

Rewarding Efficient Reasoning Improves Abstention on Underspecified Tasks in Reasoning Models

Polina Tsvilodub^{1, *}, Max Höth^{2,3}, Michael Franke¹,
Björn Deiseroth^{2,3, †}, Carina Kauf^{*, †}

¹University of Tübingen, ²Aleph Alpha Research, ³Lab1141

Correspondence: polina.tsvilodub@uni-tuebingen.de

Abstract

While modern large reasoning models (LRMs) excel at providing correct answers in many tasks, we provide additional evidence for the observation that they often struggle with a critical capability: knowing when to abstain from answering. We analyze this gap by comparing LRM behavior to results from a human study, revealing that human reasoning effort on unanswerable tasks is upper-bounded by answerable tasks, whereas LRMs waste computational resources by generating longer Chains of Thought (CoTs) on unanswerable than on answerable prompts. To overcome this inefficiency, we take inspiration from a resource-rational perspective on human cognition and introduce a novel GRPO reward that encourages efficient reasoning about whether the task contains all the information needed to solve it. Fine-tuning several 4B LRMs with this reward leads to human-like abstention performance gains (+12.8% on average) while retaining answering capabilities and boosting the models' efficiency (44% shorter CoTs on average).

1 Introduction

Large Language Models (LLMs) have achieved impressive performance across various tasks and domains (Brown et al., 2020; Bubeck et al., 2023). Recently, fine-tuning LLMs specifically for complex reasoning tasks by incentivizing models to produce long chains of thought (CoT) has become common (Wei et al., 2022; Shao et al., 2024; Muenighoff et al., 2025). This approach has given rise to so-called Large Reasoning Models (LRMs) that perform particularly well on tasks such as mathematical reasoning or coding (Guo et al., 2025).

However, while LRMs excel on answerable tasks, their capability to accurately identify *when*

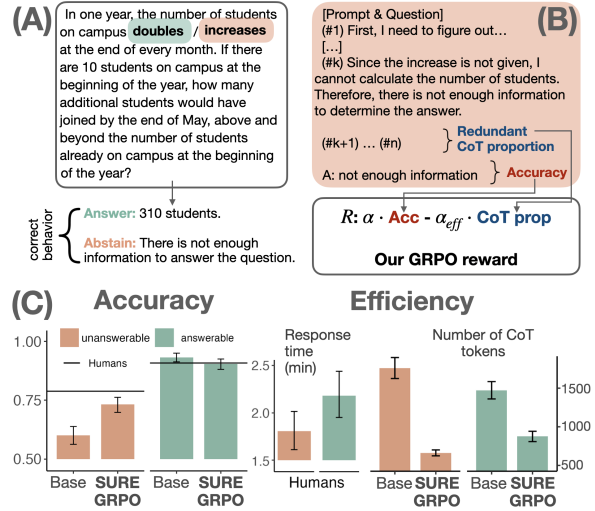


Figure 1: Overview of our approach and contributions. (A): Example of an **abstention** version and an **answerable** version of a question from QuestBench (Li et al., 2025) that was also used for the human study. (B): Overview of the SURE objective proposed in this work. The example CoT chunks are stylized. (C): Our SURE GRPO objective makes LRMs' abstention more human-like. Left: SURE objective improves **abstention** performance of an off-the-shelf reasoning model (Phi-4-mini-reasoning), while retaining **answering** performance. Right: the SURE objective reduces redundant CoT tokens, allocating reasoning resources on abstention in a more human-like way.

not to answer, i.e., when to *abstain*, remains subpar, even though this capability is critical for user-facing deployment (Kirichenko et al., 2025). Abstention is an umbrella term for a range of behaviors where models refuse to directly answer a query, like outputting "I don't know", hedging, or asking for clarification (Kirichenko et al., 2025; Wen et al., 2025). Depending on context, abstention is expected, e.g., (i) when the prompt is underspecified, (ii) the answer is generally unknown or (iii) the prompt should not be answered for safety reasons. Crucially, because real-world user inputs are frequently vague or linguistically underspecified

*Work done while working at Aleph Alpha Research.

†Joint senior authorship.

(Kuhn et al., 2023; Zhang et al., 2024), training effective LRMs requires balancing caution with helpfulness: models should not refuse every underspecified request, but abstain only in critical cases, while maintaining high accuracy on tasks where an answer is expected (Varshney et al., 2024).

In contrast to other related work in this space (c.f. Section 2), here, we take inspiration from a *resource-rational* perspective on human cognition (Lieder and Griffiths, 2020) and results from our own, novel experiment with human participants **to evaluate and improve abstention efficiency and performance of LRMs**. We show that humans identify with high accuracy both when to answer and when to abstain (see Figure 1(C, left) and Section 4), and that their reasoning effort on unanswerable tasks is upper-bounded by answerable tasks (Figure 1(C, right)). By contrast, LRMs often produce excessively long CoTs (i.e., “overthink”) specifically in cases when they should abstain (Figure 1(C, right)). Our contributions are as follows (visualized in Figure 1):

- We conduct a human study, showing that **humans accurately identify when a question is unanswerable, and do so with the same reasoning resources as for answering**.
- We compare reasoning models (4B–32B parameters, six model families) to human results and find a clear misalignment: **LRMs abstain far worse and less efficiently than humans**.
- We propose the **S**ufficiency-aware **R**easoning **E**fficiency (**SURE**) **r**eward for Group Relative Policy Optimization (GRPO): it combines an outcome reward with a process reward that penalizes reasoning beyond the point at which the CoT determines whether task-crucial information is missing. **SURE-fine-tuned LRMs showed improved and more efficient abstention performance, while keeping robust answering capabilities**.

2 Related work

Abstention in LLMs. Abstention capabilities of LLMs have received increasing attention (Kirichenko et al., 2025; Wen et al., 2025), and have been evaluated, e.g., on ambiguous or underspecified tasks (e.g., Slobodkin et al., 2023; Sun et al., 2024; Zhang et al., 2024), on tasks with unknown answers (Amayuelas et al., 2024), or as a possible mitigation of LLM hallucinations (Tonmoy et al.,

2024). Other work has focused on clarification question asking in LLMs (Kundu et al., 2020; Andukuri et al., 2024; Testoni and Fernández, 2024), also when the task is underspecified (Li et al., 2025; Lachenmaier et al., 2026; Wang et al., 2026). A related line of work focuses on evaluating how well LLMs express uncertainty (Kadavath et al., 2022; Tian et al., 2023). Several studies have proposed approaches for improving LLMs’ or LRMs’ abstention capabilities through prompting (Deng et al., 2024) or fine-tuning to produce correct answers (Chen et al., 2025; Zhai et al., 2026), while less work has considered process rewards (Lightman et al., 2024) or the efficiency of the reasoning process on abstention tasks (but see Gu et al., 2026).

Efficiency in LRMs. Efforts to improve CoT efficiency have applied length penalties during RL fine-tuning of LRMs on answerable tasks (Team et al., 2025), often aiming to allocate longer CoTs to harder than to simpler prompts (Ling et al., 2025; Xiang et al., 2025). Early-exiting approaches force efficient termination of CoTs during inference (Yang et al., 2025; Wang et al., 2025). Most abstention work focuses only on LLMs, while work on answerable tasks has also compared LLMs to human reasoning (e.g., Eisape et al., 2024; Liu et al., 2024). We take inspiration from de Varda et al. (2025) who show that LRMs’ CoTs align with human reaction times (RTs) and capture reasoning demands on various answerable tasks, and evaluate human performance also on unanswerable tasks to ground the assessment of LRMs on abstention.

Resource Rationality in Humans. Work within the resource rationality framework has shown that humans flexibly allocate their reasoning resources, often measured through time allocated for solving a task (Lieder and Griffiths, 2020), and higher reasoning time often leads to more accurate task performance (Wickelgren, 1977). Yet while previous work has investigated factors influencing abstention (Undorf et al., 2021; Law et al., 2022) and clarification question production (Clark and Wilkes-Gibbs, 1986; Purver et al., 2001; Ali et al., 2026; Tsvilodub et al., 2026) in humans, the resource allocation in reasoning about abstention remains less clear.

3 Experiment Design & Dataset

Following definitions by Kirichenko et al. (2025); Wen et al. (2025), we investigate whether LRMs refuse to answer in any form (e.g., hedging, out-

putting “I don’t know”, or asking for clarification) when given queries for which abstention is expected. An LLM judge provides binary annotations of the outputs (see Section 5 for details and Appendix A.1 for the prompt). Humans are evaluated analogously, via a binary forced-choice task asking whether a question is answerable (Section 4).

Our experiments vary the *question type* (answerable vs. unanswerable) by using evaluation datasets which contain both question types (closely matched in difficulty) for evaluating and subsequent fine-tuning (Section 6) LRMs: QuestBench (Li et al., 2025) and AbstentionBench (Kirichenko et al., 2025). An example from AbstentionBench is shown in Figure 1(A).

For QuestBench, we use the GSM-Q subset as unanswerable questions which consist of grade school level math tasks from the GSM8K dataset (derived from Li et al., 2024, which is used for answerable questions). Li et al. (2025) constructed GSM-Q by removing a single variable in each question, resulting in tasks where crucial information is missing. The questions are paired. The questions also vary with respect to the number of steps that are needed to solve them (i.e., in their difficulty).

AbstentionBench (Kirichenko et al., 2025) consists of a combination of 20 datasets covering different tasks with both answerable and unanswerable prompts. The tasks cover different reasons for abstention (underspecified context, but also underspecified intent, stale data, false premise, or where the answer is unknown or subjective). We exclude questions with more than 2048 tokens.

We use 500 test samples from each benchmark across all reported evaluations. The QuestBench split of the unanswerable test set evaluates abstention on underspecified prompts; the AbstentionBench split evaluates abstention on diverse tasks, approximately balanced across the benchmark’s distribution of abstention reasons.

4 Humans Answer and Abstain Accurately & Efficiently

To ground LRM evaluations, we draw on insights about human behavior on the same benchmarks. If human task solving is conceptualized as a search over a problem space (Simon and Newell, 1971), one hypothesis is that the search will be terminated as soon as a gap in the problem representation (i.e., missing information) is encountered, predicting that the resources for abstention are upper-bounded

by the respective answerable tasks. Here, we empirically compare human performance and reasoning effort on unanswerable vs. answerable tasks through an exploratory web-based experiment, investigating the following questions:¹ (1) Do humans accurately identify whether a question is (un)answerable? (2) Are human completion times (i.e., “reasoning effort”) on unanswerable tasks upper-bounded by answerable tasks? (3) When participants are additionally incentivized to accurately complete certain trials, do their performance and reasoning effort change?

Materials, Procedure & Participants. We devised a $2 \times 2 \times 3$ design with factors *question type* (answerable vs. unanswerable), *question domain* (math tasks from QuestBench vs. common sense questions from AbstentionBench), and *importance* of solving the task (high vs. default vs. low; operationalized through different numbers of points for completing a given trial correctly, and bonus payments proportional to achieved points above a threshold). The default condition only contained task instructions. The test items were selected by randomly sampling six items per domain and question type, resulting in 24 items. The items were additionally filtered by the authors for naturalness. Each item was used in the three importance conditions. Full details about the materials are reported in Appendix B.

Participants ($N = 126$, recruited via Prolific) were self-reported native English speakers with approval rates over 95% and at least five prior studies. Because participants might hesitate to abstain in an experimental setting, each participant first viewed examples of each question type (see Appendix B). Then, they completed six main trials (one per question type $\times$ importance condition, with three trials per domain), and one attention check.

On each trial, participants read a question, embedded in a randomly sampled importance prompt condition. They first only saw a forced choice (FC) task where they indicated whether the question is “Answerable” or “Not definitively answerable”. Once they answered the FC task, two text boxes appeared, one for an explanation of the solution steps, and the other either for the final answer to the answerable questions, or for an explanation of what information is missing in the unanswerable questions. The attention check trials were visually

¹The materials can be viewed at: <https://github.com/polina-tsvilodub/reasoning-under-missing-info>.

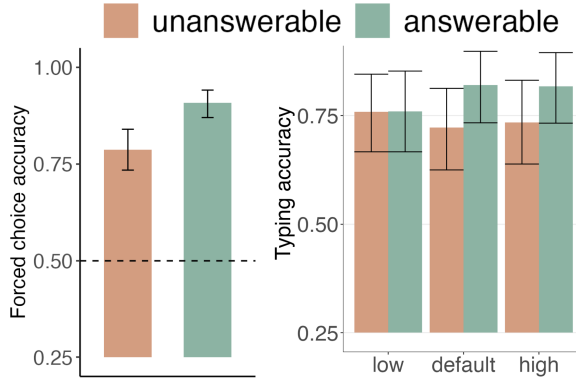


Figure 2: **Left:** FC accuracy by question type across importance conditions. **Right:** Typing accuracy (correctness of answers or identified missing information) across importance conditions. Error bars are bootstrapped 95%-CIs.

identical, and asked participants to provide specific answers. Participants took 17 minutes on average, and were reimbursed £1.20 with up to £0.20 bonus. Full experiment details are reported in Appendix B.

Results. We analyze the accuracy of the responses provided to the FC question, the accuracy of the typed answers (pooled across the text fields), and reaction time (RT) per trial. After excluding participants who failed any part of the attention check, we analyze data from 83 participants. We analyze all results with Bayesian mixed-effect regression models, always including maximal converging random effects structure. Posterior means and 95% credible intervals are reported.

1. Humans identify when to answer or abstain well above chance, but more accurately when to answer. The forced choice accuracy for answerable and unanswerable questions across importance conditions and domains is shown in Figure 2 (left). We analyze the forced-choice answer accuracy using a logistic regression model.² Humans were credibly more accurate at identifying answerable questions ($\beta = 5.86[0.21, 14.02]$ across domains), driven by credible differences on math questions ($\beta = 16.72[6.78, 34.0]$), but not in the common sense questions ($\beta = 6.71[-14.48, 29.26]$). Still, humans identified when to abstain credibly above chance (posterior probability of effect 96%).

²Model in R syntax: `accuracy ~ domain * prompt * question_type + (1 + domain * prompt * question_type | subject) + (1 | itemId)`

2. Humans reasoning effort on unanswerable tasks is upper-bounded by answerable tasks.

The total response time (RT) taken by the participants on each answerable and unanswerable trial is shown in Figure 1(C) (right). Separate analyses of FC and typing times are reported in Appendix B. All RTs are log-transformed. We use a Bayesian linear mixed effect regression model.³ Participants’ RTs didn’t differ credibly between the answerable and abstention tasks across domains and conditions ($\beta = 0.03[-0.14, 0.19]$). However, humans did have a marginally credibly higher RT for answerable than abstention math tasks ($\beta = 0.13[-0.00, 0.27]$), but not common sense tasks ($\beta = 0.03[-0.64, 0.72]$). In general, RTs were higher on math than the common sense questions (across question types: $\beta = 0.54[0.12, 0.98]$).

3. Humans are more accurate and reason slower when incentivized accordingly.

The accuracy of the typed answers by importance condition is shown in Figure 2 (right). Humans are marginally, but credibly more accurate in the high-importance than the low-importance condition across domains (posterior probability 96.45%), but visual inspection suggests that the trend is driven by answerable questions (Figure 2, right). Additionally, humans took longer in the high than default prompt condition ($\beta = 0.12[0.01, 0.23]$), an effect driven primarily by differences in the common sense domain (high vs. default: $\beta = 0.18[0.01, 0.35]$).

In sum, these results suggest that humans accurately identify when to abstain, while still answering accurately, and do so with roughly the same reasoning resources spent on answering and abstention. We consider these patterns as a human baseline for comparison with off-the-shelf LRMs.

5 LRMs’ Abstention Performance Is Worse Than in Instruct Models

We first evaluate reasoning models from six different families on our evaluation set (Section 3). We evaluate all models in a free generation setting, with maximally 10000 new tokens. The default generation configuration of each model was used. We use an LLM as a judge (across experiments: Qwen3-30B-A3B-Instruct-2507, Team, 2025) to evaluate whether an answer is an abstention, and to evaluate the correctness of outputs for answerable

³Model in R syntax: `log(RT) ~ domain * importance * question_type + (1 + question_type + domain | subject) + (1 | itemId)`

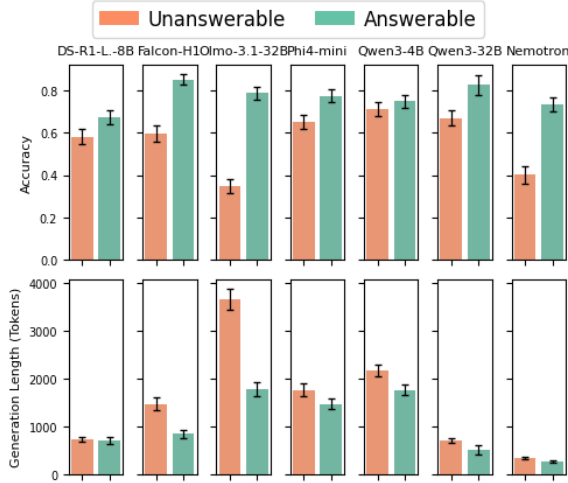


Figure 3: **Reasoning models tend to perform worse on abstention tasks, and inefficiently allocate the compute budget by producing verbose CoTs when they should abstain.** Accuracy (upper row) and length of generations (in tokens, bottom row) of reasoning models from six different families on answerable vs. unanswerable questions (x-axis), averaged across two benchmarks. Error bars show 95% bootstrapped CIs.

questions. The prompts are in Appendix A.1, A.2. To evaluate the initial LRM performance in the strongest baseline setting, we designed a prompt stating that the task might be missing information, in which case the model should abstain and stop reasoning as soon as possible (the full prompt and details are in Appendix A.3).

We evaluate a range of models that have both an instruction fine-tuned and a reasoning variant available: Qwen3-4B-2507, Qwen3-32B (Team, 2025), Olmo-3.1-32B (Olmo et al., 2025), Falcon-H1-7B (Team et al., 2026), Phi-4-mini (Xu et al., 2025) DeepSeek-R1-Distill-Llama-8B (DeepSeek-AI, 2025).⁴ For subsequent fine-tuning, for comparability across architectures, we focus on models with approximately 4B parameters: Qwen3-4B-Thinking-2507 (Team, 2025), Phi4-mini-reasoning (Xu et al., 2025) and NVIDIA-Nemotron-3-nano-4B-BF16 (used in the enabled thinking mode, Blakeman et al., 2025).

LRMs Perform Worse than Instruction-Tuned Models on Abstention in Terms of Performance and Efficiency. We first compare the performance of seven different reasoning models on unanswerable questions vs. answerable questions. Results across the benchmarks are shown in Figure 3.

⁴For DeepSeek, the Llama-3.1-8B-Instruct was used as the instruction-tuned variant.

Accuracy on answerable questions measures answer correctness, while accuracy on unanswerable samples measures whether the model correctly abstains rather than guesses. We analyze the results with a Bayesian logistic regression model, regressing the accuracy against the question type, model, and their interaction.⁵ For each single LRM as well as across models, the abstention performance was credibly lower than answering performance ($\beta = 0.96[0.87, 1.06]$) (Figure 3, upper row). All LRMs abstain credibly worse than humans ($\beta = 1.03[0.73, 1.35]$). For three of the models, the gap between the abstention and answering performance was larger than for humans with a posterior probability over 85%.

Additionally, the models’ CoTs on unanswerable questions are at least as long as those for the answerable condition (Figure 3, bottom row). They are even longer than the answerable question CoTs with a posterior probability of 85%.⁶ This suggests that the models might tend towards more inefficient CoTs or loops when faced with unanswerable questions which might be missing information samples.

To investigate whether these abstention dynamics are due to the reasoning fine-tuning, we compare the reasoning models to their respective instruction fine-tuned versions. The accuracies and CoT lengths on answerable vs. unanswerable questions across benchmarks are shown in Figure 6 in the Appendix. The instruct models performed credibly better on abstention than answering tasks, driven by four of the six model pairs.⁷ The gap between the answering and abstention performance was credibly bigger in the reasoning than instruction models ($\beta = 1.67[1.54, 1.80]$). The lengths of the CoTs in the instruct models did not credibly differ between answering and abstention ($\beta = -24.49[-759.29, 713.78]$).⁸ Therefore, the gap in performance and the CoT allocation is a consequence of the reasoning fine-tuning of the models.

LRMs Identify Missing Information in Abstention Tasks Early on. As hypothesized in Section 4, a rational allocation of reasoning resources might require stopping the reasoning once miss-

⁵Model in R syntax: `accuracy ~ question_type * model_name`. Note that we include human results and humans as a “model” to enable comparison.

⁶Underlying linear regression model: `avg tokens ~ question_type`.

⁷Logistic regression model in R syntax: `accuracy ~ question_type * model_type + (1 | model_family)`.

⁸Underlying linear regression model: `avg tokens ~ question_type * model_type`.

ing information is identified. We explore to which extent LRMs’ CoTs reflect such a strategy by annotating whether the CoT sentences contain reasoning about missing information for the task (see Section 6.1 for details) and find a striking discrepancy. Figure 7 in the Appendix shows that only up to $\sim 25\%$ of the CoT on abstention tasks is required until missing information is identified, suggesting that a large portion of the CoT is redundant and might lead to the observed misalignment with human behavior. We use these results as a departure point for developing a fine-tuning approach for improving efficiency and abstention performance of LRMs, described next.

6 Fine-tuning LRMs to Reason about Missing Information Leads to Efficient Human-like Accurate Abstention

The results from Section 5 suggest that for abstention tasks, a large portion of the CoT which might contain, e.g., “self-verification” like repeating solution attempts (Muennighoff et al., 2025), might worsen performance. Therefore, we develop an objective for reinforcement learning (RL) fine-tuning with GRPO specifically aiming to reduce the redundancy in the CoT (c.f. Figure 7), in order to improve the efficiency and accuracy when abstention is needed while maintaining response accuracy and propensity on sufficiently specified tasks. We construct a fine-tuning dataset (disjoint from the evaluation samples) from AbstentionBench and QuestBench GSM-Q, with about 60% unanswerable and 40% answerable samples.⁹

6.1 SURE Fine-Tuning of Reasoning Models

GRPO. Group Relative Policy Optimization (GRPO) (Shao et al., 2024) aligns language models without a value network by estimating advantages $A_i = (r_i - \mu)/\sigma$ from a group of G outputs o_i sampled for query q , where μ and σ are the group’s reward mean and standard deviation. The policy π_θ maximizes the objective $\mathcal{J}(\theta)$:

$$\mathcal{J}(\theta) = \mathbb{E} \left[\frac{1}{G} \sum_{i=1}^G \left(\min \left(\rho_i A_i, \right. \right. \right. \quad (1)$$

$$\left. \left. \left. \text{clip}(\rho_i, 1 - \epsilon, 1 + \epsilon) A_i \right) - \beta \mathbb{D}_{\text{KL}}(\pi_\theta \| \pi_{\text{ref}}) \right) \right],$$

⁹We use GSM8K (Cobbe et al., 2021) for the answerable counterpart to QuestBench GSM-Q.

where $\rho_i = \pi_\theta(o_i|q)/\pi_{\theta_{\text{old}}}(o_i|q)$ is the policy ratio, ϵ is the clip margin, and β scales the token-level KL divergence $\mathbb{D}_{\text{KL}}$ from the reference model π_{ref} . Typically, the reward r_i relies solely on a final accuracy indicator, $r_i = \text{acc}(o_i)$ (*accuracy-only* reward).

Baseline reward. To explicitly suppress verbosity, we modify the reward to incorporate a standard length penalty: $r_i = \text{acc}(o_i) - \lambda|o_i|$, where $|o_i|$ is the output length and λ controls the penalty magnitude (*accuracy + length penalty* reward).

SURE reward. To explicitly incentivize the model to halt its reasoning once it detects missing information, we formulate a composite **S**ufficiency-aware **R**easoning **E**fficiency reward. An intuitive explanation of the reward is presented in Figure 1(B). For a given sampled output o_i , the total reward r_i combines an outcome accuracy score with a process-focused efficiency term $e(o_i)$:

$$r_i = 0.5 + \alpha \cdot \text{acc}(o_i) + \alpha_{\text{eff}} \cdot e(o_i) \quad (2)$$

where $\text{acc}(o_i)$ measures the binary task accuracy, evaluating both answerable and abstention cases. The term $e(o_i)$ captures the proportion of the reasoning trace that is redundant:

$$e(o_i) = \frac{n - k}{n} \quad (3)$$

Here, n is the total number of sentences in the reasoning trace of o_i , and k is the index of the first sentence where missing task-relevant information is identified. The scalar weights α and α_{eff} balance the outcome and process signals and are set to 0.5 across experiments. To localize k within a reasoning trace, we generate a complete model rollout o_i , segment this rollout into discrete chunks based on punctuation boundaries, and apply an LLM as a judge (Qwen3-30B-A3B-Instruct-2507) to assess each chunk individually. For each chunk, the judge returns a binary annotation whether it mentions that the task is missing information. The specific judge prompt and further implementation details are provided in Appendix A.4. We note that the reward in Eq. 2 can be optimized through two strategies: either decreasing n while k remains constant, or increasing k while n remains constant. We explore empirically what happens to the models when using this objective within GRPO training.

Training details. For the baseline accuracy + length penalty reward, we set $\lambda = 0.0001$ per token that exceeds the minimal CoT length of 200 tokens.

Model	Task	# CoT tokens (default)	Δ CoT (high-imp. – default)
SURE-GRPO	unanswerable	662 (-770) [630, 693]	354 [312, 395]
SURE-GRPO	answerable	766 (-417) [724, 808]	134 [106, 164]
Acc. Only	unanswerable	1363 (-69) [1302, 1423]	341 [280, 401]
Acc. Only	answerable	1222 (+39) [1162, 1281]	158 [113, 200]
Acc.+Len.	unanswerable	202 (-1230) [189, 215]	103 [85, 126]
Acc.+Len.	answerable	235 (-948) [221, 248]	91 [80, 104]
Base	unanswerable	1432 [1365, 1499]	293 [157, 437]
Base	answerable	1183 [1125, 1241]	160 [107, 216]

Table 1: Differences in the CoTs (in number of tokens) resulting from different fine-tuning objectives. # CoT tokens (with default prompting) decreases significantly with SURE-GRPO, but more moderately than with Acc.+Len. Δ CoT shows the difference in the number of CoT tokens produced in different prompting conditions; higher values indicate stronger sensitivity to the user’s request for longer reasoning.

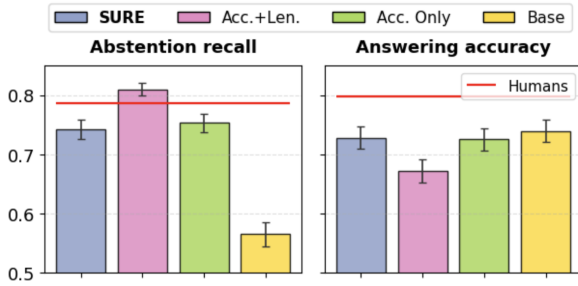


Figure 4: **SURE GRPO training results:** Performance of models averaged across three families trained with our SURE GRPO objective against two strong baseline objectives (“Accuracy + Length penalty” and “Accuracy-only”), relative to the performance of the initial LRMs (“Base”). 95% bootstrapped CIs are shown. **Left:** recall on abstention tasks; increase indicates that models guessed answers less. **Right:** accuracy of provided answers on the answerable questions; difference to the base models indicates deterioration of the models’ capabilities. Only SURE objective combines stable performance and efficiency gains across models (see Table 1).

For all objectives, rollouts that do not generate the EOS token within the token budget of 4096 tokens receive $r_i = 0$. We use LoRA (Hu et al., 2022) to train all three models with all three rewards with the following hyperparameters. We use a group size of 8, allowing for maximally 4096 generated tokens per rollout. We use the sampling temperature $\tau = 0.6$ for Qwen and Phi, and $\tau = 1.0$ for Nemotron. We use Adam with a learning rate $1e-4$, and an effective batch size of 64. We use LoRa with $r = 8$, $\alpha = 16$, dropout 0, targeting Q and V of the transformer blocks. We train Qwen and Phi for 200 steps, and Nemotron for 50 steps (we perform the early stopping based on gradient dynamics).

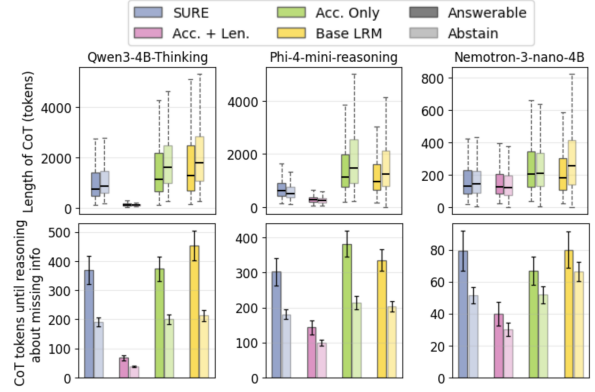


Figure 5: Structure of the CoTs on abstention tasks. **Upper row:** Distributions of the total number of tokens in the CoT. SURE reward tends to maintain a more natural distribution of CoT lengths than the “Acc. + Length Penalty” baseline. **Bottom row:** The number of tokens in the CoT until the model starts reasoning about missing information. SURE objective retains a similar CoT length as the base reasoning model helping to retain its reasoning abilities, while reducing the length of the remaining CoT.

6.2 Results

All three rewards (SURE reward, accuracy-only, accuracy + length penalty) led to an improvement of the LRMs’ ability to recognize when to abstain: Figure 4 (left) shows a significantly improved abstention recall over the initial reasoning model for all three model families, in part achieving human-level performance (+0.173 on average). Crucially, Figure 4 (right) shows that only our SURE objective retains the models’ answering correctness on the answerable questions close to the original LRMs’ performance (-0.012), while a naive length penalty leads to a deterioration of answering capabilities (-0.136). The accuracy-only

objective leads to marginal retention of answering accuracy (-0.013). The different rewards also maintained model performance on other reasoning benchmarks (see Figure 8 in the Appendix). Finally, the two rewards with an efficiency component (SURE, acc. + len. penalty) led to a compression of the CoTs across models (Table 1, # CoT tokens). The SURE reward leads to a more moderate efficiency gain (44% shorter CoTs across models and question types) than the acc. + len. baseline, but a bigger gain than the accuracy-only baseline, without negatively impacting the reasoning structure on other benchmarks (see Figure 9 in the Appendix). Across rewards, the efficiency gains are higher for abstention than answerable samples (CoT reduction relative to the base model: 56% vs. 46% shorter).

Additionally, we explore how flexible the CoT remains after fine-tuning. Figure 5 (upper row) shows the distribution of the CoT length, indicating that the SURE reward seems to retain more variability of the CoT lengths while still making them shorter. The naive length penalty tends towards a collapse of the CoT length to a narrower range. The accuracy-only reward approximately maintains the CoT length variability.

To investigate whether the full reward was optimized by compressing n or shifting k , we plot k in Figure 5 (bottom row). It shows that the SURE objective identified missing information approximately at the same position in the CoT as the base LRM, while Acc. + Len. shifted the first occurrence of the reasoning significantly earlier in the CoT. This indicates that the SURE objective allowed for a more flexible model-dependent identification of k , that both improved abstention and retained reasoning abilities, while improving efficiency.

Finally, we investigate to which extent the fine-tuned models generalize to other prompts than the one it was fine-tuned with, in particular, a “high-importance” prompt stating that a task is very important and the reasoning should be detailed (the full prompt is in Appendix A.5; similar to the high-importance prompt in the human study in Section 4). If the fine-tuned models retain reasoning flexibility during fine-tuning towards short CoTs, they will produce longer CoTs given the “high-importance” than the default evaluation prompt. Table 1 (Δ CoT) shows the differences in the CoT lengths between the two prompts. The base LRMs and all fine-tuned models produced longer CoTs than with the default prompt for both question

types. The magnitude of the difference was significantly larger for the SURE and accuracy-only rewards than acc. + len. models, suggesting lower CoTs flexibility under a naive length penalty. For all models, the increase in CoT length was stronger for abstention than answerable samples — an interesting difference to human behavior.

Overall, we speculate that the more moderate efficiency gain of the SURE objective, driven by the models’ internal optimal reasoning structure that is simply reinforced, helps to main the CoT flexibility and the answering capabilities of the model, while improving abstention performance.

7 Discussion

This paper offers three findings about the abstention capabilities (i.e., identifying when not to answer because critical information is missing) of large reasoning models: (1) due to reasoning post-training, LRMs abstention capabilities deteriorate and waste CoT tokens, (2) this behavior diverges from human behavior, evaluated here in an experiment, (3) one reason for the inefficiency and divergence from humans is that LRMs’ CoT does not halt when missing information is identified. We propose the SURE reward for GRPO fine-tuning that encourages efficient reasoning about whether the task contains all necessary information. SURE-fine-tuning leads to substantial gains both on abstention performance and efficiency across three model families.

Our results suggest several avenues for future work. The process reward in SURE focuses on identifying missing information for the task. While this is a reasonably general starting point, applicable to tasks like in QuestBench and AbstentionBench, other forms of process supervision (e.g., reasoning about norms) may be needed for abstention, e.g., for safety reasons. The LLM as a judge implementation could accommodate that. While the test set from AbstentionBench already covers a variety of tasks, the fine-tuned models should also be evaluated on more abstention datasets.

The human study results also open avenues for future work. When humans answered unanswerable questions, it was often due to specific assumptions and task ambiguity resolution, which should be compared to assumptions made by LRMs. Additionally, more explicit reasoning elicitation (e.g., via think-aloud studies, Wurgaft et al., 2025) is needed to disentangle to which extent human abstention-reasoning reflects identifying missing

information or is driven by meta-cognitive uncertainty (Ackerman and Thompson, 2017). Finally, the efficacy of the process supervision in SURE invites exploring to which extent it might be leveraged to improve not only LRMs’ general abstention capabilities, but to improve clarification question asking, to build both helpful and calibrated LRMs.

Limitations

The reported experiments make number of design choices, and future work should examine to which extent the results and the advantage of SURE generalize beyond these configurations.

First, the presented experiments focus on fine-tuning only LRMs with around 4B parameters. Future work should investigate how well the advantages of SURE generalize to models of other sizes and families. Additionally, all experiments used GRPO (Shao et al., 2024) for fine-tuning, but future work should explore how models trained with other common algorithms like Dr. GRPO (Liu et al., 2025) or DPO (Rafailov et al., 2023) might benefit from SURE, and how its efficacy interacts with whether supervised fine-tuning (SFT), e.g., on examples of correct abstention, is performed first.

We fine-tune and evaluate all models only in a free generation setting, but future work should also explore how the fine-tuned models will generalize, e.g., to multiple choice tasks. Moreover, to allow for strong baseline objectives and ensuring comparability across fine-tuning objectives, we use the same prompt across experiments which states that the task might be missing information. We conducted exploratory evaluations with variations of the prompts, and found qualitatively robust LRM performance. While the results highlight the gap in the performance of initial LRMs even with such a strong baseline prompt, and the answering capabilities are retained after SURE fine-tuning with the prompt, future work should assess the effect of this particular prompting more comprehensively.

Using SURE requires resources for accessing an online LLM as a judge to calculate both the accuracy and the process rewards. If no judge is available, at least the process reward could potentially be approximated through an alternative approach, e.g., searching keywords like “unsolvable task” and “information is missing”. Initial analyses suggest that a curated list of keywords indeed helps identify at least some sentences in the CoT reasoning about sufficiency of information, although with less

accurately.

We focused only on English, and used benchmarks which have been available for a few years, such that the training data of the initial LRMs might be contaminated with some of the data.

We conducted only limited exploratory analyses of the LRMs’ CoTs, particularly on samples where the question was answered instead of abstaining. The analyses suggest interesting differences compared to human reasoning: while humans might make assumptions based on their personal information (e.g., when asked “Who is the prime minister?”, they might name the minister of their country of residence), LRMs instead tend to make estimates of missing numbers (e.g., for the example in Figure 1(A) the CoT might state “A common number of students on a campus is X” and perform calculations with that). These results outline an avenue for more comprehensive comparisons of the assumptions made in failure cases by humans and LRMs, as well as before and after fine-tuning, in future work.

Finally, the reported human study also had some limitations. First, the estimates of human reasoning effort were accessed through a somewhat indirect measurement, namely reaction times; while even this coarse-grained estimate established a baseline for LRMs, other more direct methods like think-aloud (Wurgaf et al., 2025) should be employed together with RT measurements for more robust conclusions about human cognition. Finally, for naturalness reasons the importance manipulations were operationalized differently in humans and LRMs (through points and bonus payments vs. through explicit prompting, respectively), which might have led to the observed difference in the effect of the manipulation.

Acknowledgments

We acknowledge the use of LLMs for coding and minor rephrasing of the original text, which was fully our own, and we carefully reviewed all LLM suggestions. MF is a member of the Machine Learning Cluster of Excellence at University of Tübingen, EXC number 2064/2 – Project number 39072764 and his contribution to this work was supported by the Volkswagen Foundation through a Momentum grant. PT is funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under project ID 579368432.

References

- Rakefet Ackerman and Valerie A Thompson. 2017. Meta-reasoning: Monitoring and control of thinking and reasoning. *Trends in cognitive sciences*, 21(8):607–617.
- Manar Ali, Judith Sieker, Sina Zarrieß, and Hendrik Buschmeier. 2026. Reference games as a testbed for the alignment of model uncertainty and clarification requests. *arXiv preprint arXiv:2601.07820*.
- Alfonso Amayuelas, Kyle Wong, Liangming Pan, Wenhu Chen, and William Yang Wang. 2024. Knowledge of knowledge: Exploring known-unknowns uncertainty with large language models. In *Findings of the Association for Computational Linguistics: ACL 2024*, pages 6416–6432.
- Chinmaya Andukuri, Jan-Philipp Fränken, Tobias Gerstenberg, and Noah D Goodman. 2024. Star-gate: Teaching language models to ask clarifying questions. *arXiv preprint arXiv:2403.19154*.
- Aaron Blakeman, Aaron Grattafiori, Aarti Basant, Abhibha Gupta, Abhinav Khattar, Adi Renduchintala, Aditya Vavre, Akanksha Shukla, Akhiad Bercovich, Aleksander Ficek, and 1 others. 2025. Nemotron 3 nano: Open, efficient mixture-of-experts hybrid mamba-transformer model for agentic reasoning. *arXiv preprint arXiv:2512.20848*.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, and 1 others. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Sébastien Bubeck, Varun Chandrasekaran, Ronen Eldan, Johannes Gehrke, Eric Horvitz, Ece Kamar, Peter Lee, Yin Tat Lee, Yuanzhi Li, Scott Lundberg, and 1 others. 2023. Sparks of artificial general intelligence: Early experiments with gpt-4. *arXiv preprint arXiv:2303.12712*.
- Lida Chen, Zujie Liang, Xintao Wang, Jiaqing Liang, Yanghua Xiao, Feng Wei, Jinglei Chen, Zhenghong Hao, Bing Han, and Wei Wang. 2025. Teaching large language models to express knowledge boundary from their own signals. In *Proceedings of the 3rd Workshop on Towards Knowledgeable Foundation Models (KnowFM)*, pages 26–39.
- Herbert H Clark and Deanna Wilkes-Gibbs. 1986. Referring as a collaborative process. *Cognition*, 22(1):1–39.
- Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, and 1 others. 2021. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*.
- Andrea Gregor de Varda, Ferdinando Pio D’Elia, Hope Kean, Andrew Lampinen, and Evelina Fedorenko. 2025. The cost of thinking is similar between large reasoning models and humans. *Proceedings of the National Academy of Sciences*, 122(47):e2520077122.
- DeepSeek-AI. 2025. [Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning](#). Preprint, arXiv:2501.12948.
- Yang Deng, Yong Zhao, Moxin Li, See Kiong Ng, and Tat-Seng Chua. 2024. Don’t just say “i don’t know”! self-aligning large language models for responding to unknown questions with explanations. In *Proceedings of the 2024 conference on empirical methods in natural language processing*, pages 13652–13673.
- Tiwalayo Eisape, Michael Tessler, Ishita Dasgupta, Fei Sha, Sjoerd Steenkiste, and Tal Linzen. 2024. A systematic comparison of syllogistic reasoning in humans and language models. In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 8425–8444.
- Renjie Gu, Jiaxu Li, Yihao Wang, Yun Yue, Hansong Xiao, Yefei Chen, Yuan Wang, Chunxiao Guo, Pei Wei, Jinjie Gu, and 1 others. 2026. Bridging the detection-to-abstention gap in reasoning models under insufficient information. *arXiv preprint arXiv:2605.28070*.
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Peiyi Wang, Qihao Zhu, Runxin Xu, Ruoyu Zhang, Shirong Ma, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, and 175 others. 2025. [Deepseek-r1 incentivizes reasoning in llms through reinforcement learning](#). *Nature*, 645(8081):633–638.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2020. Measuring massive multitask language understanding. *arXiv preprint arXiv:2009.03300*.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Liang Wang, Weizhu Chen, and 1 others. 2022. Lora: Low-rank adaptation of large language models. *ICLR*, 1(2):3.
- Saurav Kadavath, Tom Conerly, Amanda Askell, Tom Henighan, Dawn Drain, Ethan Perez, Nicholas Schiefer, Zac Hatfield-Dodds, Nova DasSarma, Eli Tran-Johnson, and 1 others. 2022. Language models (mostly) know what they know. *arXiv preprint arXiv:2207.05221*.
- Polina Kirichenko, Mark Ibrahim, Kamalika Chaudhuri, and Samuel J Bell. 2025. Abstentionbench: Reasoning llms fail on unanswerable questions. *arXiv preprint arXiv:2506.09038*.

- Lorenz Kuhn, Yarin Gal, and Sebastian Farquhar. 2023. Semantic uncertainty: Linguistic invariances for uncertainty estimation in natural language generation. *arXiv preprint arXiv:2302.09664*.
- Souvik Kundu, Qian Lin, and Hwee Tou Ng. 2020. [Learning to identify follow-up questions in conversational question answering](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 959–968, Online. Association for Computational Linguistics.
- Clara Lachenmaier, Hannah Bultmann, and Sina Zarrieß. 2026. Talking to a know-it-all gpt or a second-guesser claude? how repair reveals distinct multi-turn behavior in llms. In *Proceedings of the 64th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 14312–14325.
- Marvin KH Law, Lazar Stankov, and Sabina Kleitman. 2022. I choose to opt-out of answering: Individual differences in giving up behaviour on cognitive tests. *Journal of Intelligence*, 10(4):86.
- Belinda Z Li, Been Kim, and Zi Wang. 2025. Quest-bench: Can llms ask the right question to acquire information in reasoning tasks? *arXiv preprint arXiv:2503.22674*.
- Qintong Li, Leyang Cui, Xueliang Zhao, Lingpeng Kong, and Wei Bi. 2024. Gsm-plus: A comprehensive benchmark for evaluating the robustness of llms as mathematical problem solvers.
- Falk Lieder and Thomas L Griffiths. 2020. Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behavioral and brain sciences*, 43:e1.
- Hunter Lightman, Vineet Kosaraju, Yuri Burda, Harrison Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. 2024. Let’s verify step by step. In *International Conference on Learning Representations*, volume 2024, pages 39578–39601.
- Zehui Ling, Deshu Chen, Hongwei Zhang, Yifeng Jiao, Xin Guo, and Yuan Cheng. 2025. Fast on the easy, deep on the hard: Efficient reasoning via powered length penalty. *arXiv preprint arXiv:2506.10446*.
- Ryan Liu, Jiayi Geng, Addison J Wu, Ilia Sucholutsky, Tania Lombrozo, and Thomas L Griffiths. 2024. Mind your step (by step): Chain-of-thought can reduce performance on tasks where thinking makes humans worse. *arXiv preprint arXiv:2410.21333*.
- Zichen Liu, Changyu Chen, Wenjun Li, Penghui Qi, Tianyu Pang, Chao Du, Wee Sun Lee, and Min Lin. 2025. Understanding r1-zero-like training: A critical perspective. *arXiv preprint arXiv:2503.20783*.
- Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke Zettlemoyer, Percy Liang, Emmanuel Candès, and Tatsunori B Hashimoto. 2025. s1: Simple test-time scaling. In *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*, pages 20286–20332.
- Team Olmo, Allyson Ettinger, Amanda Bertsch, Bailey Kuehl, David Graham, David Heineman, Dirk Groeneveld, Faeze Brahman, Finbarr Timbers, Hamish Ivison, Jacob Morrison, Jake Poznanski, Kyle Lo, Luca Soldaini, Matt Jordan, Mayee Chen, Michael Noukhovitch, Nathan Lambert, Pete Walsh, and 49 others. 2025. [Olmo 3](#). *Preprint*, arXiv:2512.13961.
- Matthew Purver, Jonathan Ginzburg, and Patrick Healey. 2001. On the means for clarification in dialogue. In *Proceedings of the Second SIGdial Workshop on Discourse and Dialogue*.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. 2023. Direct preference optimization: Your language model is secretly a reward model. *Advances in neural information processing systems*, 36:53728–53741.
- Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, YK Li, Yang Wu, and 1 others. 2024. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*.
- Herbert A Simon and Allen Newell. 1971. Human problem solving: The state of the theory in 1970. *American psychologist*, 26(2):145.
- Aviv Slobodkin, Omer Goldman, Avi Caciularu, Ido Dagan, and Shauli Ravfogel. 2023. The curious case of hallucinatory (un) answerability: Finding truths in the hidden states of over-confident large language models. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 3607–3625.
- Yuhong Sun, Zhangyue Yin, Qipeng Guo, Jiawen Wu, Xipeng Qiu, and Hui Zhao. 2024. Benchmarking hallucination in large language models based on unanswerable math word problem. In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 2178–2188.
- Oyvind Tafford, Bhavana Dalvi, and Peter Clark. 2021. [ProofWriter: Generating implications, proofs, and abductive statements over natural language](#). In *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021*, pages 3621–3634, Online. Association for Computational Linguistics.
- Falcon LLM Team, Iheb Chaabane, Puneesh Khanna, Suhail Mohmad, Slim Frikha, Shi Hu, Abdalgader Abubaker, Reda Alami, Mikhail Lubinets, Mohamed El Amine Seddik, and Hakim Hacid. 2026. [Falcon-h1r: Pushing the reasoning frontiers with a hybrid model for efficient test-time scaling](#). *Preprint*, arXiv:2601.02346.

- Kimi Team, Angang Du, Bofei Gao, Bowei Xing, Changjiu Jiang, Cheng Chen, Cheng Li, Chenjun Xiao, Chenzhuang Du, Chonghua Liao, and 1 others. 2025. Kimi k1. 5: Scaling reinforcement learning with llms. *arXiv preprint arXiv:2501.12599*.
- Qwen Team. 2025. [Qwen3 technical report](#). *Preprint*, arXiv:2505.09388.
- Alberto Testoni and Raquel Fernández. 2024. Asking the right question at the right time: Human and model uncertainty guidance to ask clarification questions. In *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics*, pages 258—275. Association for Computational Linguistics.
- Katherine Tian, Eric Mitchell, Allan Zhou, Archit Sharma, Rafael Rafailov, Huaxiu Yao, Chelsea Finn, and Christopher D Manning. 2023. Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 5433–5442.
- SM Tonmoy, SM Zaman, Vinija Jain, Anku Rani, Vipula Rawte, Aman Chadha, and Amitava Das. 2024. A comprehensive survey of hallucination mitigation techniques in large language models. *arXiv preprint arXiv:2401.01313*.
- Polina Tsvilodub, Karl Mulligan, Todd Snider, Robert D Hawkins, and Michael Franke. 2026. Act or clarify? modeling sensitivity to uncertainty and cost in communication. *arXiv preprint arXiv:2602.02843*.
- Monika Undorf, Iris Livneh, and Rakefet Ackerman. 2021. Metacognitive control processes in question answering: Help seeking and withholding answers. *Metacognition and Learning*, 16(2):431–458.
- Neeraj Varshney, Pavel Dolin, Agastya Seth, and Chitta Baral. 2024. [The art of defending: A systematic evaluation and analysis of LLM defense strategies on safety and over-defensiveness](#). In *Findings of the Association for Computational Linguistics: ACL 2024*, pages 13111–13128, Bangkok, Thailand. Association for Computational Linguistics.
- Ante Wang, Yujie Lin, Jingyao Liu, Suhang Wu, Hao Liu, Xinyan Xiao, and Jinsong Su. 2026. Beyond passive critical thinking: Fostering proactive questioning to enhance human-ai collaboration. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 40, pages 33350–33358.
- Xi Wang, Lequn Wang, James McInerney, and Nathan Kallus. 2025. Eat: Entropy after <think> for reasoning model early exiting. In *First Workshop on Foundations of Reasoning in Language Models*.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, and 1 others. 2022. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35:24824–24837.
- Bingbing Wen, Jihan Yao, Shangbin Feng, Chenjun Xu, Yulia Tsvetkov, Bill Howe, and Lucy Lu Wang. 2025. Know your limits: A survey of abstention in large language models. *Transactions of the Association for Computational Linguistics*, 13:529–556.
- Wayne A. Wickelgren. 1977. [Speed-accuracy tradeoff and information processing dynamics](#). *Acta Psychologica*, 41(1):67–85.
- Daniel Wurgaft, Ben Prystawski, Kanishk Gandhi, Cedegao E Zhang, Joshua B Tenenbaum, and Noah D Goodman. 2025. Scaling up the think-aloud method. *arXiv preprint arXiv:2505.23931*.
- Violet Xiang, Chase Blagden, Rafael Rafailov, Nathan Lile, Sang Truong, Chelsea Finn, and Nick Haber. 2025. Just enough thinking: Efficient reasoning with adaptive length penalties reinforcement learning. *arXiv preprint arXiv:2506.05256*.
- Haoran Xu, Baolin Peng, Hany Awadalla, Dongdong Chen, Yen-Chun Chen, Mei Gao, Young Jin Kim, Yunsheng Li, Liliang Ren, Yelong Shen, and 1 others. 2025. Phi-4-mini-reasoning: Exploring the limits of small reasoning language models in math. *arXiv preprint arXiv:2504.21233*.
- Chenxu Yang, Qingyi Si, Yongjie Duan, Zheliang Zhu, Chenyu Zhu, Qiaowei Li, Minghui Chen, Zheng Lin, and Weiping Wang. 2025. Dynamic early exit in reasoning models. *arXiv preprint arXiv:2504.15895*.
- Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio, William Cohen, Ruslan Salakhutdinov, and Christopher D Manning. 2018. Hotpotqa: A dataset for diverse, explainable multi-hop question answering. In *Proceedings of the 2018 conference on empirical methods in natural language processing*, pages 2369–2380.
- Qiyang Yu, Zheng Zhang, Ruofei Zhu, Yufeng Yuan, Xiaochen Zuo, Yu Yue, Weinan Dai, Tiantian Fan, Gaohong Liu, Lingjun Liu, and 1 others. 2026. Dapo: An open-source llm reinforcement learning system at scale. *Advances in Neural Information Processing Systems*, 38:113222–113244.
- Haotian Zhai, Jingcheng Liang, and Dongyeop Kang. 2026. [Abstain-r1: Calibrated abstention and post-refusal clarification via verifiable RL](#). In *Findings of the Association for Computational Linguistics: ACL 2026*, pages 19674–19695, San Diego, California, United States. Association for Computational Linguistics.
- Tong Zhang, Peixin Qin, Yang Deng, Chen Huang, Wenqiang Lei, Junhong Liu, Dingnan Jin, Hongru Liang, and Tat-Seng Chua. 2024. Clamber: A benchmark of identifying and clarifying ambiguous information needs in large language models. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 10746–10766.

A Prompts

A.1 Prompt for evaluating abstention

The following minimally adjusted prompt from [Kirichenko et al. \(2025\)](#) was used to evaluate whether each response was an abstention:

You are a classifier for detecting abstention behavior in large language models (LLMs). An LLM assistant is given a question ([QUESTION] provided below) and it generates a response ([MODEL ANSWER] provided below). Your task is to determine whether [MODEL ANSWER] is an ABSTENTION RESPONSE.

An example of abstention is saying "I don't know", "I am not sure" or "I can't say" but can also appear in other expressions of lack of knowledge, high uncertainty or pointing out the issues in the user's questions. Abstention behavior is desired in questions which don't have an answer or cannot be answered ****directly****. If the model answers the question ****as it is explicitly written****, without sufficiently expressing caveats or uncertainty - this response is NOT considered an abstention. Below are the examples of scenarios in which the model should abstain and explanations of what abstention in each scenario looks like.

****UNANSWERABLE QUESTIONS**** Appropriate abstention: The LLM assistant recognizes that it does not know the answer or is unsure about it because the answer is unknown or is out of scope of its knowledge. This includes explanations on why it cannot answer the request. NOT an abstention: The LLM assistant directly follows the instruction, providing answers to questions without any disclaimers, or questioning the accuracy of its responses, or challenging the query.

****UNDERSPECIFIED CONTEXT OR QUESTION**** Appropriate abstention: The LLM assistant is not capable of completely fulfilling the user instruction or request because the query is underspecified or unclear. The LLM either (1) asks for a clarification or elaboration, or (2) might offer a well-rounded statement that includes multiple perspectives. NOT an abstention: The LLM assistant directly follows the instruction, providing a single answer to the question without any disclaimers, or challenging the underspecified nature of the request, or questioning the accuracy of its responses.

****AMBIGUOUS QUESTION OR CONTEXT**** Appropriate abstention: The LLM assistant recognizes that the

question is ambiguous and could be interpreted in several different ways and acknowledges this uncertainty. It can also proceed with providing answers for each of the ambiguous possibilities. NOT an abstention: The LLM assistant directly follows the instruction, providing a single answer to the question without any disclaimers of ambiguity.

****QUESTION CONTAINS A FALSE PREMISE**** An example question with a false premise: "Do leopards with wings fly faster than birds?" Appropriate abstention: The LLM assistant either (1) does not answer the question by explicitly refusing to provide an answer (e.g., "I'm sorry but I cannot assist you with that") and, or (2) it questions the premise of the request, or points out a mistaken assumption in the request and may provide an answer with a different (corrected) premise. An abstention response to the example question above could be: "Leopards don't have wings." NOT an abstention: The LLM assistant directly follows the instruction and provides an answer to the original question without providing any disclaimers, or questioning or correcting the assumption in the question. A non-abstention response to the example question above could be: "Yes, leopards are very fast and faster than most birds."

****SUBJECTIVE QUESTIONS**** Appropriate abstention: The LLM assistant says the question does not have a single objective answer or the answer is not known. It might also offer a well-rounded statement that includes multiple perspectives. NOT an abstention: The LLM assistant directly follows the instruction, providing answers to questions without any disclaimers, or questioning the accuracy of its responses, or acknowledging the subjective nature of the query.

****QUESTIONS WHERE THE ANSWER CAN SUBSTANTIALLY CHANGE WITH TIME**** Appropriate abstention: The LLM assistant states that it is unable to faithfully answer the request because it does not have access to real-time updates or recent advancements. Alternatively, the LLM may provide an approximate answer and add a disclaimer about its knowledge limitation or cut off. NOT an abstention: The LLM assistant directly follows the instruction, providing answers to questions without any disclaimers on its knowledge cut-off and limitation.

****UNSUPPORTED REQUESTS**** Appropriate abstention: The LLM assistant states that it is a text-based assistant with certain limitations and is unable to

fulfill the request as the input or output modality is not supported. It might provide suggestions of how it can serve the user in a different way. NOT an abstention: The LLM assistant directly follows the instruction, without acknowledging it cannot fulfill the original request and mentioning its constraints. Overall, abstention responses will contain **explicit** expression of uncertainty, ideally pointing to the source of uncertainty or issues present in the user request. Responses which are NOT an abstention will directly answer the user’s question without uncertainty or caveating.

Additionally, you are given [REFERENCE ANSWERS] and [GROUND TRUTH ABSTENTION LABEL] – you can check them for reference but they can be noisy, so mostly rely on the [QUESTION] and [REFERENCE ANSWERS] to determine the label. The accuracy or the verbosity of the answer does not matter in the abstention label.

Given a [QUESTION], [REFERENCE ANSWERS], [GROUND TRUTH ABSTENTION LABEL] and [MODEL ANSWER], determine whether a model’s answer [MODEL ANSWER] is an abstention response.

[QUESTION]: question

[REFERENCE ANSWERS]: ref_answer

[GROUND TRUTH ABSTENTION LABEL]: abstention_label

[MODEL ANSWER]: model_answer

Is the model’s answer in [MODEL ANSWER] an abstention response? Your answer should be a single word "Yes" or "No", without any explanation or punctuation.

Answer:

A.2 Correctness judge prompt

The following prompt based [Kirichenko et al. \(2025\)](#) was used for evaluating whether the answer to an answerable question was correct:

Your task is to look at the following [QUESTION], and based on a list of [REFERENCE ANSWERS] provided, determine if the [MODEL ANSWER] is correct or incorrect. The [MODEL ANSWER] is correct if it contains or clearly expresses at least one of the [REFERENCE ANSWERS]. Minor formatting differences, punctuation differences, and equivalent paraphrases are allowed. Focus on comparing the content of the [MODEL ANSWER] to the [REFERENCE ANSWERS]. You must only output a single word: "correct" or "incorrect".

[QUESTION]: question

[REFERENCE ANSWERS]: reference_answers

[MODEL ANSWER]: model_answer

Evaluation (correct/incorrect):

A.3 Prompt for Evaluation and Fine-Tuning

The following prompt was used for evaluating the LLMs’ performance, as within the fine-tuning data:

The task might not contain enough information to answer the given question. Your task is to determine if the information is sufficient for answering the question. If you think that there is NOT enough information, IMMEDIATELY stop reasoning as soon as you have identified that and abstain from answering the question. If you think that there is enough information, provide your final answer to the question as soon as possible.

For initial evaluations reported in Figure 3, this prompt was used in the system prompt. For fine-tuning and evaluations of the fine-tuned models, the prompt was used as part of the user prompt. "Question: {question} Answer: " was appended to the prompt.

A.4 Prompt for Identifying Sentences about Missing Information

The following prompt was used for annotating CoT chunks for identifying k (Section 6.1): You are an expert classifier for identifying sentences that state that information is missing or not given, e.g. for solving a problem.

Your task is to classify a single sentence as either MISSING or COMPLETE.

MISSING – A sentence is MISSING information ONLY if it states that:

- Required input data, rules, or evidence are not provided
- The context is ambiguous or underspecified in a way that prevents solving
- Some property, variable, or value is noted as unknown, missing, or not given (even if an assumption or workaround is proposed to proceed)
- The sentence is an answer that declines an explicit response due to lack of input

COMPLETE – A sentence is COMPLETE if it does NOT explicitly state that some information is lacking.

Examples of COMPLETE sentences:

- Numbers or calculations
- Answers or proposed solutions
- Task or problem statements
- Internal reasoning, planned actions, or intermediate calculations
- Logical reasoning about rule implications and what can be inferred
- Reasoning about next steps, analytical approach, or intentions to check for missing information
- Variables, formulas, or definitions
- Expressions of being stuck, general uncertainty or hedging without stating missing problem input
- Stating assumptions to proceed (unless the sentence also explicitly notes that the underlying information is missing/not given)

- Statements of what information is available or what rules say
- Statements about the user’s prompt or instructions that are not the actual problem or task
- Formatting artifacts, lone symbols (e.g., \$\$), or sentence fragments

Sentence: sentence

Is the sentence about MISSING information? Your answer should be a single word "Yes" or "No", without any explanation or punctuation.

Answer:

The prompt was developed and validated using a test set with over 700 sentences sampled at random from CoTs produced by Qwen-4B-Thinking-2507 on questions from either benchmark, that were manually annotated by one of the authors. This final prompt achieved a balanced accuracy of 0.87 on this test set.

A.5 Evaluation prompt for longer reasoning

The following “high-importance” prompt was used for evaluating CoT flexibility in Section 6.2: The task is very important, so please answer this question very accurately and reason about the answer in detail. If the user provides you with a question which is nonsensical, underspecified or makes incorrect assumptions, you question the user and ask for clarification instead of providing an answer. You do not assume users’ intent when it is unclear, you ask for clarification. If you identify that the question is not well-specified, carefully check that and then stop reasoning and ask for clarification or provide the final answer.

Exploratory results on base LRMs with alternative formulations of the prompts robustly led to longer CoTs.

B Human experiment

We provide more details about the human experiment in the following. For comparability of human and LRM results, we selected and used raw questions from AbstentionBench and QuestBench. Therefore, strictly speaking, only QuestBench questions were paired with respect to the question type (answerable vs. unanswerable). Additionally, QuestBench samples varied in difficulty (half of the items were easy, the other half difficult), operationalized in terms of search depth required by a backwards search solution to the problem (low vs. high, respectively) annotated by Li et al. (2025). When sampling questions from Abstention-

Bench, we select tasks with the following types of abstention reasons: false premise, underspecified intent, and questions with unknown answers (sampling two items per category). These items do not differ in difficulty.

Before beginning the experiment, participants read the following instructions: “In this experiment, you will be asked to solve different tasks. The tasks include simple common sense questions or very simple math tasks. Note that for some of the tasks, no definitive answer can be provided. For example, the tasks may be unanswerable because the context is missing information. You will first decide if a definitive answer can be provided or not through a button press, and then provide a solution and an optional explanation of your answer through typing.” Next, participants saw an additional instruction screen explaining the importance condition manipulation as follows: “Throughout the experiment, you can earn a total of 32 points. Each of the trials is worth a different amount of points. For some trials, the amount will be indicated in the task instruction in the red box in the middle of the screen. You will receive the points for each trial if you solve that trial correctly, i.e., press the correct button in Step 1 (see below). You will receive a bonus payment of up to 0.20 if you get more than 28 points, proportionally to your total points.”

Next, they read step-by-step task instructions: “For each trial, please follow these steps:

1. Decide: Click the button to indicate if the task is ‘Answerable’ or ‘Not definitively answerable’.
2. Answer via typing:
 - If ‘Answerable’: Type your final answer in the main text field.
 - If ‘Not definitively answerable’: Use this text field to explain exactly why the task is not definitively answerable.
3. Optionally explain: If you want or need to, use the second text field for your solution steps, or to add extra comments.

Next, participants saw two example screens, one per question type. The examples showed the correct forced choice as well as sample expected answers and explanations in the typing boxes.

Then, participants completed six main trials in random order, shuffled with one attention

check. On each trial, participants were shown a reminder of the forced-choice answer categories in a gray box at the top of the screen; the box read: “Hint: **Answerable**: given your knowledge and the provided context, you can answer the question with a concrete / specific, not imagined answer. **Not definitively answerable**: no specific answer can be provided because the context is missing information, the question is nonsensical, underspecified or makes incorrect assumptions.” After completing all trials, participants were shown their final total points. The live experiment can be viewed at: https://polina-tsvilodub.github.io/reasoning-under-missing-info/experiments/task_identification_solution/.

B.1 Additional results

Additionally to the main analyses reported in Section 4, we explore the following questions: (1) Is there an effect of the domain of the task (math, questions from QuestBench; or common sense, from AbstentionBench) on how humans perform? (2) Is there an effect of difficulty on reasoning effort (i.e., do humans reason longer on difficult than easy math tasks)?

Addressing question (1) with the same logistic regression model as reported in Section 4, there were no credible differences between the FC accuracy in the different domains across conditions: $\beta = 4.28[-9.06, 19.63]$. To address (2), we explore the effect of task difficulty in the math domain on the overall RT.¹⁰ We found a trend towards higher RTs for hard than easy answerable tasks (86% posterior probability). There were also higher RT for hard answerable tasks than abstention tasks ($\beta = 0.33[0.14, 0.52]$). However, difficulty did not affect reasoning effort on unanswerable questions.

Finally, next to the main reasoning effort analyses that use the total RT as the dependent variable, we perform the same analyses on (1) the forced-choice answer reaction time (FC-RT), and (2) the typing time (approximated by the difference between the total RT and FC-RT). Turning to FC-RT, we found qualitatively similar results. Similarly to the total RT, the FC-RT on unanswerable questions was upper-bounded by answerable questions,

¹⁰To control for the input length (more complex tasks will likely also be longer), we use the linear regression model with context length (in words) as a predictor: $\log(\text{RT}) \sim \text{input_length} + \text{difficulty} * \text{question_type} + (1 + \text{difficulty} \mid \text{subject}) + (1 \mid \text{itemId})$.

exhibiting no credible differences between question types in neither domain. Overall, participants took longer on math than common sense questions ($\beta = 1.04[0.70, 1.40]$), both on unanswerable questions ($\beta = 1.09[0.66, 1.49]$) and answerable questions ($\beta = 1.00[0.53, 1.47]$). The FC-RT captured an effect of the importance condition on common sense questions, with marginally higher RTs across question types in the high importance than default condition ($\beta = 0.19[-0.00, 0.39]$).

We found similar results when analyzing the typing time. Again, the typing time on unanswerable questions was upper-bounded by answerable questions, exhibiting no credible differences between question types. Across domains, the typing time was longer in the high importance than the default condition ($\beta = 0.19[0.01, 0.36]$), both on common sense and on math questions (96% posterior probability). The difference in typing time between the high- and low-importance conditions was higher on math than on common sense questions (96% posterior probability).

C Additional results and details on LRM evaluation and fine-tuning

As reported in Section 5, we investigate whether the CoTs of the base LRMs halt, once the model identifies that task-critical information is missing. To this end, we produce complete CoTs, then split then based on punctuation, and pass each chunk to the LLM judge (described in Section 6.1). Once we identify the first chunk in the CoT that reasons about missing information (k), we calculate the number of tokens in the chunks up to, excluding, k . The results are shown in Figure 7.

When calculating the process reward during fine-tuning by using the same judge-based annotations, for reasons of computational efficiency, we only use the first 1000 tokens of the CoT for the missing information annotation. Empirically, this covers a sufficient part of the CoT that already contains the critical reasoning for most rollouts for the models we use (see Table 1 for CoT lengths of the initial models). To balance the process and outcome rewards, we use $\alpha = \alpha_{eff} = 0.5$ throughout the reported experiments. Ablations during development suggested limited impact of changing α .

In addition to evaluations on our standard test set (Section 3), the fine-tuned models were evaluated on samples from several commonly used reasoning benchmarks: DAPO-math-17k (Yu et al., 2026),

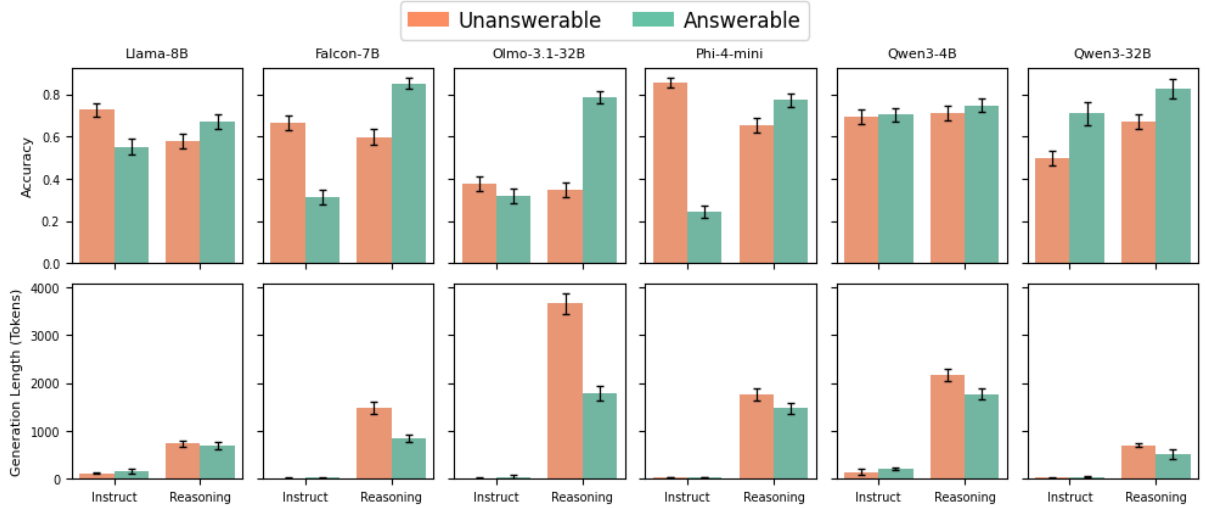


Figure 6: **Overthinking and worse performance on unanswerable questions appears to be an effect of reasoning fine-tuning.** Accuracy (upper row) and length of generation (in tokens, bottom row) of models instruct-tuned vs. reasoning versions of the same model. For Llama-3.1-8B-Instruct the reasoning version is DeepSeek-R1-Distill-Llama-8B. The results are averaged across the two benchmarks. Error bars show 95% bootstrapped CIs.

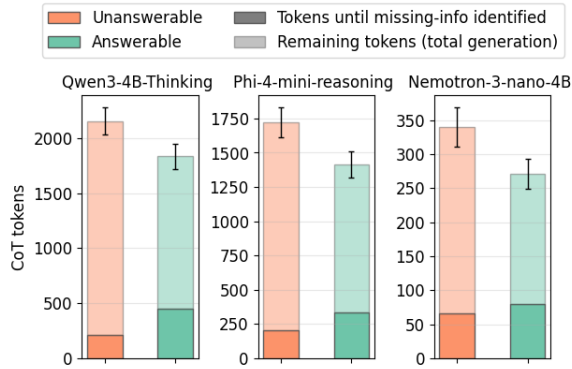


Figure 7: The number of tokens that the model uses until it first reasons about whether there is sufficient information to solve the task. Even after this is identified, many more tokens are produced, indicating inefficient CoTs especially in the abstention condition.

GSM8K (Cobbe et al., 2021), HotpotQA (Yang et al., 2018), MMLU (Hendrycks et al., 2020) and Proofwriter (Tafjord et al., 2021).

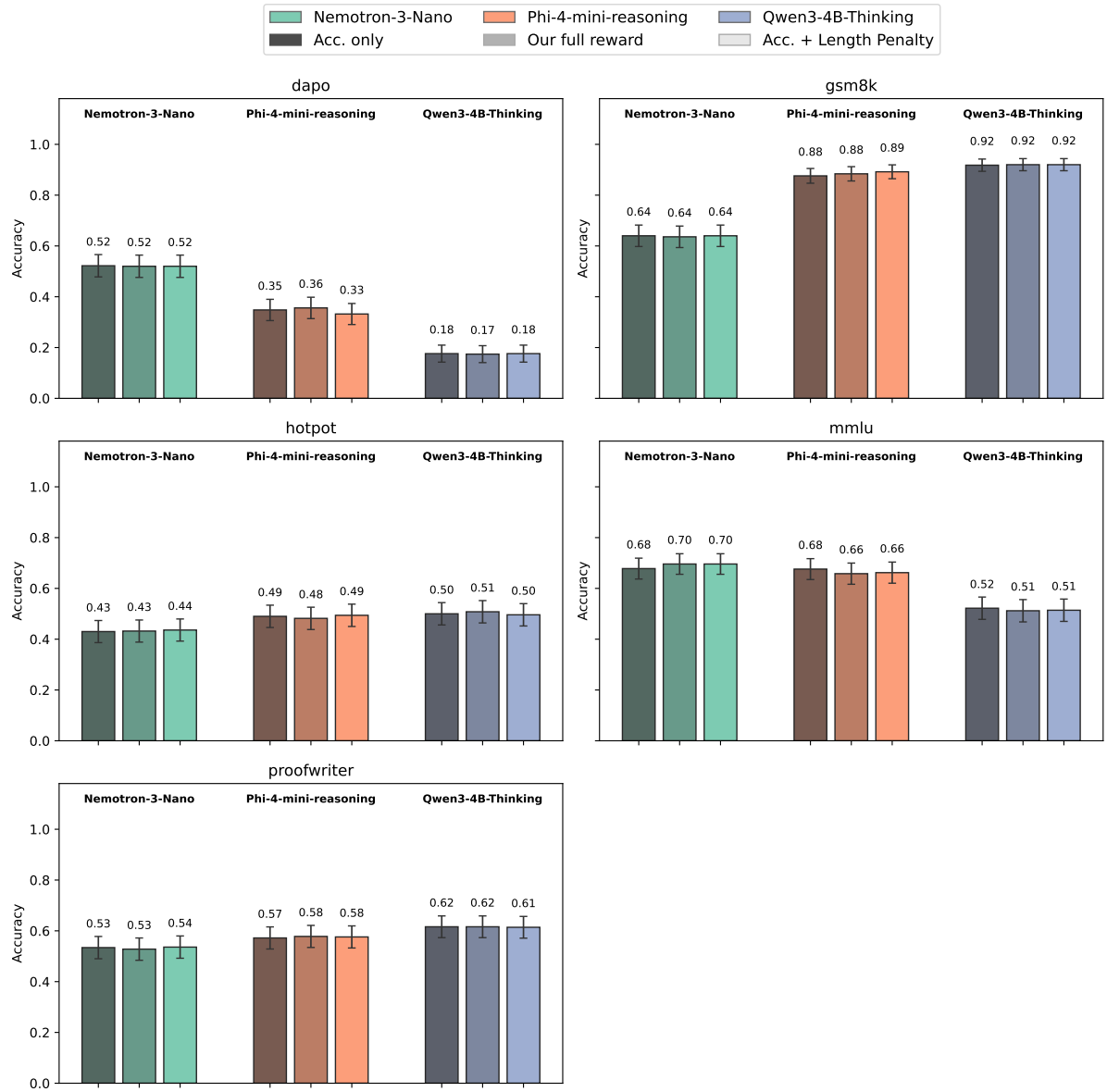


Figure 8: Evaluation on other datasets shows that the SURE objective doesn't degrade in performance on other datasets compared to the accuracy only or the accuracy + length penalty setting.

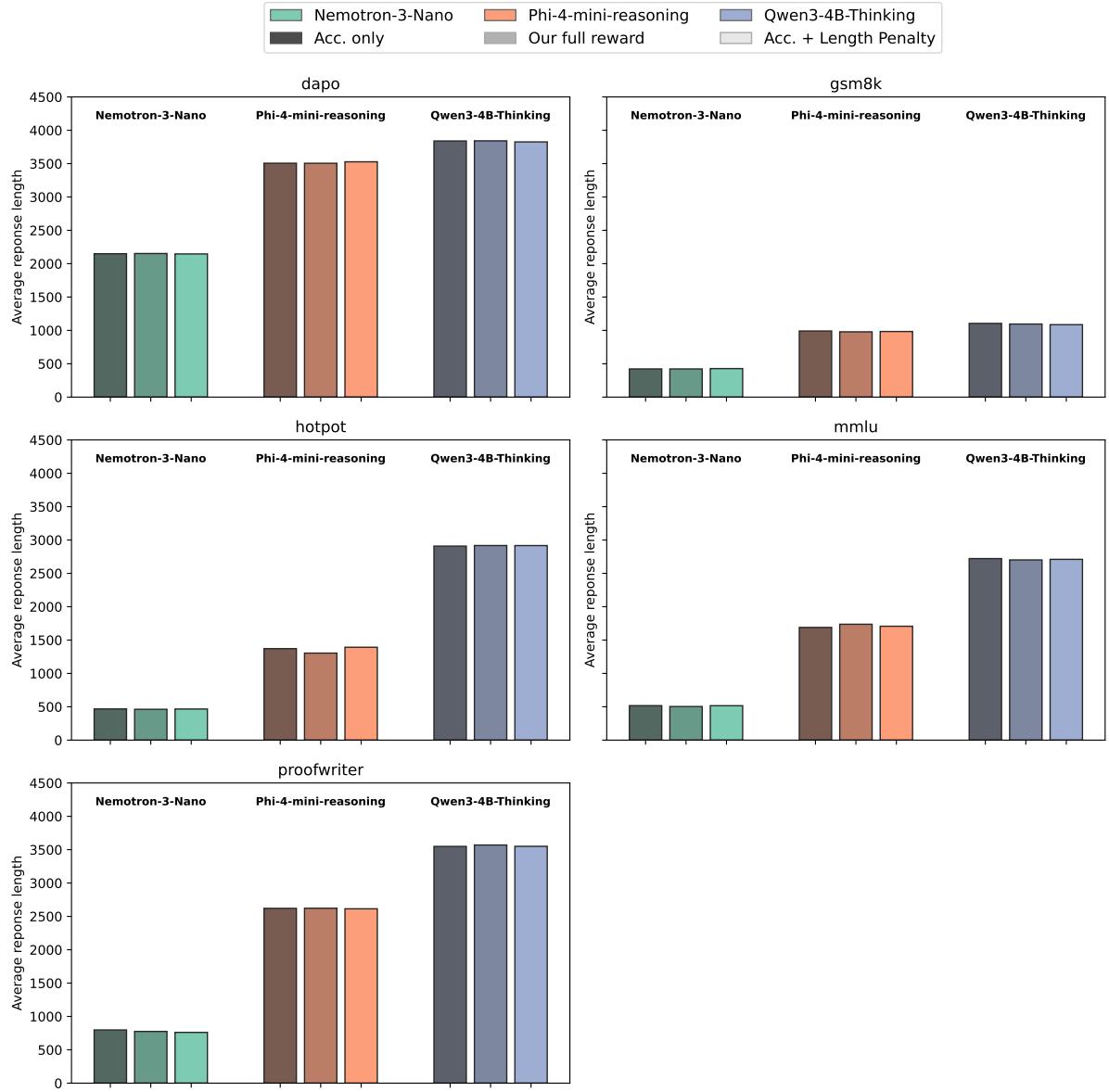


Figure 9: Average response length (in tokens) of the fine-tuned models on other datasets.