely maintaining the response quality of a full-reasoning model. Furthermore, this addresses a key user experience concern: for simpler queries, users often prioritize low latency, and AdaCoT’s ability to provide fast, direct answers improves overall satisfaction.

Table 7: Human evaluation results. Pairwise comparison of AdaCoT (Adaptive) against Full CoT and No CoT baselines on 100 real-world queries. G:S:B denotes Good:Similar:Bad win rates for AdaCoT.

Comparison	Win Rate (G:S:B)
AdaCoT vs. Full CoT	23% : 51% : 26%
AdaCoT vs. No CoT	45% : 39% : 16%

G USE OF LARGE LANGUAGE MODELS

In accordance with the disclosure policy regarding the use of Large Language Models (LLMs), we detail their role in the preparation of this manuscript.

The use of LLMs was strictly limited to post-writing assistance. Specifically, we utilized an LLM for tasks such as proofreading for grammatical errors and rephrasing sentences to improve clarity.

Crucially, LLMs did not contribute to any of the core intellectual aspects of this research. All conceptual development, including the formulation of the AdaCoT framework, the design of the Pareto optimization approach, the development of the Selective Loss Masking technique, the experimental design, data analysis, and the interpretation of results, were conducted exclusively by the human authors.

Therefore, the LLM functioned as a sophisticated editing and language refinement tool. The fundamental research ideas, experimental work, and conclusions presented in this paper are the sole intellectual product of the authors.