

---

# RL on Incorrect Synthetic Data Scales the Efficiency of LLM Math Reasoning by Eight-Fold

---

Amrith Setlur<sup>\*,1</sup> Saurabh Garg<sup>1</sup> Xinyang (Young) Geng<sup>2</sup> Naman Garg<sup>3</sup>  
Virginia Smith<sup>1</sup> Aviral Kumar<sup>1,2</sup>

<sup>1</sup>Carnegie Mellon University <sup>2</sup>Google DeepMind <sup>3</sup>MultiOn

## Abstract

Training on model-generated synthetic data is a promising approach for finetuning LLMs, but it remains unclear when it helps or hurts. In this paper, we investigate this question for math reasoning via an empirical study, followed by building a conceptual understanding of our observations. First, we find that while the typical approach of finetuning a model on synthetic correct or *positive* problem-solution pairs generated by capable models offers modest performance gains, sampling more correct solutions from the finetuned learner itself followed by subsequent fine-tuning on this self-generated data **doubles** the efficiency of the same synthetic problems. At the same time, training on model-generated positives can amplify various spurious correlations, resulting in flat or even inverse scaling trends as the amount of data increases. Surprisingly, we find that several of these issues can be addressed if we also utilize *negative* responses, *i.e.*, model-generated responses that are deemed incorrect by a final answer verifier. Crucially, these negatives must be constructed such that the training can appropriately recover the utility or advantage of each intermediate step in the negative response. With this *per-step* scheme, we are able to attain consistent gains over only positive data, attaining performance similar to amplifying the amount of synthetic data by **8×**. We show that training on per-step negatives can help to unlearn spurious correlations in the positive data, and is equivalent to advantage-weighted reinforcement learning (RL), implying that it inherits robustness benefits of RL over imitating positive data alone.

## 1 Introduction

Training large language models (LLMs) relies on the ability to train on large amounts of high-quality data. It is predicted that we will run out of high-quality internet data by 2026 [32, 58], necessitating training on model-generated data, or what is commonly referred to as *synthetic data*. Recent trends illustrate that scaling up synthetic data can lead to improvements [8, 29] on hard reasoning problems, while other results illustrate that training on synthetic data can steer the performance of the model into a downward spiral [3, 17, 51]—amplifying biases, misinformation, and undesired stylistic properties. Thus while *in principle*, synthetic data could potentially address data scarcity, it must be designed in an appropriate manner to be effective. However, this has been hard due to a lack of an understanding of how synthetic data contributes to LLM behavior.

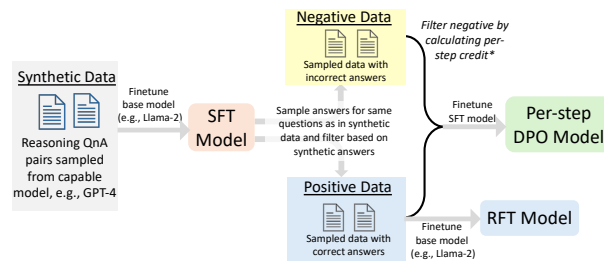
To provide clarity on how synthetic data contributes to performance, we aim to understand its impact on LLM capabilities via a study on math reasoning, a prevalent scenario where synthetic data is used. Typically, in this setting, synthetic data corresponds to correct or *positive* model-generated responses for a novel set of initial problems synthesized by prompting capable models [29, 31]. The resulting model is then evaluated on a held-out set of problems drawn from a test set. Perhaps as expected, we find that performance improves when finetuning models on positive synthetic responses, though the scaling rates for performance improvement are often substantially slower than those observed during pretraining. Concretely, we find that under the scaling law of Zhang et al. [71], the error rate scales as  $\approx D^{-0.05}$  to  $D^{-0.15}$  in the size  $D$  of synthetic dataset. Second, we observe that not all types

---

\*Corresponding author. Please send correspondences to asetlur@andrew.cmu.edu.

of positive synthetic data are equally effective: often positive responses self-generated by the learner itself are as effective as  $2\times$  synthetic data from bigger models in improving performance. This is because responses from a similar model are “easier-to-fit” than those from a more capable model, resulting in reduced memorization [26, 56] during finetuning. We also observe that if the positive response contains incorrect/irrelevant intermediate steps, training on such data often incentivizes the model to overfit on spurious correlations, leading to a flat or even inverse scaling with more data.

Perhaps surprisingly, we find that the aforementioned pathologies of training on positive data only can be addressed if we also utilize synthetic *negative* responses: responses generated by the model that do not result in obtaining a correct final answer. One way to utilize negative responses is via methods such as direct preference optimization (DPO) [41]. While performance of standard DPO [41] largely flatlines as the number of synthetic problems are scaled up (Figure 5), we are able to attain consistent improvements if the negative data is generated appropriately. A solution trace for a math problem typically comprises of multiple *reasoning steps* corresponding to intermediate results. Our insight is that instead of contrasting arbitrary correct and incorrect responses, we should contrast those positive and negative responses that depict good and bad choices for the more “critical” intermediate steps: steps that the model must carefully produce so as to succeed at the problem. In other words, critical steps are those which the model is unable to recover from, and hence, must be emphasized. With this scheme, we are able to attain consistent gains over only positive data, **attaining performance similar to scaling up positive synthetic data by  $8\times$** . We show that training on this sort of negative data evades spurious steps amplified by training on positive data alone.



**Figure 1: Positive and negative synthetic data:** Pictorial representation of positive/negative synthetic data definitions we use and how they are fed to SFT, RFT and DPO.

To theoretically understand our findings, we build a conceptual model of how training on this data benefits performance. Formally, we show that this construction of negative data, which emphasizes “critical” tokens (Figure 6) enables us to perform credit assignment, and is equivalent to training the model with per-step advantage-weighted reinforcement learning (RL) [40] on a mixture of positive and negative synthetic data. Specifically, these advantage values are computed under an optimal value function induced by sampling multiple responses under the SFT policy obtained by training on only the positive data. This reduction of using negative data to advantage-weighted RL enables us to conceptually compare it to training on positive data, which corresponds to imitation learning (*i.e.*, behavioral cloning) on positive data. First, we are able to argue for the generalization gains of advantage-weighted RL through the lens of distribution robust objectives. Second, building on theoretical results in RL [27], we are also able to show that when advantages can be estimated reliably, advantage-weighted RL will be significantly more sample-efficient compared to imitation. Overall, this model explains the utility of negative data over only positive data.

Our contribution is a study of the role of synthetic data in improving math reasoning capabilities of LLMs. We derive scaling laws for positive and negative data on common reasoning benchmarks and observe that: **(a)** training on positive synthetic data from capable models results in scaling rates that are significantly slower than standard empirical risk minimization; **(b)** training on model-generated positive synthetic data can improve sample efficiency by  $2\times$  but also amplifies spurious correlations; **(c)** appropriate ways of constructing learner-specific negative data with emphasis on critical steps, results in a performance boost equivalent to scaling up positive data  $8\times$ ; **(d)** training with negative data provides a mechanism to unlearn spurious correlations; and **(e)** we present a conceptual model inspired from RL to explain our observations on synthetic data and the generalization benefits we see.

## 2 Related Work

A standard procedure to finetune a pretrained LLM is teacher-forcing on expert data, *i.e.*, maximizing the likelihood of the next token given all previous tokens [7, 61]. In Appendix G we discuss some failure modes of this procedure for math reasoning that positive or negative synthetic data can address.

**Positive synthetic data.** Learning theory dictates that the SFT policy trained on more SFT data (*e.g.*, 1.5M for DeepSeek-Math [5]) would have improved math reasoning capabilities. Thus, a

common goal for generating synthetic data as close as possible to the SFT data [29, 31, 32]. That said, generating high quality math data can be challenging, since verification can often be hard. When synthetic data is verified by larger models [50, 59], recent works [33, 66] observe scaling similar to finetuning LLMs on expert data [69, 71], while another work [14] notes the compositional gains from SFT data for code generation. Common sources of “good” synthetic data include responses from stronger teachers [29, 30], or data generated by the SFT policy itself, in the framework of reinforced self-training (ReST) and STaR [8, 52, 69, 70]. In our work, we study and compare the performance scaling with positive synthetic data from bigger models like GPT-4 and Gemini 1.5 Pro with self-generated positive data. We connect our findings to evidence showing “ease of learning” generalizable features on self-generated completions [26] which often prevents undesirable memorization [56]. Finally, our work also sheds light on several concerns about training on synthetic positive data amplifying biases [48, 63], and leading to model collapse [13, 17], especially due to overfitting on “spurious” intermediate steps. We conceptually explain this phenomenon and also discuss how negative model-generated responses can help identify and unlearn those spurious steps.

**Benefits and nuances of negative synthetic data.** While most works on synthetic data [29, 32, 66, 69] train only on correct answers, our work also studies complementary gains from incorrect completions generated by the SFT policy [23, 38, 39, 68]. To leverage sub-optimal negative data, we adopt the framework of offline preference optimization [16, 41, 73], where a preference pair is constructed using correct and incorrect responses for the same problem [38]. Despite numerous studies on preference data composition [8–10, 37, 54, 55, 60], it’s unclear how to pose a reasoning problem as a preference optimization problem. Randomly pairing correct and incorrect completions in a preference pair can lead to poor performance [21, 38, 39, 64] due to objective mismatch [55, 72] and requires auxiliary losses to perform well. Another option is to use negative data for training verifiers [22, 65] but this line of work still only trains the policy using positive data. We introduce a conceptual model of negative data, where we understand how certain choices of negative data can assign per-step credits, which we use to establish the equivalence of preference optimization to advantage weighted RL. Self-explore method in Hwang et al. [23] can be viewed as an special instance of our general framework. Other works [34, 59] exploit per-step credit assignment through tree-based sampling. They identify the reasoning subsequence that led to the most incorrect answers under the SFT policy for training a reward model. While this is related, our conceptual model and analysis also understands why assigning per-step credits can generalize better by unlearning spurious correlations, *e.g.*, when the credits are given by the Q-function of the “best-of-K” SFT policy.

### 3 Problem Setup and Synthetic Data Generation Pipeline

Building on the recipe of Li et al. [29], Liu et al. [31], we use GSM8K [11] and MATH [19] to collect synthetic data consisting of both novel problems designed by capable models such as GPT4 [1] and Gemini 1.5 Pro [44], and responses to these problems, obtained from the same models.

**Synthetic data pipeline.** First, given a dataset  $\mathcal{D}_{\text{real}} = \{(\mathbf{x}_i^r, \mathbf{y}_i^r)\}$  of problems  $\mathbf{x}_i^r \sim p_{\text{real}}(\mathbf{x})$  and solution traces  $\mathbf{y}_i^r \sim p_{\text{real}}(\mathbf{y} \mid \mathbf{x}_i)$ , we prompt one of the highly-capable models with a uniformly random sample  $(\mathbf{x}_i^r, \mathbf{y}_i^r) \in \mathcal{D}_{\text{real}}$  and ask the model to generate a new problem  $\mathbf{x}_i$  such that it is similar to the real problem  $\mathbf{x}_i^r$ , in a way that a feasible solution exists. Second, we ask the model to provide a solution trace answer  $\mathbf{y}_i$  with step-by-step reasoning (exact prompts for  $\mathbf{x}_i, \mathbf{y}_i$  are borrowed from Li et al. [29], shown in Appendix H). We assume that the answers generated via this process are accurate, and perform lightweight filtering step to remove duplicates, badly-formatted answer traces, and model failures. Based on the above, for any synthetic problem and solution pair  $(\mathbf{x}, \mathbf{y})$ , we can define a binary reward function  $r(\mathbf{y}, \hat{\mathbf{y}}) \mapsto \{0, 1\}$ , which verifies if a new solution trace  $\hat{\mathbf{y}}$  is correct or not. This is implemented with a set of answer extraction and string matching tools borrowed from [29, 66]. We say that a new trace  $\hat{\mathbf{y}}$  is a *positive* trace if it produces the correct final answer *i.e.*,  $r(\hat{\mathbf{y}}, \mathbf{y}) = 1$ , and *negative* if it produces an incorrect final answer, *i.e.*,  $r(\hat{\mathbf{y}}, \mathbf{y}) = 0$ . By definition,  $r(\mathbf{y}, \mathbf{y}) = 1$ , and the original trace  $\mathbf{y}$  is always positive.

**Positive and negative datasets.** The above process induces a joint distribution  $p_{\text{syn}}(\mathbf{x}, \mathbf{y})$ , *iid* samples from which yields positive synthetic dataset  $\mathcal{D}_{\text{syn}}$ . We note that the sampling process for  $\mathcal{D}_{\text{syn}}$  is designed to ensure that the induced marginal distribution over synthetic problems  $p_{\text{syn}}(\mathbf{x})$  is close to  $p_{\text{real}}(\mathbf{x})$ . We will use  $\mathcal{D}_{\pi}^{+}$  to denote the positive dataset of  $(\mathbf{x}, +\hat{\mathbf{y}})$  where  $+\hat{\mathbf{y}}$  is a positive solution trace generated from some policy  $\pi(\cdot \mid \mathbf{x})$ . For a positive  $+\hat{\mathbf{y}}$  and negative  $-\hat{\mathbf{y}}$  trace, sampled from the same policy  $\pi(\cdot \mid \mathbf{x})$ , we denote a dataset over problems and solution pairs:  $(\mathbf{x}, +\hat{\mathbf{y}}, -\hat{\mathbf{y}})$  as  $\mathcal{D}_{\pi}^{\pm}$ .

**Reasoning steps.** The trace  $y_i$  consists of several intermediate steps,  $y_i = [y_{i,1}, \dots, y_{i,L}]$ . We assume each trace has at most  $L$  steps, and use  $y_{1:t}$  to denote the subsequence of first  $t$  steps. Since mathematical reasoning problems require step-by-step computation, simply arriving at an incorrect final answer does not mean that all steps in a negative  $\hat{y}$  are incorrect. Similarly, a positive  $\hat{y}$  may also have incorrect reasoning steps. In fact, even the original answers generated by more capable models in  $\mathcal{D}_{\text{syn}}$  may also contain incorrect reasoning steps, and training on such traces may actually lead to unintended consequences (Section 5).

## 4 Learning from Synthetic Data

In this section, we discuss various algorithms for learning from the synthetic dataset  $\mathcal{D}_{\text{syn}}$  discussed in the previous section, as well as positive and negative solution traces generated using a model.

**Supervised and rejection finetuning (SFT and RFT).** Given positive synthetic  $\mathcal{D}_{\text{syn}}$ , perhaps the most straightforward approach (and the most prevalent) is to learn  $\pi_{\text{sft}}$  on this data via supervised next-token prediction:  $\pi_{\text{sft}}(\cdot | x) := \arg \max_{\pi} \mathbb{E}_{x, y \sim \mathcal{D}_{\text{syn}}} [\log \pi(y | x)]$ . Another option is to train via supervised next-token prediction on problems in  $\mathcal{D}_{\text{syn}}$ , but when using a positive solution trace  $\hat{y}$  sampled from  $\pi_{\text{sft}}(\cdot | x)$ , instead of positive synthetic responses from the capable models in  $\mathcal{D}_{\text{syn}}$ . Akin to rejection finetuning (RFT [69] or STaR [70]), sampling from  $\pi_{\text{sft}}(\cdot | x)$  once is not guaranteed to give a positive response, and we instead sample  $M$  times for each  $x$  and construct the dataset  $\mathcal{D}_{\pi_{\text{sft}}}^+$  of SFT policy generated positive responses. Then, we apply the next-token prediction loss on  $\mathcal{D}_{\pi_{\text{sft}}}^+$ .

**Preference optimization.** Beyond positive data, we can also learn from negative synthetic data generated from the SFT policy, especially when contrasted with positive responses. However, learning from negative data presents multiple open design questions pertaining to the construction of negative traces, and the choice of the loss function, and simple supervised fine-tuning will not be a good choice since it will incentivize the model to produce more errors. Therefore, we use a contrastive training approach, direct preference optimization (DPO [41]) for incorporating negative data from  $\pi_{\text{sft}}$ . In a nutshell, DPO trains a policy using the following preference optimization objective:

$$\min_{\pi} \mathcal{L}_{\text{DPO}}(\pi) := \mathbb{E}_{(x, +y, -y) \sim \mathcal{D}_{\pi_{\text{sft}}}^{\pm}} \left[ \sigma \left( \beta \log \frac{\pi(+y | x)}{\pi_{\text{sft}}(+y | x)} - \beta \log \frac{\pi(-y | x)}{\pi_{\text{sft}}(-y | x)} \right) \right]. \quad (1)$$

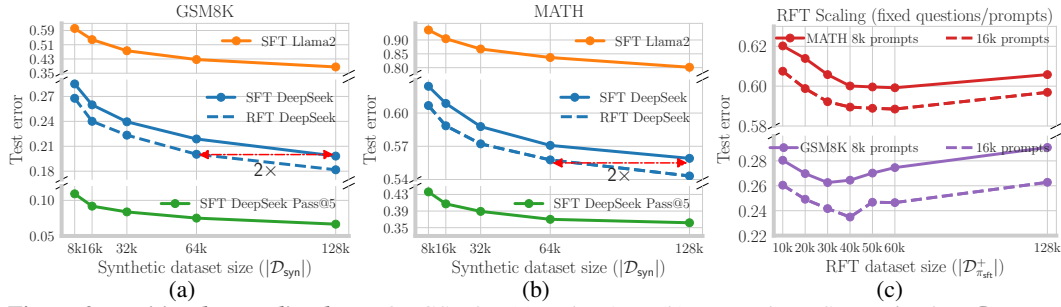
We consider two objectives that construct negative data and subsequently optimize Equation 1. The first variant is **standard DPO** [41], which samples negative data  $-y$  from the  $\pi_{\text{sft}}$  (with rejection sampling) and adds  $(x, y, -y)$  to  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm}$ . The second variant is **per-step DPO** [23], which first samples a complete solution trace  $\hat{y}_{1:L}$  from  $\pi_{\text{sft}}$  and then determines the “first pit”  $\hat{y}_c$ . The first pit  $\hat{y}_c$  is the step where any completion following the step:  $\hat{y}_{c+1:L} \sim \pi_{\text{sft}}(\cdot | x, \hat{y}_{1:c})$  leads to incorrect answers in expectation under  $\pi_{\text{sft}}$ . The triplet  $(x, y, \hat{y}_{1:c})$  is added to the preference dataset  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm}$ .

## 5 Positive Data Improves Coverage, But Amplifies Spurious Correlations

We first analyze the influence of scaling up positive synthetic data on GSM8K and MATH. In this experiment, we fine-tune DeepSeek-Math-7B [5] and LLaMA2-7B [57] models (details in Appendix J) on varying sizes of  $\mathcal{D}_{\text{syn}}$ , constructed out of a 5:1 mixture of GPT-4-turbo [1] and Gemini-1.5 Pro [44]. We obtain a series of SFT policies on this data scaling ladder. We then train a series of models by running one iteration of RFT on data obtained from the SFT policies at each step.

**Scaling results with positive synthetic data GPT-4 and Gemini 1.5 Pro.** Since we assume that the more capable models generate correct solutions for new problems, by scaling  $\mathcal{D}_{\text{syn}}$  we are increasing *coverage* under  $p_{\text{real}}$ , i.e., adding new  $x, y$  with non-zero probability under  $p_{\text{real}}$ . In Figures 2(a,b), we plot the test error rate of the SFT policy as  $\mathcal{D}_{\text{syn}}$  is scaled. As expected, we observe that the test error rate on both GSM8K and MATH improves with more positive data. Further, by simply fitting the parametric scaling law from [71], for  $D := |\mathcal{D}_{\text{syn}}|$ , we find that the scaling trends decay as  $\approx D^{-0.15}$  on GSM8K and  $\approx D^{-0.05}$  on the harder MATH dataset, with similar trends for the corresponding pass@5 error rates. Since these scaling trends are much more underwhelming than those for pre-training [20], this perhaps implies that samples in  $\mathcal{D}_{\text{syn}}$  are indeed improving coverage over samples in  $p_{\text{real}}(x, y)$ , but maybe not as efficiently as sampling *iid* samples directly from it.

**Scaling results with positive synthetic data from 7B SFT policy.** Previously, we scaled problems in  $\mathcal{D}_{\text{syn}}$  by querying GPT-4 and Gemini-1.5. Now, for existing problems in  $\mathcal{D}_{\text{syn}}$  we generate new responses by sampling from the  $\pi_{\text{sft}}$  trained on problems+olutions in  $\mathcal{D}_{\text{syn}}$ . For any  $(x, y) \in \mathcal{D}_{\text{syn}}$



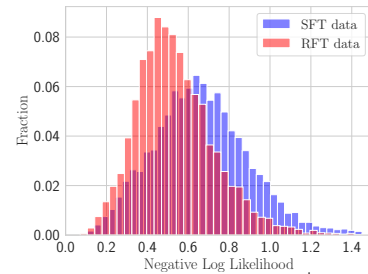
**Figure 2: Positive data scaling laws:** On GSM8K (a) and MATH (b), we evaluate SFT trained on  $\mathcal{D}_{\text{syn}}$  and RFT that uses SFT policy generated positives ( $\mathcal{D}_{\pi_{\text{sft}}}^+$ ), as we scale  $\mathcal{D}_{\text{syn}}$ , observing  $\mathcal{D}_{\pi_{\text{sft}}}^+$  to be  $2\times$  as effective as  $\mathcal{D}_{\text{syn}}$ . In (c), we plot performance of RFT the number of correct solutions in  $\mathcal{D}_{\pi_{\text{sft}}}^+$  are scaled, for a fixed set of 8k/16k problems from  $\mathcal{D}_{\text{syn}}$ , observing that scaling model positives can amplify spurious correlations.

we generate verified positive solution traces  $\hat{\mathbf{y}} \sim \pi_{\text{sft}}$  s.t.  $r(\hat{\mathbf{y}}, \mathbf{y}) = 1$ . Following Yuan et al. [67], to ensure we sample enough correct responses, we sample 100 times from  $\pi_{\text{sft}}$  and generate RFT datasets  $\mathcal{D}_{\pi_{\text{sft}}}^+$ , where each problem has atmost 4 correct and diverse solutions. Next, we finetune the pretrained DeepSeek-Math-7B model on these new series of RFT datasets and plot the performance on GSM8K and MATH (Figure 2(a,b)). First, **we observe that for any size of  $\mathcal{D}_{\text{syn}}$ , the performance of the RFT model is better than the corresponding SFT model**, and the difference remains consistent as we scale  $\mathcal{D}_{\text{syn}}$ . Surprisingly, this indicates that training on positive answer traces from the 7B  $\pi_{\text{sft}}(\mathbf{y} | \mathbf{x})$  can lead to better performing policies than capable models.

**What is the value of positives from  $\pi_{\text{sft}}(\mathbf{y} | \mathbf{x})$ ?** If sampling from  $\pi_{\text{sft}}$  also improves coverage and performance, then should we scale problems and solutions in  $\mathcal{D}_{\text{syn}}$ , or just solutions in  $\mathcal{D}_{\pi_{\text{sft}}}^+$ ? To answer this, we need to assign a value to the RFT dataset  $\mathcal{D}_{\pi_{\text{sft}}}^+$  in terms of  $|\mathcal{D}_{\text{syn}}|$ . We do this by training SFT policies on  $\mathcal{D}_{\text{syn}}$  of sizes 8k and 16k, and then generating RFT datasets from the corresponding SFT policies where we only add more correct solution traces (for the same problems) and scale RFT data from 10k to 128k (unlike RFT data in Figure 2(a,b) where both questions and answers scale). In Figure 2(c) we plot the error rate of DeepSeek-Math-7B finetuned on the different sizes of  $\mathcal{D}_{\pi_{\text{sft}}}^+$ . Comparing the lowest values of the curves in Figure 2(c) with  $\mathcal{D}_{\text{syn}}$  scaling in Figure 2(a,b), we note that **performance from  $\mathcal{D}_{\pi_{\text{sft}}}^+$  is  $2\times$  the size of  $\mathcal{D}_{\text{syn}}$  used to train  $\pi_{\text{sft}}$** . We also note that performance can plateau (or worsen in the case of GSM8K) as we scale up  $\mathcal{D}_{\pi_{\text{sft}}}^+$  by a lot. This is because  $r(\cdot, \mathbf{y})$  is unable to verify the correctness of each step in the positive solution traces in  $\mathcal{D}_{\pi_{\text{sft}}}^+$ . Later, we see how incorrect steps induce spurious correlations that get amplified as we scale positive data, explaining this drop. See Appendix C for more discussion.

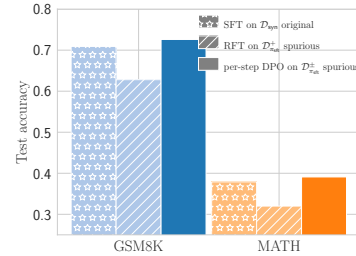
### Why is self-generated positive data more sample-efficient?

From our result above, we find that solutions sampled from  $\pi_{\text{sft}}$  (trained on  $\mathcal{D}_{\text{syn}}$ ) yield better models, as good as those trained on  $2 \times |\mathcal{D}_{\text{syn}}|$ . This finding is surprising since one might expect more capable GPT-4/Gemini models to present better solutions, training on which should lead to good performance, akin to distillation [50], but this is not the case. Our results are consistent with the study of memorization in LLMs [18, 26, 56], which shows that pretrained (base) LLMs tend to memorize “hard-to-fit” and “out-of-pretraining-distribution” responses during finetuning, resulting in imperfect generalization. In contrast, correct response traces produced by  $\pi_{\text{sft}}$  on problems from  $\mathcal{D}_{\text{syn}}$  are not as hard-to-fit or as out-of-distribution, since they are obtained from a model that is “close” to the base LLM. We confirm this hypothesis with a histogram of negative log-likelihood values of the SFT and RFT data under the base LLM (Figure 3). Hence, we expect STaR/RFT to alleviate the memorization problem on a large chunk of examples. This finding also corroborates Yuan et al. [69]’s result that lower the perplexity of SFT data under the base model, the smaller the gap between SFT and RFT performance. Note that one may also attribute better performance of RFT to improved coverage from multiple answers in  $\mathcal{D}_{\pi_{\text{sft}}}^+$  for each question in  $\mathcal{D}_{\text{syn}}$ . But, we find that even when RFT data is restricted to one solution per question, LLM trained on it outperforms SFT consistently by  $> 1\%$ . Since verification is cheap, we can sample more solutions and also benefit from coverage.



**Figure 3:** Under base LLM,  $\mathcal{D}_{\pi_{\text{sft}}}^+$  has higher likelihood than  $\mathcal{D}_{\text{syn}}$ .

**SFT/RFT policy suffers from spurious correlations in positive synthetic data.** While RFT data maybe “easier-to-fit”, in Figure 2(c) we also note that continuing to scale RFT data leads to test error saturation, or even worse test error. This is unlike scaling of problems and solutions in SFT data (in Figure 2(a,b)). This failure can be attributed to the presence of incorrect/irrelevant steps that are not detected by our verifier, since it only verifies the final answer (see Appendix J, K for examples). For a problem  $x$ , when the LLM is trained with supervised next-token prediction on some positive sub-optimal  $y$  in the RFT data, with incorrect step  $y_k$ , it is likely to overfit on spurious correlations between the sub-optimal subsequence  $y_{1:k}$ , and the following valid step  $y_{k+1}$ , when trying to maximize  $\pi(y_{k+1} | y_{1:k}, x)$ . To verify this hypothesis, we amplify the presence of these spurious steps. Specifically, for each question in  $\mathcal{D}_{\text{syn}}$  we sample “spurious steps” from  $\pi_{\text{sft}}$  trained on it, *i.e.*, steps which lead to the incorrect answer with high probability under  $\pi_{\text{sft}}$  (we sample multiple completions conditioned on the same spurious step to check how likely it leads to the correct final answer). Then, we interleave the solution traces in the RFT data with these spurious steps. Note, that all traces in the RFT data are still positive since, they all lead to the correct answer eventually. We find that the LLM trained on this sub-optimal spurious RFT data performs worse than the  $\pi_{\text{sft}}$  policy itself.



**Figure 4:** Spurious correlations in RFT data hurt performance.

#### Takeaways for scaling positive synthetic data

- While positive data from GPT-4/Gemini-1.5 improves coverage over new problems and solutions, positive data from SFT policy trained on it is 2× more sample efficient.
- Scaling positive data ( $\sim \pi_{\text{sft}}$ ) that contains spurious steps, leads to worse test errors.

## 6 Negative Synthetic Data Enables Per-Step Credit Assignment

The spurious correlations from Section 5 correspond to intermediate irrelevant or incorrect steps that are able to still steer the model towards the correct response on some training problems, but derail it otherwise. In this section, we present a conceptual model for constructing negatives that enables us to perform *per-step credit assignment*, and show that this approach can help us address these failure modes of positive data. We show that per-step DPO from Section 3 is a variant of this more general approach. We will then analyze scaling laws with negative data and empirically demonstrate that carefully constructed negative data can address issues with memorization. Finally, we theoretically prove that negative data improves sample-efficiency of  $\mathcal{D}_{\text{syn}}$ .

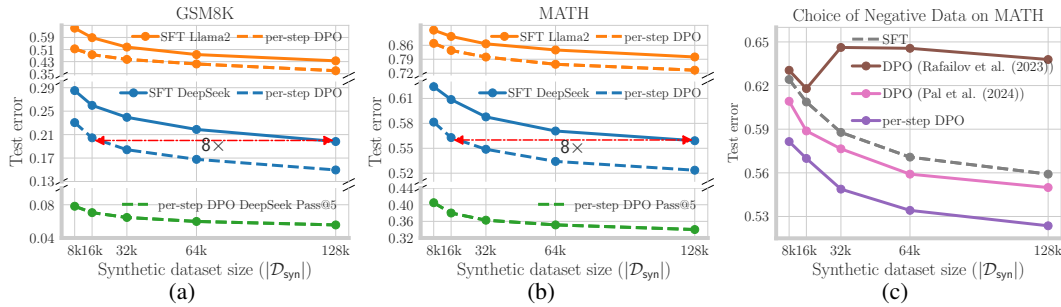
### 6.1 Conceptual Model: Constructing Negatives to Enable Per-Step Credit Assignment

While naïvely contrasting an entire positive response  $+y$  against an entire negative response  $-y$  will increase the likelihood of *each* step that appears in  $+y$  (even when incorrect or irrelevant) and reduce likelihood on each step appearing in  $-y$  (even when accurate and relevant), it does not account for the importance of each step. Formally, given a negative solution trace  $-y$ , we would want to identify the first *critical* step where the model introduces a flaw  $-y$ , and emphasize alternate correct completions from this step that the model could have still produced. Likewise, given a positive solution trace,  $+y$ , we would like to identify if a given step  $+y_i$  does not make progress towards the solution by identifying if there exist alternatives from its predecessor step,  $+y_{1:i-1}$ , which now presents a key decision-making point. **What are these critical steps and how can we identify them procedurally?**

**Value functions.** We can formalize this notion of a critical step under the notion of value functions from reinforcement learning (RL). Recall that both  $+y$  and  $-y$  are sampled from  $\pi_{\text{sft}}$ . For problem  $x$ , with correct solution  $y$ , a response  $\hat{y}$  with a sequence of steps  $\hat{y}_{1:i-1}$ , and a candidate step  $\hat{y}_i$ , we define the value function for step  $y_i$ , and previous steps under some policy  $\tilde{\pi}$  as:

$$Q_{\tilde{\pi}}(\underbrace{x, \hat{y}_{1:i-1}}_{\text{state}}, \underbrace{\hat{y}_i}_{\text{action}}) = \underbrace{\mathbb{E}_{\mathbf{y}_{i+1:L}^{\text{new}} \sim \tilde{\pi}(\cdot | x, \hat{y}_{1:i})}}_{\text{expected future reward under new actions sampled by policy } \tilde{\pi}} \left[ r([\hat{y}_{1:i}, \mathbf{y}_{i+1:L}^{\text{new}}], y) \right] \quad (2)$$

Intuitively, for any partial solution upto  $i$  steps, this Q-function evaluates the probability of succeeding at solving the problem given the remaining budget of  $L - i$  more steps, in expectation over all possible futures sampled from some policy  $\tilde{\pi}$ . Our conceptual model treats the policy  $\tilde{\pi}$  as an algorithmic design choice that can differ for algorithms using negative data. As we see later, choosing  $\tilde{\pi}$  as the Best-of-K distribution around  $\pi_{\text{sft}}$  (denoted as BoK( $\pi_{\text{sft}}$ )) enables a particularly interesting tradeoff between Q-value estimation and policy improvement. Another common choice is  $\pi_{\text{sft}}$  itself. Now,



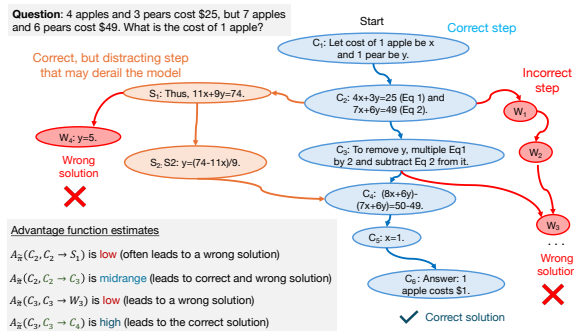
**Figure 5: Negative data scaling laws:** We evaluate algorithms that consume negative data as we scale  $\mathcal{D}_{\text{syn}}$ , and compare them with only positive training (SFT) on  $\mathcal{D}_{\text{syn}}$ . On GSM8K (a) and MATH (b), we observe an 8 $\times$  gain from per-step DPO (Section 4) which aligns with our model of negative data that enables per-step credit assignment. In (c) we compare different negative data construction algorithms, and particularly note that naïvely pairing positives and negatives [41] leads to worse performance as we scale  $\mathcal{D}_{\text{syn}}$ .

for any given step  $\hat{y}_i$ , we can define its *advantage* as the relative change in  $Q_{\bar{\pi}}$  when adding step  $\hat{y}_i$  in comparison with other possible candidates for step  $i$  as follows:

$$A_{\bar{\pi}}(\mathbf{x}, \hat{y}_{1:i-1}; \hat{y}_i) = Q_{\bar{\pi}}(\mathbf{x}, \hat{y}_{1:i-1}, \hat{y}_i) - Q_{\bar{\pi}}(\mathbf{x}, \hat{y}_{1:i-2}, \hat{y}_{i-1}). \quad (3)$$

Equation 3 is identical to the definition of advantage of an action (*i.e.*,  $\hat{y}_i$ ) at a state  $(\mathbf{x}, \hat{y}_{1:i-1})$  from RL [53], in that it is the gap between the Q-value of a state-action pair and the value function of the state (which itself is equal to the Q-value of the *previous* step due to deterministic dynamics).

**Critical steps, per-step DPO, and advantage-weighted RL.** We can use advantages (Equation 3) to characterize critical steps. Steps that attain a higher advantage value than others are **critical** since they need to be generated more precisely to solve the problem. In contrast, steps that with very low advantage values are likely worse and must be **unlearned**. Our definition of the advantage function implies that one can calculate advantages for each step in a response via additional Monte Carlo rollouts starting from prefixes defined by partial solutions. One could then use these advantage estimates



**Figure 6:** Illustration of advantage estimation from negative data on a didactic example in synthetic model generations. Critical steps are those with high advantage values.

(Equation 3) for training the model, for example, by running advantage-weighted reinforcement learning [40]. An alternate option would be to skip the computation of advantage estimates but instead rely on implicit approaches that optimize the advantage-weighted objective without computing their values. Theorem 6.1 shows that DPO performed over a precise pair distribution contrasting positive and negative traces obtained via additional rollouts from  $\bar{\pi}$ , on prefixes of a response sampled from  $\pi_{\text{sft}}$  is equivalent to advantage-weighted RL. A proof of Theorem 6.1 is in Appendix E. Note that unlike the standard reduction of DPO to the RL objective under *some* reward function [41, 42], Theorem 6.1 is stronger in that it identifies the value function induced by per-step DPO.

**Theorem 6.1** (Equivalence of advantage-weighted RL and DPO with per-step pairs). *The optimal policy from Equation 1 with  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm}$  given by  $(\mathbf{x}, [\mathbf{y}_{1:i}, +\mathbf{y}_{i+1}], [\mathbf{y}_{1:i}, -\mathbf{y}_{i+1}])$  where the positive and negative traces share prefix  $\mathbf{y}_{1:i} \sim \pi_{\text{sft}}$ , and  $-\mathbf{y}_{i+1} \sim \pi_{\text{sft}}(\cdot | \mathbf{x}, \mathbf{y}_{1:i})$ ,  $+\mathbf{y}_{i+1} \sim \sigma(A_{\bar{\pi}}(\mathbf{x}, \mathbf{y}_{1:i}; \cdot) - A_{\bar{\pi}}(\mathbf{x}, \mathbf{y}_{1:i}; -\mathbf{y}_{i+1}))$ , is identical to the optima of the advantage-weighted RL objective:*

$$\max_{\pi} \mathbb{E}_{\mathbf{x} \sim p_{\text{syn}}(\mathbf{x}), \mathbf{y} \sim \pi_{\text{sft}}(\cdot | \mathbf{x})} \left[ \sum_{i=1}^L \log \pi(\mathbf{y}_i | \mathbf{x}, \mathbf{y}_{0:i-1}) \cdot \exp(A_{\bar{\pi}}(\mathbf{x}, \mathbf{y}_{0:i-1}; \mathbf{y}_i) / \beta) \right]. \quad (4)$$

**Practical instantiation of DPO with per-step pairs.** In most of our experiments, we instantiate a practical version of the above framework, following the scheme in Hwang et al. [23]. This is a special case (Part 1) of the complete algorithm shown in Algorithm 1 (see Appendix B). Unless otherwise mentioned, we use “per-step DPO” to refer to this version (Part 1 only) in practice. We will also experiment with the complete version (parts 1 and 2) later in Section 6.3.3. Instead of computing

advantage estimates for each step, and then sampling preference pairs, as described in Theorem 6.1, we approximate this by only Q-value estimates on 8 negative responses for each question in the synthetic dataset, with  $\tilde{\pi}$  chosen to be the best-of-K policy,  $\text{BoK}(\pi_{\text{sft}})$  where  $K = 5$ . There are two benefits associated with this choice of  $\tilde{\pi}$ , especially a higher value of  $K$ : (i) estimating the advantage in Equation 3 with Monte Carlo rollouts exhibits lower variance when  $K$  is large, since a larger budget  $K$  would lead most steps to have higher Q-values and the variance of Bernoulli reduces as  $Q\text{-value} \rightarrow 1$ ; and (ii)  $Q_{\text{BoK}(\pi_{\text{sft}})}$  is a non-decreasing function in  $K$  for any state-action, which implies that the solution of advantage-weighted RL objective, in principle, can now improve over a better policy  $\text{BoK}(\pi_{\text{sft}})$ , compared to  $\pi_{\text{sft}}$ . Next, we discuss scaling results for negative data, and then in Section 6.3 show how per-step credit assignment improves generalization and suppresses irrelevant and incorrect steps appearing in a response, extracting more gains from the same synthetic data.

## 6.2 Scaling Results for Negative Data

Observe in Figure 5(a,b), that for both DeepSeek-Math-7B and LLaMA2-7B models, per-step DPO improves performance beyond the SFT policy and the performance continues to scale favorably as data size increases. In fact, for any given size of  $\mathcal{D}_{\text{syn}}$ , per-step DPO also substantially improves over RFT (Figure 2) on both datasets, and overall, **while RFT improved effective data size of  $\mathcal{D}_{\text{syn}}$  by 2 $\times$ , additionally training on negative data extends the performance improvement to 8 $\times$  the size of  $\mathcal{D}_{\text{syn}}$** . Additionally, since per-step DPO estimates advantage of each step under the Best-of-5 policy, one might expect a saturation in the pass@5 performance of the per-step DPO solution. On the contrary, we find that pass@5 performance also improves consistently. In Appendix D we present results for a filtered version of RFT. Here, steps with high advantages from positive/negative data are cloned. This resolves the scaling issue seen when naively scaling positive data in Figure 2(c).

**Choice of negative data matters.** In Figure 5(c) we plot negative data scaling laws where the choice of negative data (thereby pairs for DPO in Equation 1) differs. Observe that standard pairing of positive and negative responses in  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm}$  for DPO [41] does not improve over the SFT policy. As such, we tuned  $\beta$  in Equation 1 for DPO but could not fully avoid performance degradation. Our conceptual model explains this result: contrasting arbitrary positives and negatives would result in an incorrect induced advantage function, training with DPO will exacerbate spurious correlations that maximize this induced advantage function [39, 46, 64]. In fact, Pal et al. [38] also find similar concerns with random pairing and instead pair positives and negatives with highest edit distance, which leads to some improvement, but still performs poorer than per-step DPO that accounts for credit.

### Takeaways for scaling negative synthetic data

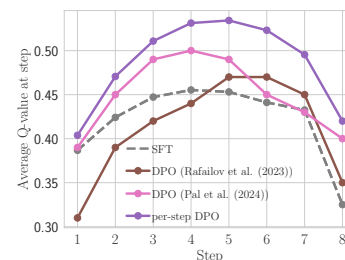
- Negative data can identify high-advantage (critical) steps in model-generated responses.
- We can construct negative data distribution that equates DPO to advantage-weighted RL. Negative data used in this way improves the sample efficiency of synthetic data by 8 $\times$ .

## 6.3 Why Does Credit Assignment from Negative Data Improve Model Generalization?

Our conceptual model illustrates that per-step DPO can perform credit assignment, and identify critical steps over irrelevant ones via advantage estimates. We saw that this improves test performance and scaling. Now, we attempt to understand why per-step credit assignment should improve generalization by understanding the generalization properties of advantage-weighted RL. We present two empirical studies below, and a formal theoretical guarantee combining these insights is shown in Appendix F.

### 6.3.1 Advantage-Weighted RL De-Emphasizes Spurious Steps and Emphasizes Critical Steps

Our key insight is that spurious correlations emerge in monolithic SFT or RFT due to the well-known issue of causal confusion [12] in imitation learning: by memorizing incorrect or irrelevant steps and associating them with the correctness of the final answer, the model fails to generalize on novel problems, as we saw in Figure 4. We now explain how *online* model-specific interventions and advantage estimation would address this issue. Consider  $\tilde{\pi} = \pi_{\text{sft}}$ . As we show later, in under-trained models memorized steps are imperfectly cloned under  $\pi_{\text{sft}}$ , implying that while teacher-forcing loss is low for some spurious, memorized step  $y_s$ , sampling paths from  $\pi_{\text{sft}}$ , conditioned on  $y_{1:s}$  is likely to generate incorrect responses. This means  $y_s$  attains a low advantage. On the other



**Figure 7:** Per-step DPO improves Q-values at each step, standard DPO only improves at irrelevant steps.

hand, for a correct step, *whp* estimated advantage is higher. Thus, training the model with advantage weighted RL would de-emphasize spurious steps and emphasize critical steps. Running per-step DPO on data generated by the RFT model that has overfit on spurious correlations improves accuracy by >6% (Figure 4). We visualize advantages in Appendix K. In Figure 7, we plot the average Q-value of a step for different negative data schemes, and note that only per-step DPO improves over SFT at each step, as expected based on the connection to advantage-weighted RL (Theorem 6.1). Standard DPO fails to improve performance since it has poor success rate at earlier (critical) steps.

### 6.3.2 Why Does Generalization Improve?: Connecting Advantage-Weighted RL to DRO

In the previous section, we discussed how advantage-weighted RL preferentially weighs the next-token prediction loss at each step. Now, we attempt to conceptually understand why this could improve generalization. For this, we present an intuitive explanation by drawing a connection between advantage-weighted RL and a distributionally robust optimization (DRO) algorithm, named Group DRO, commonly used to improve worst-group robustness in supervised learning [43].

**Intuitive explanation.** During inference, the SFT policy can fail even on training problems, especially in scenarios where the SFT policy has failed to perfectly clone the next step at each intermediate step in the SFT data. As previously discussed, these steps also present with low advantage values. One way to reduce the chance of compounding inference time errors [45] is to preferentially minimize the negative log-likelihood loss *more* for the critical step, i.e., those steps from where the model is more likely to arrive at a wrong answer. If we iteratively update the policy with gradient steps computed over a re-weighted next-step prediction objective where each step is weighted by its advantage estimate, then the resulting algorithm intuitively exhibits this characteristic similarly to distributionally robust optimizers (DRO) [28]. Similar to how DRO solutions guarantee that all subpopulations – both majority and minority subpopulations – in the training data achieve low loss values, the solution for the advantage-weighted RL objective guarantees that the negative log-likelihood of the critical steps with high advantage estimates under  $\tilde{\pi}$  (which of per-step DPO is  $\text{BoK}(\pi_{\text{sft}})$ ) is also low, to a similar extent as the other more prevalent non-critical steps.

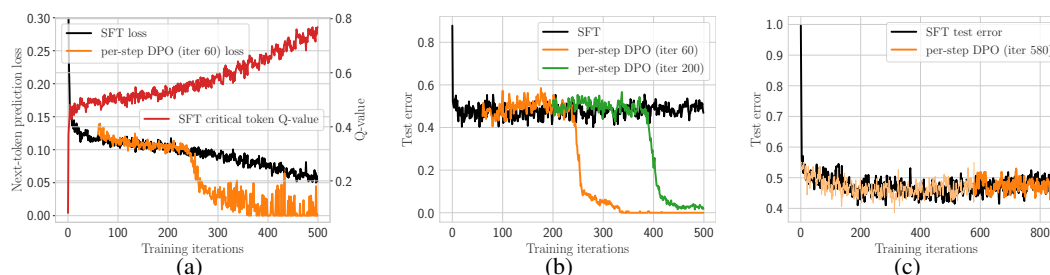
In other words, **our insight** is that weighting steps using advantages in Equation 4 upweights the likelihood of the underrepresented critical states while down-weighting it for the spurious ones. The guarantees on the training data ( $\mathcal{D}_{\text{syn}}$ ) also translate to the population level objective when the weights for on-policy samples (advantage estimates) are accurate [49] and the policy is sufficiently regularized [47]. Since correct behavior at critical steps determine the correctness of the overall solution, an elevated degree of correctness at executing critical steps at the population level implies a higher test accuracy on the reasoning task.

### 6.3.3 But, Attaining Low Generalization Error Requires Low Advantage Estimation Error

The practical efficacy of algorithms that use negative data for credit assignment requires the advantage estimation error to be low with fewer rollouts from  $\tilde{\pi}$ . For discussion, consider  $\tilde{\pi} = \pi_{\text{sft}}$ . When the initial advantage of a spurious step is incorrectly over-estimated, negative data algorithms up-weight the likelihood further. This only leads to further memorization. Hence, most Monte-Carlo rollouts from  $\pi_{\text{sft}}$  would rely upon the memorized feature. Since the model generates the correct answer from the memorized feature, it would estimate higher  $A_{\pi_{\text{sft}}}$ , and this downward spiral of training with increasing weights on the spurious step leads to test-time model collapse. On the other hand, when  $\tilde{\pi} = \text{BoK}(\pi_{\text{sft}})$  for a higher value of  $K$ , the Monte-Carlo advantage estimator has a lower variance (and error). This discussion also justifies the choice of  $K=5$ , an intermediate value, in per-step DPO.

### 6.3.4 Validating Claims About Generalization: Controlled Analysis on a Didactic Problem

With the above insights, we now study the influence of  $\pi_{\text{sft}}$  on the generalization effects of per-step DPO. For our analysis, we consider a didactic star graph problem (Appendix I) from Bachmann and Nagarajan [4], where given a graph in the shape of a star and a query (center/end node), the model is asked to output the full path between the start/end nodes. This task highlights the failure of SFT at planning problems (akin to math reasoning). They show that  $\pi_{\text{sft}}$  minimizes SFT loss by memorizing the “hard-to-predict” node adjacent to the center, and copying the rest from the input graph. It is clear that the failure stems from not being able to identify the critical adjacent token. We will show how credit assignment with negative data accurately upweights the critical token and unlearns the memorized token. To vary the choice of  $\pi_{\text{sft}}$ , we choose several intermediate checkpoints obtained during supervised finetuning for synthetic negative data generation. We consider three initializations: (1) an under-trained SFT model with a large training and test loss, and (2) an SFT model obtained by



**Figure 8: Didactic analysis on star graph:** In (a) we plot the SFT loss and Q-value of the critical token (adjacent node) for SFT and per-step DPO (starting from iter 60). Indicative of memorization SFT loss decreases at a slow rate, matching the slow rate of increase in the Q-value. In contrast per-step DPO loss sharply decreases during training. In (b) we notice a corresponding phase transition in the test error of per-step DPO starting from different under-trained SFT checkpoints, which does not happen for an over-trained SFT checkpoint in (c).

early-stopping based on a held-out validation set, where the validation loss is the lowest, and (3) an over-trained SFT checkpoint, with a low training but high validation loss.

**(1) & (2): Training on negative data from an under-trained or early-stopped  $\pi_{\text{sft}}$  improves both training loss and test performance.** As shown in Figure 8(a,b), we find that when training with negative data from iteration 60 (under-trained  $\pi_{\text{sft}}$ ) and iteration 200 (early-stopped  $\pi_{\text{sft}}$ ), utilizing per-step DPO reduces the training loss very aggressively. These benefits translate to test losses and performance as well (Figure 8(b), orange and green). In contrast, supervised finetuning exhibits a nearly-flat test loss landscape, although the train loss reduces slowly. Upon a closer inspection, we find that training on positive data via SFT only tends to memorize the critical token in the training data using non-generalizable features, and hence, the resulting model does not generalize to novel problems. More training with SFT is unable to “unlearn” this spurious correlation and does not reduce the loss function. On the other hand, per-step DPO with negative data is able to unlearn this spurious feature and drives improvement, as evident by the drastic improvement on train and test.

**(3) Training on negative data from an over-trained SFT initialization leads to model collapse.** When training with negative data on an over-trained  $\pi_{\text{sft}}$  (iteration 580) in Figure 8(c), we observe that both SFT and per-step DPO exhibit identical test errors since training with more negative data simply exacerbates the model’s dependence on memorizing the critical token, which manifests in the form of lower test losses. This is also an example where Monte-Carlo samples from the over-trained checkpoint estimates a high advantage since Q-value is already high at iteration 500 (in (a)). This means that when the SFT policy has sufficiently memorized the training data using a spurious feature, training further is unable to unlearn this dependence. Hence, we find that in this regime, negative data leads to no improvement, capping performance at what was attained by fine-tuning on positive data.

#### Takeaways for generalization and spurious correlations with negative data

Advantage-weighted RL unlearns spurious steps and improves generalization when: (i) advantage estimation error is low; and (ii) the model is under-trained enough that imperfectly cloned spurious steps have low advantage, which can then be estimated with negative data.

## 7 Discussion and Conclusion

Our work studies the role of synthetic data for improving math reasoning capabilities of LLMs. We find that while the typical approach of collecting new questions and corresponding positive (correct) solutions from capable models like GPT-4/Gemini-1.5 presents underwhelming data scaling. The sample efficiency of the same data can be improved up to  $2\times$  by sampling more positive traces from the 7B sized models SFT-ed on the original data. However, training on positive self-generated synthetic data alone often amplifies the model’s dependence on spurious steps, that erroneously appear to lead to a good solution but do not generalize to novel problems and hurt test performance. That said, surprisingly, we show that negative (incorrect) traces sampled from the same SFT model can be used to address the failure modes of training on only positive data. In particular, negative data can be used to estimate advantage values for every step, and using these advantage estimates via RL enables us to address this problem. We show how the advantages can be used implicitly by preference optimization objectives. We show how training on an instance of this objective leads to  $8\times$  improvements in sample efficiency of the synthetic data used.

## Acknowledgements

This work was done at CMU. We thank Vaishnavh Nagarajan, Yi Su, Aleksandra Faust, Hyeonbin Hwang, Christina Baek, Charlie Snell, Seohong Park, Gaurav Ghosal, Aditi Raghunathan, Katie Kang, Don Dennis, Dhruv Malik, and Pratiksha Thaker for informative discussions and feedback on an earlier version of this paper. This work was supported by compute donations from Google Cloud (TRC) and MultiOn. AS thanks OpenAI and Google respectively for providing GPT4-Turbo and Gemini-1.5 Pro credits for academic use. AK and YG thank Tianhe Yu for feedback on the paper. This work was supported in part by the National Science Foundation grants IIS2145670 and CCF2107024, and funding from Amazon, Apple, Google, Intel, Meta, and the CyLab Security and Privacy Institute. Any opinions, findings and conclusions expressed in this material are those of the author(s) and do not necessarily reflect the views of any of these funding agencies.

## References

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [2] Alekh Agarwal, Nan Jiang, Sham M Kakade, and Wen Sun. Reinforcement learning: Theory and algorithms. *CS Dept., UW Seattle, Seattle, WA, USA, Tech. Rep*, 2019.
- [3] Sina Alemohammad, Josue Casco-Rodriguez, Lorenzo Luzi, Ahmed Imtiaz Humayun, Hossein Babaei, Daniel LeJeune, Ali Siahkoobi, and Richard G Baraniuk. Self-consuming generative models go mad. *arXiv preprint arXiv:2307.01850*, 2023.
- [4] Gregor Bachmann and Vaishnavh Nagarajan. The pitfalls of next-token prediction, 2024.
- [5] Xiao Bi, Deli Chen, Guanting Chen, Shanhuang Chen, Damai Dai, Chengqi Deng, Honghui Ding, Kai Dong, Qiushi Du, Zhe Fu, et al. Deepseek llm: Scaling open-source language models with longtermism. *arXiv preprint arXiv:2401.02954*, 2024.
- [6] Ralph Allan Bradley and Milton E Terry. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952.
- [7] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901, 2020.
- [8] Zixiang Chen, Yihe Deng, Huizhuo Yuan, Kaixuan Ji, and Quanquan Gu. Self-play fine-tuning converts weak language models to strong language models. *arXiv preprint arXiv:2401.01335*, 2024.
- [9] Pengyu Cheng, Yifan Yang, Jian Li, Yong Dai, and Nan Du. Adversarial preference optimization. *arXiv preprint arXiv:2311.08045*, 2023.
- [10] Wei-Lin Chiang, Zhuohan Li, Zi Lin, Ying Sheng, Zhanghao Wu, Hao Zhang, Lianmin Zheng, Siyuan Zhuang, Yonghao Zhuang, Joseph E. Gonzalez, Ion Stoica, and Eric P. Xing. Vicuna: An open-source chatbot impressing gpt-4 with 90%\* chatgpt quality, March 2023. URL <https://lmsys.org/blog/2023-03-30-vicuna/>.
- [11] Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*, 2021.
- [12] Pim De Haan, Dinesh Jayaraman, and Sergey Levine. Causal confusion in imitation learning. *Advances in neural information processing systems*, 32, 2019.
- [13] Elvis Dohmatob, Yunzhen Feng, and Julia Kempe. Model collapse demystified: The case of regression, 2024.

- [14] Guanting Dong, Hongyi Yuan, Keming Lu, Chengpeng Li, Mingfeng Xue, Dayiheng Liu, Wei Wang, Zheng Yuan, Chang Zhou, and Jingren Zhou. How abilities in large language models are affected by supervised fine-tuning data composition. *arXiv preprint arXiv:2310.05492*, 2023.
- [15] Nouha Dziri, Ximing Lu, Melanie Sclar, Xiang Lorraine Li, Liwei Jiang, Bill Yuchen Lin, Sean Welleck, Peter West, Chandra Bhagavatula, Ronan Le Bras, et al. Faith and fate: Limits of transformers on compositionality. *Advances in Neural Information Processing Systems*, 36, 2024.
- [16] Kawin Ethayarajh, Winnie Xu, Niklas Muennighoff, Dan Jurafsky, and Douwe Kiela. Kto: Model alignment as prospect theoretic optimization. *arXiv preprint arXiv:2402.01306*, 2024.
- [17] Matthias Gerstgrasser, Rylan Schaeffer, Apratim Dey, Rafael Rafailov, Henry Sleight, John Hughes, Tomasz Korbak, Rajashree Agrawal, Dhruv Pai, Andrey Gromov, Daniel A. Roberts, Diyi Yang, David L. Donoho, and Sanmi Koyejo. Is model collapse inevitable? breaking the curse of recursion by accumulating real and synthetic data, 2024.
- [18] Valentin Hartmann, Anshuman Suri, Vincent Bindschaedler, David Evans, Shruti Tople, and Robert West. Sok: Memorization in general-purpose large language models, 2023.
- [19] 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. *NeurIPS*, 2021.
- [20] Jordan Hoffmann, Sebastian Borgeaud, Arthur Mensch, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, et al. Training compute-optimal large language models. *arXiv preprint arXiv:2203.15556*, 2022.
- [21] Jiwoo Hong, Noah Lee, and James Thorne. Reference-free monolithic preference optimization with odds ratio. *arXiv preprint arXiv:2403.07691*, 2024.
- [22] Arian Hosseini, Xingdi Yuan, Nikolay Malkin, Aaron Courville, Alessandro Sordoni, and Rishabh Agarwal. V-star: Training verifiers for self-taught reasoners. *arXiv preprint arXiv:2402.06457*, 2024.
- [23] Hyeonbin Hwang, Doyoung Kim, Seungone Kim, Seonghyeon Ye, and Minjoon Seo. Self-explore to avoid the pit: Improving the reasoning capabilities of language models with fine-grained rewards. *arXiv preprint arXiv:2404.10346*, 2024.
- [24] Matti Kääriäinen. Lower bounds for reductions. In *Atomic Learning Workshop*, 2006.
- [25] Sham Kakade and John Langford. Approximately optimal approximate reinforcement learning. In *International Conference on Machine Learning (ICML)*, volume 2, 2002.
- [26] Katie Kang, Eric Wallace, Claire Tomlin, Aviral Kumar, and Sergey Levine. Unfamiliar finetuning examples control how language models hallucinate, 2024.
- [27] Aviral Kumar, Joey Hong, Anikait Singh, and Sergey Levine. When Should We Prefer Offline Reinforcement Learning over Behavioral Cloning? *ICLR*, 2022.
- [28] Daniel Levy, Yair Carmon, John C Duchi, and Aaron Sidford. Large-scale methods for distributionally robust optimization. *Advances in Neural Information Processing Systems*, 33: 8847–8860, 2020.
- [29] Chen Li, Weiqi Wang, Jingcheng Hu, Yixuan Wei, Nanning Zheng, Han Hu, Zheng Zhang, and Houwen Peng. Common 7b language models already possess strong math capabilities. *arXiv preprint arXiv:2403.04706*, 2024.
- [30] Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step, 2023.
- [31] Hao Liu, Matei Zaharia, and Pieter Abbeel. Exploration with principles for diverse ai supervision. *arXiv preprint arXiv:2310.08899*, 2023.

- [32] Ruibo Liu, Jerry Wei, Fangyu Liu, Chenglei Si, Yanzhe Zhang, Jinmeng Rao, Steven Zheng, Daiyi Peng, Diyi Yang, Denny Zhou, and Andrew M. Dai. Best practices and lessons learned on synthetic data for language models, 2024.
- [33] Haipeng Luo, Qingfeng Sun, Can Xu, Pu Zhao, Jianguang Lou, Chongyang Tao, Xiubo Geng, Qingwei Lin, Shifeng Chen, and Dongmei Zhang. Wizardmath: Empowering mathematical reasoning for large language models via reinforced evol-instruct, 2023.
- [34] Liangchen Luo, Yinxiao Liu, Rosanne Liu, Samrat Phatale, Harsh Lara, Yunxuan Li, Lei Shu, Yun Zhu, Lei Meng, Jiao Sun, et al. Improve mathematical reasoning in language models by automated process supervision. *arXiv preprint arXiv:2406.06592*, 2024.
- [35] R Thomas McCoy, Shunyu Yao, Dan Friedman, Matthew Hardy, and Thomas L Griffiths. Embers of autoregression: Understanding large language models through the problem they are trained to solve. *arXiv preprint arXiv:2309.13638*, 2023.
- [36] Ida Momennejad, Hosein Hasanbeig, Felipe Vieira Fruejri, Hiteshi Sharma, Nebojsa Jojic, Hamid Palangi, Robert Ness, and Jonathan Larson. Evaluating cognitive maps and planning in large language models with cogeal. *Advances in Neural Information Processing Systems*, 36, 2024.
- [37] Rémi Munos, Michal Valko, Daniele Calandriello, Mohammad Gheshlaghi Azar, Mark Rowland, Zhaohan Daniel Guo, Yunhao Tang, Matthieu Geist, Thomas Mesnard, Andrea Michi, et al. Nash learning from human feedback. *arXiv preprint arXiv:2312.00886*, 2023.
- [38] Arka Pal, Deep Karkhanis, Samuel Dooley, Manley Roberts, Siddartha Naidu, and Colin White. Smaug: Fixing failure modes of preference optimisation with dpo-positive. *arXiv preprint arXiv:2402.13228*, 2024.
- [39] Richard Yuanzhe Pang, Weizhe Yuan, Kyunghyun Cho, He He, Sainbayar Sukhbaatar, and Jason Weston. Iterative reasoning preference optimization. *arXiv preprint arXiv:2404.19733*, 2024.
- [40] Xue Bin Peng, Aviral Kumar, Grace Zhang, and Sergey Levine. Advantage-weighted regression: Simple and scalable off-policy reinforcement learning. *arXiv preprint arXiv:1910.00177*, 2019.
- [41] Rafael Rafailov, Archit Sharma, Eric Mitchell, Stefano Ermon, Christopher D Manning, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *arXiv preprint arXiv:2305.18290*, 2023.
- [42] Rafael Rafailov, Joey Hejna, Ryan Park, and Chelsea Finn. From  $r$  to  $q^*$ : Your language model is secretly a q-function, 2024.
- [43] Hamed Rahimian and Sanjay Mehrotra. Distributionally robust optimization: A review. *arXiv preprint arXiv:1908.05659*, 2019.
- [44] Machel Reid, Nikolay Savinov, Denis Teplyashin, Dmitry Lepikhin, Timothy Lillicrap, Jean-baptiste Alayrac, Radu Soricut, Angeliki Lazaridou, Orhan Firat, Julian Schrittwieser, et al. Gemini 1.5: Unlocking multimodal understanding across millions of tokens of context. *arXiv preprint arXiv:2403.05530*, 2024.
- [45] Stéphane Ross and Drew Bagnell. Efficient reductions for imitation learning. In *International Conference on Artificial Intelligence and Statistics (AISTATS)*, pages 661–668, 2010.
- [46] Amir Saeidi, Shivanshu Verma, and Chitta Baral. Insights into alignment: Evaluating dpo and its variants across multiple tasks. *arXiv preprint arXiv:2404.14723*, 2024.
- [47] Shiori Sagawa, Pang Wei Koh, Tatsunori B Hashimoto, and Percy Liang. Distributionally robust neural networks. In *International Conference on Learning Representations*, 2019.
- [48] Mohamed El Amine Seddik, Swei-Wen Chen, Soufiane Hayou, Pierre Youssef, and Merouane Debbah. How bad is training on synthetic data? a statistical analysis of language model collapse, 2024.

- [49] Amrith Setlur, Don Dennis, Benjamin Eysenbach, Aditi Raghunathan, Chelsea Finn, Virginia Smith, and Sergey Levine. Bitrate-constrained dro: Beyond worst case robustness to unknown group shifts. *arXiv preprint arXiv:2302.02931*, 2023.
- [50] Archit Sharma, Sedrick Keh, Eric Mitchell, Chelsea Finn, Kushal Arora, and Thomas Kollar. A critical evaluation of ai feedback for aligning large language models, 2024.
- [51] Ilia Shumailov, Zakhar Shumaylov, Yiren Zhao, Yarin Gal, Nicolas Papernot, and Ross Anderson. The curse of recursion: Training on generated data makes models forget. *arXiv preprint arXiv:2305.17493*, 2023.
- [52] Avi Singh, John D. Co-Reyes, Rishabh Agarwal, Ankesh Anand, Piyush Patil, Xavier Garcia, Peter J. Liu, James Harrison, Jaehoon Lee, Kelvin Xu, Aaron Parisi, Abhishek Kumar, Alex Alemi, Alex Rizkowsky, Azade Nova, Ben Adlam, Bernd Bohnet, Gamaleldin Elsayed, Hanie Sedghi, Igor Mordatch, Isabelle Simpson, Izzeddin Gur, Jasper Snoek, Jeffrey Pennington, Jiri Hron, Kathleen Kenealy, Kevin Swersky, Kshiteej Mahajan, Laura Culp, Lechao Xiao, Maxwell L. Bileschi, Noah Constant, Roman Novak, Rosanne Liu, Tris Warkentin, Yundi Qian, Yamini Bansal, Ethan Dyer, Behnam Neyshabur, Jascha Sohl-Dickstein, and Noah Fiedel. Beyond human data: Scaling self-training for problem-solving with language models, 2024.
- [53] Richard S Sutton and Andrew G Barto. *Reinforcement learning: An introduction*. The MIT Press, second edition, 2018.
- [54] Gokul Swamy, Christoph Dann, Rahul Kidambi, Zhiwei Steven Wu, and Alekh Agarwal. A minimaximalist approach to reinforcement learning from human feedback. *arXiv preprint arXiv:2401.04056*, 2024.
- [55] Fahim Tajwar, Anikait Singh, Archit Sharma, Rafael Rafailov, Jeff Schneider, Tengyang Xie, Stefano Ermon, Chelsea Finn, and Aviral Kumar. Preference Fine-Tuning of LLMs Should Leverage Suboptimal, On-Policy Data, ICML 2024.
- [56] Kushal Tirumala, Aram Markosyan, Luke Zettlemoyer, and Armen Aghajanyan. Memorization without overfitting: Analyzing the training dynamics of large language models. *Advances in Neural Information Processing Systems*, 35:38274–38290, 2022.
- [57] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.
- [58] Pablo Villalobos, Jaime Sevilla, Lennart Heim, Tamay Besiroglu, Marius Hobbhahn, and Anson Ho. Will we run out of data? an analysis of the limits of scaling datasets in machine learning. *arXiv preprint arXiv:2211.04325*, 2022.
- [59] Peiyi Wang, Lei Li, Zhihong Shao, R. X. Xu, Damai Dai, Yifei Li, Deli Chen, Y. Wu, and Zhifang Sui. Math-shepherd: Verify and reinforce llms step-by-step without human annotations, 2024.
- [60] Yuanhao Wang, Qinghua Liu, and Chi Jin. Is rlhf more difficult than standard rl? *arXiv preprint arXiv:2306.14111*, 2023.
- [61] Ronald J Williams and David Zipser. A learning algorithm for continually running fully recurrent neural networks. *Neural computation*, 1(2):270–280, 1989.
- [62] Tianhao Wu, Banghua Zhu, Ruoyu Zhang, Zhaojin Wen, Kannan Ramchandran, and Jiantao Jiao. Pairwise proximal policy optimization: Harnessing relative feedback for llm alignment. *arXiv preprint arXiv:2310.00212*, 2023.
- [63] Sierra Wyllie, Ilia Shumailov, and Nicolas Papernot. Fairness feedback loops: Training on synthetic data amplifies bias, 2024.
- [64] Haoran Xu, Amr Sharaf, Yunmo Chen, Weiting Tan, Lingfeng Shen, Benjamin Van Durme, Kenton Murray, and Young Jin Kim. Contrastive preference optimization: Pushing the boundaries of llm performance in machine translation. *arXiv preprint arXiv:2401.08417*, 2024.

- [65] Fei Yu, Anningzhe Gao, and Benyou Wang. Outcome-supervised verifiers for planning in mathematical reasoning. *arXiv preprint arXiv:2311.09724*, 2023.
- [66] Longhui Yu, Weisen Jiang, Han Shi, Jincheng Yu, Zhengying Liu, Yu Zhang, James T. Kwok, Zhenguo Li, Adrian Weller, and Weiyang Liu. Metamath: Bootstrap your own mathematical questions for large language models, 2024.
- [67] Lifan Yuan, Ganqu Cui, Hanbin Wang, Ning Ding, Xingyao Wang, Jia Deng, Boji Shan, Huimin Chen, Ruobing Xie, Yankai Lin, et al. Advancing llm reasoning generalists with preference trees. *arXiv preprint arXiv:2404.02078*, 2024.
- [68] Weizhe Yuan, Richard Yuanzhe Pang, Kyunghyun Cho, Sainbayar Sukhbaatar, Jing Xu, and Jason Weston. Self-rewarding language models. *arXiv preprint arXiv:2401.10020*, 2024.
- [69] Zheng Yuan, Hongyi Yuan, Chengpeng Li, Guanting Dong, Chuanqi Tan, and Chang Zhou. Scaling relationship on learning mathematical reasoning with large language models. *arXiv preprint arXiv:2308.01825*, 2023.
- [70] Eric Zelikman, Yuhuai Wu, Jesse Mu, and Noah Goodman. Star: Bootstrapping reasoning with reasoning. *Advances in Neural Information Processing Systems*, 35:15476–15488, 2022.
- [71] Biao Zhang, Zhongtao Liu, Colin Cherry, and Orhan Firat. When scaling meets llm finetuning: The effect of data, model and finetuning method, 2024.
- [72] Ruiqi Zhang, Licong Lin, Yu Bai, and Song Mei. Negative preference optimization: From catastrophic collapse to effective unlearning. *arXiv preprint arXiv:2404.05868*, 2024.
- [73] Yao Zhao, Mikhail Khalman, Rishabh Joshi, Shashi Narayan, Mohammad Saleh, and Peter J Liu. Calibrating sequence likelihood improves conditional language generation. In *The Eleventh International Conference on Learning Representations*, 2022.

# Appendices

## A Limitations of our Work

While our work provides some results and conceptual models to understand the role of synthetic data for reasoning, there are still many open questions that need to be answered to fully understand its utility. While synthetic data from LLMs like Gemini and GPT-4 holds great potential, for more complex reasoning problems (more complicated than the datasets evaluated in our work), synthetic data generated from more capable models can contain errors. Generating negative/positive data by referencing synthetic data answers can reinforce unwanted spurious correlations highlighted in our work. This means that novel recipes for generating synthetic problems may be utilized in the future, and our analysis might need to be re-done. That said, we believe that our insights about algorithmic behavior with synthetic data are still quite general and should transfer to these novel settings as well. Ultimately, we would want that training on synthetic data improves transfer and generalization abilities of the model in general reasoning scenarios, and to this end, an evaluation of transfer capabilities is an important avenue that future work should focus on.

## B Per-step DPO Algorithm

---

**Algorithm 1** Per-step DPO ([Part 1](#): Practical version for most experiments; [Parts 1 + 2](#): Complete version)

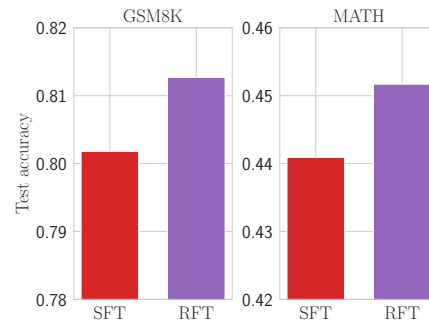
---

**Require:** Synthetic dataset:  $\mathcal{D}_{\text{syn}}$ , SFT policy finetuned on  $\mathcal{D}_{\text{syn}}$ :  $\pi_{\text{sft}}$ , sampling policy  $\tilde{\pi}$ .

- 1: Initialize per-step DPO dataset  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm} \leftarrow \{\}$ .
  - 2: **for**  $(x, y) \in \mathcal{D}_{\text{syn}} \cup \mathcal{D}_{\pi_{\text{sft}}}^+$  **do**
  - 3:   # [Part 1: Identify critical steps in incorrect responses](#)
  - 4:   Sample multiple incorrect answers  $-\hat{y} \sim \pi_{\text{sft}}(\cdot | x)$ , and collect them in set  $\mathcal{C}(x)$ .
  - 5:   **for**  $-\hat{y} := [-\hat{y}_1, \dots, -\hat{y}_L] \in \mathcal{C}(x)$  **do**
  - 6:     Compute the Monte Carlo estimate for  $Q_{\tilde{\pi}}(x, -\hat{y}_{1:i-1}; -\hat{y}_i)$  for each step  $-\hat{y}_i$ .
  - 7:     If  $-\hat{y}_c$  is the first step with least  $Q_{\tilde{\pi}}(x, -\hat{y}_{1:i-1}; -\hat{y}_i)$ , then  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm} \leftarrow \mathcal{D}_{\pi_{\text{sft}}}^{\pm} \cup \{(x, y, -\hat{y}_{1:c})\}$ .
  - 8:   **end for**
  - 9:   # [Part 2: Identify spurious steps in correct responses](#)
  - 10:   Sample multiple correct answers  $+\hat{y} \sim \pi_{\text{sft}}(\cdot | x)$ , and collect them in set  $\mathcal{C}'(x)$ .
  - 11:   **for**  $+\hat{y} := [+ \hat{y}_1, \dots, + \hat{y}_L] \in \mathcal{C}'(x)$  **do**
  - 12:     Compute the Monte Carlo estimate for  $Q_{\tilde{\pi}}(x, +\hat{y}_{1:i-1}; +\hat{y}_i)$  for each step  $+\hat{y}_i$ .
  - 13:     If  $+\hat{y}_c$  is the first step with least  $Q_{\tilde{\pi}}(x, +\hat{y}_{1:i-1}; +\hat{y}_i)$ , then  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm} \leftarrow \mathcal{D}_{\pi_{\text{sft}}}^{\pm} \cup \{(x, y, +\hat{y}_{1:c})\}$ .
  - 14:   **end for**
  - 15: **end for**
  - 16: Optimize DPO loss in Equation (1) on  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm}$  with  $\pi_{\text{sft}}$  as the reference policy.
- 

## C Additional Experiments using Positive Synthetic Data (Section 5)

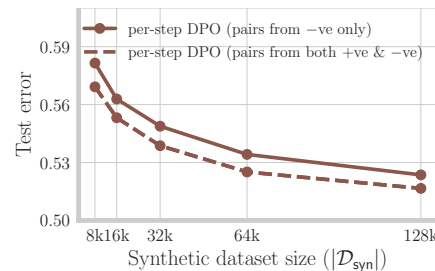
Recall from Section 5 we observed a  $2\times$  boost in sample efficiency (over  $\pi_{\text{sft}}$ ) of the question/answer pairs in the synthetic data when we cloned positive solutions sampled from  $\pi_{\text{sft}}$ . Note that one may also attribute better performance of RFT to improved coverage from multiple responses in  $\mathcal{D}_{\pi_{\text{sft}}}^+$  for each question in  $\mathcal{D}_{\text{syn}}$ . We find that even when RFT data is restricted to one solution per question, base LLMs finetuned on it outperform SFT consistently by  $> 1\%$ . In Figure 9, we plot the performance of DeepSeek-Math-7B finetuned on SFT data  $\mathcal{D}_{\text{syn}}$  and RFT data  $\mathcal{D}_{\pi_{\text{sft}}}^+$  where  $\mathcal{D}_{\pi_{\text{sft}}}^+$  has the same questions as  $\mathcal{D}_{\text{syn}}$ , and only one positive solution per question, sampled from  $\pi_{\text{sft}}$  finetuned on  $\mathcal{D}_{\text{syn}}$ . Thus, both SFT and RFT datasets are of the same size. This means that a significant portion of the  $2\times$  sample efficiency gains we observe for RFT in Figure 2(a,b) can be attributed to RFT data



**Figure 9:** RFT data with a single (self-generated) correct solution per problem outperforms SFT data (from highly-capable models) of the same size.

from  $\pi_{\text{sft}}$  being easier-to-fit, and not purely because RFT data improves coverage by finetuning on multiple solution traces per question.

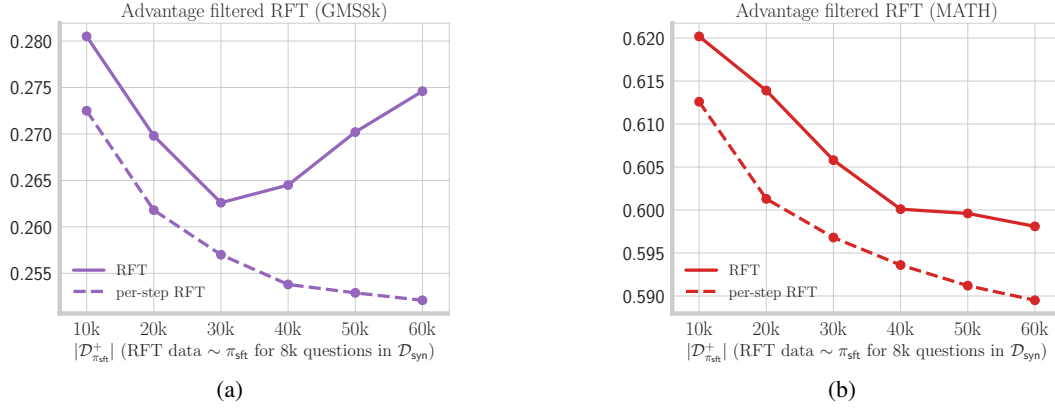
## D Additional Experiments using Negative Synthetic Data (Section 6)



**Figure 10:** On MATH, improving advantage estimates by computing advantages over both positive and negative traces sampled from  $\pi_{\text{sft}}$  improves estimation error and final performance for per-step DPO.

**Improving per-step DPO [23] with a closer approximation of advantage-weighted RL (Parts 1 + 2 in Algorithm 1).** Now, we discuss an experiment that improves the performance of per-step DPO [23] by running the full version of Algorithm 1. In particular, we add new preference pairs to the dataset of per-step DPO algorithm starting from positive samples. Recall from Section 4 and Algorithm 1, that for a problem  $x$ , with correct answer  $y$  given by SFT or RFT data, Part 1 of per-step DPO estimates the expected accuracy (Q-value) of each step in a negative rollout  $-\hat{y} \sim \pi_{\text{sft}}(\cdot | x)$  sampled from the SFT policy. For each step  $-\hat{y}_c$  the Q-value is computed conditioned on  $x$  and previous steps  $-\hat{y}_{1:c-1}$ . It then adds the triplet  $(x, y, -\hat{y}_{1:c})$  to the DPO dataset. We improve the coverage and accuracy of advantage estimates via Part 2, i.e., repeating this process for steps appearing on a positive trace  $+\hat{y} \sim \pi_{\text{sft}}(\cdot | x)$  as well. Specifically, we add  $(x, y, +\hat{y}_{1:c})$  to the DPO dataset, where the step  $+\hat{y}_c$  is the first step in the positive trace to have a low Q-value (as dictated by a relative threshold on the Q-value), which indicates that  $+\hat{y}_c$  is likely a spurious step that the SFT model generated. For individual steps that are more likely to occur in either positive or negative traces we improve coverage of alternate possible steps, and for steps that occur in both positive and negative traces, we lower the variance (and error) of the advantage estimate. In Figure 10, we compare the performance of per-step DPO runs with the datasets constructed from only negative vs. both positive and negative traces, and find that the latter has a lower test error for all sizes of  $\mathcal{D}_{\text{syn}}$ .

**Advantage filtered per-step RFT.** We ran an experiment with advantage filtering on all the steps present in both positive and negative data from the SFT policy and cloned the filtered data. For this, we cloned responses with high advantage steps from positive and negative responses sampled from the SFT policy. We filter all steps where the minimum advantage across all steps is in the bottom 50% percentile. This “per-step RFT” outperforms standard RFT (Figure 11), indicating that training on useful steps from negative data can improve beyond only training on positive data alone. While per-step RFT is worse than per-step DPO, we believe that this only further hints at the point that even using low advantage steps (that per-step RFT filters) for training, can further improve.



**Figure 11: Advantage filtered RFT:** We clone responses with high advantage steps from positive and negative responses sampled from the SFT policy. We filter all responses where the minimum advantage across all steps is in the bottom 50% percentile.

## E Proof of Theorem 6.1

We first restate the theorem statement and then provide a proof for this below. Our main goal in this theorem is to show that training with per-step DPO is equivalent to running advantage-weighted RL shown in the theoretical result.

**Theorem E.1** (Equivalence of advantage-weighted RL and DPO with per-step pairs). *The optimal policy from Equation 1 with  $\mathcal{D}_{\pi_{\text{sft}}}^{\pm}$  given by  $(\mathbf{x}, [\mathbf{y}_{1:i}, +\mathbf{y}_{i+1}], [\mathbf{y}_{1:i}, -\mathbf{y}_{i+1}])$  where the positive and negative traces share prefix  $\mathbf{y}_{1:i} \sim \pi_{\text{sft}}$ , and  $-\mathbf{y}_{i+1} \sim \pi_{\text{sft}}(\cdot | \mathbf{x}, \mathbf{y}_{1:i})$ ,  $+\mathbf{y}_{i+1} \sim \sigma(A_{\tilde{\pi}}(\mathbf{x}, \mathbf{y}_{1:i}; \cdot) - A_{\tilde{\pi}}(\mathbf{x}, \mathbf{y}_{1:i}; -\mathbf{y}_{i+1}))$ , is identical to the optima of the advantage-weighted RL objective:*

$$\max_{\pi} \mathbb{E}_{\mathbf{x} \sim p_{\text{syn}}(\mathbf{x}), \mathbf{y} \sim \pi_{\text{sft}}(\cdot | \mathbf{x})} \left[ \sum_{i=1}^L \log \pi(\mathbf{y}_i | \mathbf{x}, \mathbf{y}_{0:i-1}) \cdot \exp(A_{\tilde{\pi}}(\mathbf{x}, \mathbf{y}_{0:i-1}, \mathbf{y}_i) / \beta) \right]. \quad (5)$$

*Proof.* To prove this statement, we make the following observation: DPO [41] is equivalent to optimizing a KL-divergence penalized expected reward objective in an induced Bradley-Terry model of preferences defined by the reward function. That is, for any reward function  $r(\mathbf{x}, \mathbf{y})$  over contexts  $\mathbf{x} \sim \mu$  and responses  $\mathbf{y}$ , the optimal solution to the following RL objective:

$$\max_{\pi} \mathbb{E}_{\mathbf{x} \sim \mu, \mathbf{y} \sim \pi(\cdot | \mathbf{x})} [r(\mathbf{x}, \mathbf{y})] - \beta D_{\text{KL}}(\pi(\cdot | \mathbf{x}) || \pi_{\text{sft}}(\cdot | \mathbf{x})), \quad (6)$$

is given by the following advantage-weighted optimal policy,  $\pi^*(\cdot | \cdot)$ :

$$\forall \mathbf{x}, \mathbf{y}, \quad \pi^*(\mathbf{y} | \mathbf{x}) \propto \pi_{\text{sft}}(\mathbf{y} | \mathbf{x}) \cdot \exp\left(\frac{r(\mathbf{x}, \mathbf{y})}{\beta}\right), \quad (7)$$

and one can learn this optimal policy by running DPO on preference tuples  $(\mathbf{x}, \mathbf{y}_1, \mathbf{y}_2)$  sampled by the Bradley-Terry model [6] induced by the reward function  $r$ :

$$p(\mathbf{y}_1 \succ \mathbf{y}_2 | \mathbf{x}) = \frac{\exp(r(\mathbf{x}, \mathbf{y}_1))}{\exp(r(\mathbf{x}, \mathbf{y}_1)) + \exp(r(\mathbf{x}, \mathbf{y}_2))}. \quad (8)$$

Given this background information, we know that the optimal advantage-weighted RL policy optimizing Equation 5 is given by:

$$\forall \mathbf{x}, \mathbf{y}_{0:i}, \quad \pi(\mathbf{y}_i | \mathbf{x}, \mathbf{y}_{0:i-1}) \propto \pi_{\text{sft}}(\mathbf{y}_i | \mathbf{x}, \mathbf{y}_{0:i-1}) \cdot \exp\left(\frac{A_{\tilde{\pi}}(\mathbf{x}, \mathbf{y}_{0:i-1}, \mathbf{y}_i)}{\beta}\right). \quad (9)$$

Combining Equation 9 with the equivalence between Equation 7 and the Bradley-Terry model (Equation 8), we obtain that, if preference pairs  $(\mathbf{x}, [\mathbf{y}_{1:i}, +\mathbf{y}_{i+1}], [\mathbf{y}_{1:i}, -\mathbf{y}_{i+1}])$  were sampled from the SFT policy:  $+\mathbf{y}_{i+1} \sim \pi_{\text{sft}}(\cdot | \mathbf{x}, \mathbf{y}_{0:i})$  and  $-\mathbf{y}_{i+1} \sim \pi_{\text{sft}}(\cdot | \mathbf{x}, \mathbf{y}_{0:i})$ , and labeled according to Equation 8 applied on advantage estimates, then we obtain the desired equivalence result.  $\square$

## F Theory: Why Does Negative Data Improve Generalization?

We saw in Section 6.3 that collecting negative data from an appropriate SFT policy  $\pi_{\text{sft}}$  and an appropriate  $K$ , and training on this data improves generalization performance of the resulting model. In this section, building on the equivalence to advantage-weighted RL (Theorem 6.1), we attempt to formalize this observation into a performance guarantee. In particular, we show below that training on negative data implies that we are able to improve over the SFT policy, especially via the detection of critical steps, that attain high advantages,  $A_{\hat{\pi}}(\mathbf{x}, \mathbf{y}_{0:i-1}, \mathbf{y}_i)$ , that are otherwise not prioritized by training on positive data alone. Our theoretical result extends guarantees from the RL literature [27] comparing RL with imitation learning to show that indeed the use of RL (and hence negative data) improves over imitation alone.

**Notation and setup.** Define the policy obtained after advantage-weighted RL training as  $\pi_{\text{neg}}$ . Concretely,  $\pi_{\text{neg}}(\mathbf{y}|\mathbf{x})$  is given as:

$$\forall \mathbf{x}, \mathbf{y}_{0:j+1}, \pi_{\text{neg}}(\mathbf{y}_{j+1}|\mathbf{x}, \mathbf{y}_{0:j}) = \frac{1}{\hat{\mathbb{Z}}(\mathbf{x}, \mathbf{y}_{0:j})} \pi_{\text{sft}}(\mathbf{y}_{j+1}|\mathbf{x}, \mathbf{y}_{0:j}) \cdot \exp\left(\frac{\hat{A}_{\hat{\pi}}(\mathbf{x}, \mathbf{y}_{0:j}, \mathbf{y}_{j+1})}{\beta}\right), \quad (10)$$

where the normalization factor is given by  $\mathbb{Z}(\mathbf{x}, \mathbf{y}_{0:j})$  for each of the per-step policy distributions. This normalization factor is a critical factor that will drive the core of the theoretical result. We also note that the normalization factor in Equation 10 is derived from *empirical* advantage estimates and not from the expected estimates for the advantage value. Following Agarwal et al. [2], Kumar et al. [27], we operate in a tabular setting with a discrete (but combinatorially-large and variable-length) action space of responses, and our proof follows Theorem 4.4 in Kumar et al. [27].

**Theorem F.1** (Utility of negative data over positive data.). *Let  $\pi_{\text{neg}}$  denote the policy obtained after advantage-weighted RL (Equation 10) under an empirical distribution  $\hat{\mu}$  over prompts  $\mathbf{x}$ . Then the difference between the expected reward (i.e., task success rate),  $J(\cdot)$ , attained by  $\pi_{\text{neg}}$  and  $\pi_{\text{sft}}$  is lower-bounded as:*

$$J(\pi_{\text{neg}}) - J(\pi_{\text{sft}}) \geq \beta \cdot \mathbb{E}_{\mathbf{x}_i \sim \hat{\mu}, \mathbf{y}_{i,0:L} \sim \pi_{\text{neg}}(\cdot|\mathbf{x}_i)} \left[ \sum_{j=1}^L \log \mathbb{Z}(\mathbf{x}_i, \mathbf{y}_{i,0:j}) \right] - (\text{overestimation in } \hat{A}_{\hat{\pi}}(\mathbf{x}, \mathbf{y}_{0:i-1}, \mathbf{y}_i)) + \frac{c_0}{\sqrt{|\mathcal{D}_{\text{syn}}|}},$$

where  $\mathbb{Z}(\clubsuit, \circ)$  denotes the sum over exponentiated differences of the advantage and log likelihood values under  $\pi_{\text{sft}}$  for all possible candidate steps given a problem  $\clubsuit$  and a partial solution  $\circ$ . That is,

$$\mathbb{Z}(\clubsuit, \circ) := \sum_{\spadesuit \in \text{step candidates}} \exp\left(\frac{A_{\hat{\pi}}(\clubsuit, \circ; \spadesuit)}{\beta} + \log \pi_{\text{sft}}(\spadesuit|\clubsuit, \circ)\right),$$

$c_0$  is a constant depending upon the Rademacher complexity of the space of policies  $\pi_{\text{neg}}$  close to the SFT policy under the KL-divergence,  $|\mathcal{D}_{\text{syn}}|$  denotes the size of synthetic training prompts.

*Proof.* To begin the proof, we recall that we are operating in a discrete action space of steps  $\mathbf{y}_i$ , although this space is exponentially large. Since we operate in discrete action spaces, we invoke Lemma 5 from Agarwal et al. [2] for analyzing softmax policy gradient methods (this Lemma was extended by Lemma B.11 from Kumar et al. [27] for comparing BC vs offline RL). Denote by  $\hat{J}(\pi)$ , the reward attained by policy  $\pi$  in expectation over the empirical distribution  $\hat{\mu}$ :

$$\hat{J}(\pi_{\text{neg}}) - \hat{J}(\pi_{\text{sft}}) := \mathbb{E}_{\mathbf{x} \sim \hat{\mu}} [\widehat{V}^{\pi_{\text{neg}}}(\mathbf{x})] - \mathbb{E}_{\mathbf{x} \sim \hat{\mu}} [\widehat{V}^{\pi_{\text{sft}}}(\mathbf{x})] \geq \beta \mathbb{E}_{\mathbf{x} \sim \hat{\mu}} [\log \hat{\mathbb{Z}}(\mathbf{x})]. \quad (11)$$

We utilize the performance difference lemma [25] on the MDP induced by the set of initial problems in the training distribution  $\hat{\mu}$ , and the model induced deterministic dynamics distribution:

$$\begin{aligned}
\hat{J}(\pi_{\text{neg}}) - \hat{J}(\pi_{\text{sft}}) &= \sum_{j=1}^L \mathbb{E}_{\mathbf{x} \sim \hat{\mu}, \mathbf{y}_{0:j-1} \sim \pi_{\text{neg}}(\cdot|\mathbf{x})} \left[ \sum_{\mathbf{y}_j} \pi_{\text{neg}}(\mathbf{y}_j|\mathbf{x}, \mathbf{y}_{0:j-1}) \hat{A}_{\hat{\pi}}(\mathbf{x}, \mathbf{y}_{0:i-1}, \mathbf{y}_i) \right] \\
&= \sum_{j=1}^L \mathbb{E}_{\mathbf{x} \sim \hat{\mu}, \mathbf{y}_{0:j-1} \sim \pi_{\text{neg}}(\cdot|\mathbf{x})} \left[ \sum_{\mathbf{y}_j} \pi_{\text{neg}}(\mathbf{y}_j|\mathbf{x}, \mathbf{y}_{0:j-1}) \log \frac{\pi_{\text{neg}}(\mathbf{y}_j|\mathbf{x}, \mathbf{y}_{0:j-1}) \cdot \hat{\mathbb{Z}}(\mathbf{x}, \mathbf{y}_{0:j})}{\pi_{\text{sft}}(\mathbf{y}_j|\mathbf{x}, \mathbf{y}_{0:j-1})} \right] \\
&= \beta \cdot \sum_{j=1}^L \mathbb{E}_{\mathbf{x} \sim \hat{\mu}, \mathbf{y}_{0:j-1} \sim \pi_{\text{neg}}(\cdot|\mathbf{x})} \left[ D_{\text{KL}}(\pi_{\text{neg}}(\cdot|\mathbf{x}, \mathbf{y}_{0:j-1}), \pi_{\text{sft}}(\cdot|\mathbf{x}, \mathbf{y}_{0:j-1})) + \log \hat{\mathbb{Z}}(\mathbf{x}, \mathbf{y}_{0:j}) \right] \\
&\geq \beta \cdot \sum_{j=1}^L \mathbb{E}_{\mathbf{x} \sim \hat{\mu}, \mathbf{y}_{0:j-1} \sim \pi_{\text{neg}}(\cdot|\mathbf{x})} \left[ \log \hat{\mathbb{Z}}(\mathbf{x}, \mathbf{y}_{0:j}) \right] \\
&= \beta \cdot \mathbb{E}_{\mathbf{x} \sim \hat{\mu}, \mathbf{y}_{i,0:L} \sim \pi_{\text{neg}}(\cdot|\mathbf{x})} \left[ \sum_{j=1}^L \log \mathbb{Z}(\mathbf{x}, \mathbf{y}_{0:j}) \right].
\end{aligned}$$

Now, we can prove the desired result by accounting for the gap in the success rate between the actual distribution over  $\mathbf{x} \sim \mu$  and the empirical distribution induced by problems in the dataset  $\hat{\mu}$ :

$$\begin{aligned}
J(\pi_{\text{neg}}) - J(\pi_{\text{sft}}) &\geq \underbrace{J(\pi_{\text{neg}}) - \hat{J}(\pi_{\text{neg}})}_{(a)} + \underbrace{\hat{J}(\pi_{\text{neg}}) - \hat{J}(\pi_{\text{sft}})}_{(b)} - \underbrace{J(\pi_{\text{sft}}) - \hat{J}(\pi_{\text{sft}})}_{(c)} \\
&\geq \beta \cdot \mathbb{E}_{\mathbf{x} \sim \hat{\mu}, \mathbf{y}_{i,0:L} \sim \pi_{\text{neg}}(\cdot|\mathbf{x})} \left[ \sum_{j=1}^L \log \hat{\mathbb{Z}}(\mathbf{x}, \mathbf{y}_{0:j}) \right] - \frac{c_0}{\sqrt{|\mathcal{D}_{\text{syn}}|}} \\
&\geq \beta \cdot \mathbb{E}_{\mathbf{x} \sim \hat{\mu}, \mathbf{y}_{i,0:L} \sim \pi_{\text{neg}}(\cdot|\mathbf{x})} \left[ \sum_{j=1}^L \log \mathbb{Z}(\mathbf{x}, \mathbf{y}_{0:j}) \right] - \frac{c_0}{\sqrt{|\mathcal{D}_{\text{syn}}|}} + \Delta,
\end{aligned}$$

where  $c_0$  is a constant that depends on the Rademacher complexity of the function class of policies  $\pi_{\text{neg}}$  (coming from a uniform bound that we invoke to bound term (a), since  $\pi_{\text{neg}}$  depends on the dataset samples), and this term arises since the empirical distribution over prompts is not the same as the true population. This concentration term decays as the size of the synthetic data (number of problems) are scaled up. The term  $\Delta$  denotes the overestimation error between the estimated advantages  $\hat{A}_{\hat{\pi}}(\mathbf{x}, \mathbf{y}_{0:i-1}, \mathbf{y}_i)$  and the true advantages  $A_{\hat{\pi}}(\mathbf{x}, \mathbf{y}_{0:i-1}, \mathbf{y}_i)$ , in expectation under the distribution of the learned policy. The estimation error  $\Delta$  depends on  $\pi_{\text{sft}}$  and the value of  $K$  used if the rollout policy  $\hat{\pi}$  corresponds to the BoK( $\pi_{\text{sft}}$ ) policy. This proves the theorem.  $\square$

**Interpretation & perspectives.** Also note that the improvement in performance between  $\pi_{\text{neg}}$  and  $\pi_{\text{sft}}$  depends on the advantage estimate: if the advantage estimates are high, then this term is large, meaning that the more the fraction of high-advantage critical states, the higher the improvement. In addition, the bound also says that if the over-estimation  $\Delta$  in the advantage estimate is large, the performance improvement is small. This is perhaps expected: consider the scenario when the BoK( $\pi_{\text{sft}}$ ) policy is used to collect data, for a large  $K$ . In this scenario, the divergence between the empirical advantage estimate  $\hat{A}_{\hat{\pi}}$  and the expected estimate  $A_{\hat{\pi}}$  is likely large. In the worst case, the estimate  $\hat{A}_{\hat{\pi}}$  can arbitrarily overestimate  $A_{\hat{\pi}}$ , as it would take on a high value even if there is just *one* sequence among the  $K$  rollouts that successfully solves the problem. For example, a spurious step may be labeled incorrectly as critical in this case and training on negative data may not improve (consistent with running per-step DPO on an over-trained SFT checkpoint in Figure 8). On the other hand, if advantages are more accurate, training on negative data should improve performance.

## G Additional Related Work

**Failure modes for supervised finetuning (SFT).** First, since SFT induces an open-loop [62] next-token prediction loss, prediction errors on even a single token can snowball during inference, leading

to poor performance on the prompts appearing in the data itself [24, 45]. Second, even when an LLM has perfectly cloned the SFT data, it is prone to memorize “hard to learn” tokens [56], especially in planning and lookahead tasks [35, 36], which is critical for math reasoning. This leads to poor generalization [4, 15] and hallucination on new novel, test-time prompts [26]. In this work, we study how synthetic data methods can address these failures via: (i) maximizing likelihood on positive data generated from both the SFT policy and a stronger teacher that enjoys improved coverage over new states, and (ii) preference optimization using the negative data generated from the SFT policy.

## H Synthetic Data Generation

Prompt used for GSM8K/MATH synthetic data [29]

Please act as a professional math teacher. Your goal is to create high quality math problems to help students learn math. You will be given a math question. Please create a new question based on the Given Question and following instructions.  
To achieve the goal, you have one job.  
# Please generate a similar but new question according to the Given Question.  
You have four principles to do this. # Ensure the new question only asks for one thing, be reasonable, be based on the Given Question, and can be answered with only a number(float or integer). For example, DO NOT ask, ‘what is the amount of A, B and C?’.  
# Ensure the new question is in line with common sense of life. For example, the amount someone has or pays must be a positive number, and the number of people must be an integer.  
# Ensure your student can answer the new question without the given question. If you want to use some numbers, conditions or background in the given question, please restate them to ensure no information is omitted in your new question.  
# You only need to create the new question. Please DO NOT solve it.  
Given Question: <insert question from original dataset here>  
Your output should be in the following format:  
CREATED QUESTION: <your created question>

For GSM8K, we replace the phrase “Your goal is to create high quality math problems to help students learn math.” with “Your goal is to create high quality math *word* problems to help students learn math.”, as we found this to produce problems that were closer to GSM-style problems.

To generate the synthetic data, we used OpenAI credits worth approximately 3000 US dollars.

## I Details on Star Graph Problem

The star graph problem we study is borrowed from Bachmann and Nagarajan [4], where given a graph in the shape of a star and a query (center/end node pair), the model is asked to output the full path between the start/end nodes.

**Goal.** Bachmann and Nagarajan [4] show that  $\pi_{\text{sft}}$  minimizes SFT loss by memorizing the “hard-to-predict” node adjacent to the center, and copying the rest of the path from the input graph. This task highlights the failure of SFT at planning problems (akin to math reasoning). Thus, we use this as a case study to understand:

- when accurate advantage estimation is possible with few negative samples from the  $\pi_{\text{sft}}$  model.
- whether there are generalization benefits of advantage-weighted RL when advantage estimates are accurate
- when advantage-weighted RL can unlearn the memorized feature that causes  $\pi_{\text{sft}}$  to fail.

**SFT dataset.** The data we use for supervised fine-tuning consists of 30000 of random star graphs (see examples below) where each graph has a centre node with out degree 2. Hence, there are two paths that originate from the centre node. Each path from the center to one of the end nodes is of length 4. Each node on the path is denoted with a randomly sampled number from 0 to 20. For example, in the sample “8,3|3,10|14,13|10,11|17,14|8,17/8,13=8,17,14,13”. The graph is given by the

adjacency list “8,3|3,10|14,13|10,11|7,14|8,17/8,13”, the query is denoted by “8,13”, and the correct path is given by “8,17,14,13”.

**Test-time inference from the model.** At test time, the input to the LLM is only the graph and the query: “8,3|3,10|14,13|10,11|7,14|8,17/8,13=” and the model is expected to generate the full path from start node 8 to end node 13. When evaluating the test performance of an LLM, we calculate 0/1 accuracy averaged over 1000 test star graphs (that are different from train star graphs). The accuracy on a sample is 1 when the LLM accurately predicts all nodes in the graph.

**Failure models of the SFT model,  $\pi_{\text{sft}}$ .** A model with perfect accuracy (0 error) would be the one that has accurately learned the correct feature of backtracking the path from the end node to the start node, and then producing it in reverse. This computation is precisely what makes the adjacent token “hard-to-fit”. On the other hand, if the LLM minimizes next-token prediction loss during SFT by instead memorizing the hard-to-fit adjacent token by overfitting on the random input graph instance, at test time the accuracy would be zero. An intermediate solution that SFT model instead learns is to output a path that is adjacent to the node. At training time, it only needs to memorize which of the two possible paths to predict. Note that even this solution does not require the model to backtrack, and is thus easier to quickly learn with a few samples. This would quickly minimize the loss on all but the adjacent node, which the model memorizes as training progresses. On the test set, this model would then have 50% test accuracy. Note, that as we increase the size of the graph or the node vocabulary size it becomes easier for the model to overfit on the hard to predict adjacent token given random combinations of the input graph. Thus, we choose the vocabulary size to be 20, which is higher than what is needed to represent any input graph of this size.

Below we provide examples from degree two, path length 4, node 20 problem, where

#### Examples of 20 node path length 4 star graph problem

Example 1: 8,3|3,10|14,13|10,11|7,14|8,17/8,13=8,17,14,13  
 Example 2: 14,16|8,10|9,5|3,14|9,3|5,8/9,16=9,3,14,16  
 Example 3: 14,1|10,4|9,7|10,17|4,9|17,14/10,7=10,4,9,7  
 Example 4: 19,8|7,18|14,15|15,7|14,19|8,10/14,10=14,19,8,10  
 Example 5: 1,6|10,1|6,12|10,17|17,18|18,5/10,12=10,1,6,12

**SFT Training details.** We finetune a pretrained GPT-2 model with 125 million parameters. We train with a batch size of 128, Adam without any weight decay, and a constant learning rate of  $1e-5$ .

**Advantage estimation and per-step DPO training equivalent to advantage-weighted RL.** For a sample from  $\pi_{\text{sft}}$ , we estimate the advantage of each step by sampling 5 rollouts conditioned on the subsequence up till the step. We then pair subsequences with shared prefix,  $y_{1:i}$  differing in the last step  $+y_{i+1}$  vs.  $-y_{i+1}$ , where the former is the one with the highest estimated advantage and the latter is the one with the lowest estimated advantage. Note that this preference pair construction, closely approximates the preference pair distribution in Theorem 6.1, which would imply that the DPO objective being optimized closely approximates advantage weighted RL in Equation 4.

Given these pairs for a batch of star graph problems in SFT data, we update the model with a single gradient step on the DPO objective in Equation 1. In the next iteration, advantage is estimated and pairs are constructed on a fresh batch of star graphs. We set  $\beta = 0.1$  in the DPO objective and use the same batch size (one preference pair per star graph). Starting from an SFT checkpoint we train in the above manner for at least 200 iterations. The SFT model is trained for over 600 iterations.

## J Implementation Details

**Datasets and pretrained LLMs.** We run all our experiments on GSM8K and MATH datasets. Each dataset has about 7.5k training examples. The GSM8K has about 1.3k and MATH has 5k test examples. We conduct experiments with DeepSeek-Math-7B pretrained LLM and LLaMA2-7B, both of which have pretrained weights publicly available on Huggingface.

**Details for SFT/RFT training.** For our positive data scaling results, the SFT model is trained for 5 epochs with a learning rate of  $1e-5$ , and a batch size of 64 for all sizes of  $\mathcal{D}_{\text{syn}}$ . We use a holdout validation set to choose the checkpoint and report the performance of the best performing checkpoint across the five epochs. To generate RFT data we only train the SFT model for 2 epochs (under-trained checkpoint). For each question we sample  $M = 100$  times, with a temperature of 0.7 and following

Yuan et al. [67] we retain at most 4 most diverse (based on edit distance) and correct completions. This is for our results in Figure 2(a,b). For Figure 2(c), we sample more than 4 correct solutions and keep sampling responses until we have a dataset of size 10k, 20k, ..., 50k, when questions are given by the  $\mathcal{D}_{\text{syn}}$  of size 8k and 16k. For our experiment on the RFT dataset with purposely inserted spurious steps, as we describe in the main paper, we obtain spurious steps by computing which intermediate steps in a negative response lead to most incorrect solutions and randomly insert this in between reasoning steps for a positive solution in RFT data. See examples below.

**Details for per-step DPO training.** Training data for DPO is generated in the procedure outlined in [23]. The value of  $K$  is 5 for the  $\text{BoK}(\pi_{\text{sft}})$  policy used to estimate the advantage, and a single rollout from the best-of-K policy is used to estimate the advantage of each step in “rejected” response, where the “chosen” and “rejected” pairs for each question are generated from the RFT data, using the method in [38]. We train per-step DPO with  $\beta = 0.3$  for all  $\mathcal{D}_{\text{syn}}$  from GSM8K, and  $\beta = 0.1$  for all  $\mathcal{D}_{\text{syn}}$  from MATH. We train for three epochs, with a learning rate of  $1e-6$  and a batch size of 64 prompts. We use the rmsprop optimizer and a cosine learning rate decay with a warmup ratio of 0.1.

Overall, to replicate our experiments, we need approximately 1000 hours of compute on H100 clusters with 8 GPUs.

#### RFT data with spurious steps highlighted

##### Prompt:

Karen’s work tote bag is twice the weight of her husband Kevin’s briefcase when the briefcase is empty. When Kevin puts his laptop and work papers in his briefcase, it is twice the weight of Karen’s tote. Kevin’s work papers are a sixth of the weight of the contents of his full briefcase. If Karen’s tote weighs 8 pounds, how many more pounds does Kevin’s laptop weigh than Karen’s tote?

##### Response:

If Karen’s tote weighs 8 pounds, and it is twice the weight of Kevin’s empty briefcase, then Kevin’s empty briefcase weighs  $8/2 = \ll 8/2=4 \gg 4$  pounds. **In words, “Karen’s tote weighs 8 pounds” is “ $T = 8$ ”, so thus Keven is 2 pounds.** When Kevin puts his laptop and work papers in his briefcase, it is twice the weight of Karen’s tote at 8 pounds, so the briefcase now weighs  $8+4 = \ll 8+4=12 \gg 12$  pounds. Kevin’s work papers are a sixth of the weight of the contents of his full briefcase, which weighs 12 pounds, so the work papers weigh  $12/6 = \ll 12/6=2 \gg 2$  pounds. If Kevin’s laptop and work papers weigh 12 pounds in total, and the work papers weigh 2 pounds, then the laptop weighs  $12-2 = \ll 12-2=10 \gg 10$  pounds. Thus, Kevin’s laptop is  $10-8 = \ll 10-8=2 \gg 2$  pounds heavier than Karen’s tote. The answer is 2.

#### RFT data with spurious steps highlighted

##### Prompt:

Carl has a jar full of marbles. He takes out 12 marbles to play a game with, but he accidentally drops them and  $1/2$  the marbles get lost. So Carl takes out 10 more marbles. While he is playing his game his mother comes home with another bag of marbles for him, which has 25 marbles in it. If Carl can’t find his lost marbles, then how many marbles is he going to put in the jar after he plays his game, from both his original marbles and the new ones?

##### Response:

First find out how many marbles Carl loses:  $12 \text{ marbles} * 1/2 = \ll 12*1/2=6 \gg 6$  marbles Then add that number to the number of marbles Carl takes out of the jar to play his game to find out how many total marbles he takes out:  $6 \text{ marbles} + 10 \text{ marbles} = \ll 6+10=16 \gg 16$  marble. **Now add that number to the number of marbles in the jar to find out how many marbles are left:  $16 \text{ marbles} + 12 \text{ marbles} = \ll 16+12=28 \gg 28$  marbles.** Then add in the marbles Carl gets from his mother to find out how many marbles he has in total:  $16 \text{ marbles} + 25 \text{ marbles} = \ll 16+25=41 \gg 41$  marbles. The answer is 41.

## K Negative Data Identifies Spurious Steps with Advantage Estimates

In this section, we present several examples from the GSM8K and MATH synthetic datasets of model-generated (from  $\pi_{\text{sft}}$ ) positive data with spurious steps. Then, we show how computing

advantage values accurately identify the spurious steps. We highlight the steps that were deemed to be low or high advantage based on the estimated advantage value. The darker the intensity of the color, the lower the advantage of that step. With this qualitative probe, we see that the steps that are darkest are indeed either incorrect or correct but involve a reasoning step that is irrelevant for solving the problem at hand.

#### Spurious step identified by negative data

**Question:** Ken created a care package to send to his brother, who was away at boarding school. Ken placed a box on a scale, and then he poured into the box enough jelly beans to bring the weight to 2 pounds. Then, he added enough brownies to cause the weight to triple. Next, he added another 2 pounds of jelly beans. And finally, he added enough gummy worms to double the weight once again. What was the final weight of the box of goodies, in pounds?

#### Positive response with identified spurious step:

When the weight was originally on the scale, the box weighed 0 pounds. After Ken added the jelly beans to the box, its weight was 2 pounds - 0 pounds =  $\langle 2-0=2 \rangle$  2 pounds. Then Ken added the brownies to the box, which caused the box to weigh 2 pounds \* 3 =  $\langle 2*3=6 \rangle$  6 pounds. The next step was to add 2 pounds of jelly beans, which caused the box to weigh 6 pounds + 2 pounds =  $\langle 6+2=8 \rangle$  8 pounds. First, we need to figure out what the weight of the box was when Ken put the jelly beans in it, 2 pounds of jelly beans was enough to triple the weight, so the box was at 2 pounds / 3 =  $\langle 2/3=.6666666666666667 \rangle$ . 6666666666666667 pounds. Finally, Ken added the gummy worms, causing the box to weigh 8 pounds \* 2 =  $\langle 8*2=16 \rangle$  16 pounds. The answer is 16

#### Spurious step identified by negative data

**Question:** Betty is saving money for a new wallet which costs \$100. Betty has only half of the money she needs. Her parents decided to give her \$15 for that purpose, and her grandparents twice as much as her parents. How much more money does Betty need to buy the wallet?

#### Positive response with identified spurious step:

Betty's parents gave her 15, so her grandparents gave her  $2 * 15 = \langle 15*2=30 \rangle$  30. In total, Betty received  $30 + 15 = \langle 30+15=45 \rangle$  45. The total amount of money Betty needs is  $100 * 2 = \langle 100*2=200 \rangle$  200. Betty needs  $100 / 2 = \langle 100/2=50 \rangle$  50 in total. Betty still needs to save  $50 - 45 = \langle 50-45=5 \rangle$  5. The answer is 5

#### Spurious step identified by negative data

**Question:** To make pizza, together with other ingredients, Kimber needs 10 cups of water, 16 cups of flour, and 1/2 times as many teaspoons of salt as the number of cups of flour. Calculate the combined total number of cups of water, flour, and teaspoons of salt that she needs to make the pizza.

#### Positive response with identified spurious step:

Since 1 cup is equal to 16 teaspoons, then Kimber needs  $128/16 = \langle 128/16=8 \rangle$  8 cups of salt. So in total, she needs  $10+16+8 = \langle 10+16+8=34 \rangle$  34 cups of water, flour, and salt. In cups, the amount of salt is  $128/6 = \langle 128/6=21.33 \rangle$  21.33 cups. The answer is 34

#### Spurious step identified by negative data

**Question:** Eight friends compete in a pie eating contest. The chart displays the number of pies each person consumed. Emily ate the most pies, and Mike ate the least. How many more pies did Emily eat compared to Mike?

```
[asy]\ndefaultpen(linewidth(1pt)+fontsize(10pt));
\npair[] yaxis = new pair[10];\nfor( int i = 0 ; i < 10 ; ++i )
{\n\nyaxis[i] = (0,i);\n\nndraw(yaxis[i]--yaxis[i]+(17,0));\n}
\ndraw((0,0)--(0,9));\ndraw((17,9)--(17,0));
\n\nfill((1,0)--(1,5)--(2,5)--(2,0)--cycle,greyscale(.5));
\nfill((3,0)--(3,8)--(4,8)--(4,0)--cycle,greyscale(.5));
\nfill((5,0)--(5,3)--(6,3)--(6,0)--cycle,greyscale(.5));
\n\nlabel("0",yaxis[0],W);
\nlabel("1",yaxis[1],W);\nlabel("2",yaxis[2],W);
\nlabel("3",yaxis[3],W);\nlabel("4",yaxis[4],W);
\nlabel("5",yaxis[5],W);\nlabel("6",yaxis[6],W);
\nlabel("7",yaxis[7],W);\nlabel("8",yaxis[8],W);
\nlabel("Friends/Participants", (8.5,0),S);
\nlabel("Results of a Pie Eating Contest", (8.5,9),N);
\nlabel(rotate(90)*"$\\frac{1}{10}$ of Pies Eaten", (-1,4.5),W);[/asy]
```

#### Positive response with identified spurious step:

To solve this problem, we need to identify the number of pies Emily and Mike consumed, and then find the difference between their consumption. From the chart provided: - Emily: Consumed 8 pies. - Mike: Consumed 1 pie. To determine how many more pies Emily ate compared to Mike, we first identify how many pies each person ate. - Emily ate 5 pies. - Mike ate 1 pie. Now, calculating the difference between Emily's and Mike's consumption: [ Difference = Emily's consumption - Mike's consumption = 8 - 1 = 7. ] Therefore, Emily ate 7 more pies than Mike. The answer is 7.

#### Spurious step identified by negative data

**Question:** When the base-16 number  $1AB3_{16}$  is written in base 2, how many base-2 digits (bits) does it have?

#### Positive response with identified spurious step:

To solve this problem, we first need to convert the base-16 number  $1AB3_{16}$  to base-2. Step 1: Convert each digit of the number from base-16 to base-2. The number  $1AB3_{16}$  consists of the digits 1, A, B, and 3. Step 2: Convert each digit individually to base-2 representation. -  $1_{16} = 0001_2$  -  $A_{16} = 1010_2$  -  $B_{16} = 1011_2$  -  $3_{16} = 0011_2$  Step 3: Combine the base-2 representations of each digit. Concatenating the individual base-2 representations, we get:  $1AB3_{16} = 0001101010110011_2$  Step 4: Count the total number of digits in the base-2 representation. The base-2 representation of  $1AB3_{16}$  is:  $0001101010110011_2$  The base-16 number ( $1AB3_{16}$ ) can be written in base 10 (decimal) as:  $1 \times 16^3 + A \times 16^2 + B \times 16 + 3$  where (A) and (B) are decimal digits. This sequence of digits has a total of 16 digits. The answer is boxed{16}

## NeurIPS Paper Checklist

### 1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: Claims from the introduction and abstract are reflected in the paper.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

### 2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: Appendix [A](#).

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

### 3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Theorem statements are included in the main paper and detailed proofs are included in Appendix F and Appendix E.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

#### 4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: Experimental details are included in the main paper and in Appendix H and Appendix J

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
  - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
  - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
  - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
  - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

#### 5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [No]

Justification: While the code and dataset generated are not yet ready for anonymous open sourcing, we plan to open-source the code with appropriate licensing and the generated synthetic data with the updated version of the paper.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

## 6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: Experimental details are included in the main paper and in Appendix H and Appendix J

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

## 7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: For experiments that were bottlenecked by computational resources, we were unable to run multiple experiments. However, since we perform experiments on well-established benchmarks, we could make comparisons with existing work to understand the significance of our performance. And hence, whenever needed we make sure to discuss observed performances and significance with existing work. Moreover, since we run experiments at different scales and dataset sizes, the trends observed implicitly normalize for error bars.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

## 8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [\[Yes\]](#)

Justification: We have added computational resources needed in the Appendix. Overall, to replicate our experiments, we need approximately 1000 hours of compute on H100 clusters with 8 GPUs. To generate our synthetic data, we need approximately 3000 USD worth of credits with the current pricing of GPT-4 model.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

## 9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [\[Yes\]](#)

Justification: Our work is to study reasoning in large language models and it doesn't violate any code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

## 10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We discuss limitations of our work in Appendix A. Since our work is about empirical and theoretical study of reasoning in LLMs, we do not believe that there are any direct negative implications.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

## 11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: Our work is not focused towards releasing novel pretrained language models, image generators or scraped datasets. For the synthetic dataset obtained for reasoning tasks by prompting large models, we believe there are no risks of misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

## 12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We make sure to cite relevant and related work. Our code base builds on top of open-source code repositories that allow free access for research purposes. When we open-source our code with the updated version, we will make sure that all the credits are appropriately attributed.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, [paperswithcode.com/datasets](https://paperswithcode.com/datasets) has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

### 13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: We do not release any new assets. For the dataset generated in our work, we include prompts mentioned in Appendix H.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

### 14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: We do not do any crowdsourcing experiments.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

### 15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: We do not do any experiments with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.