

EXPLORE DATA LEFT BEHIND IN REINFORCEMENT LEARNING FOR REASONING LANGUAGE MODELS

Chenxi Liu^{1,2,*}, Junjie Liang², Yuqi Jia^{2,3}, Bochuan Cao⁴, Yang Bai², Heng Huang¹, Xun Chen²

¹University of Maryland, College Park ²ByteDance ³Duke University

⁴Pennsylvania State University

ABSTRACT

Reinforcement Learning with Verifiable Rewards (RLVR) has emerged as an effective approach for improving the reasoning abilities of large language models (LLMs). The Group Relative Policy Optimization (GRPO) family has demonstrated strong performance in training LLMs with RLVR. However, as models train longer and scale larger, more training prompts become residual prompts—those with zero-variance rewards that provide no training signal. Consequently, fewer prompts contribute to training, reducing diversity and hindering effectiveness. To fully exploit these residual prompts, we propose the **Explore Residual Prompts in Policy Optimization (ERPO)** framework, which encourages exploration on residual prompts and reactivates their training signals. ERPO maintains a history tracker for each prompt and adaptively increases the sampling temperature for residual prompts that previously produced all-correct responses. This encourages the model to generate more diverse reasoning traces, introducing incorrect responses that revive training signals. Empirical results on the Qwen2.5 series demonstrate that ERPO consistently surpasses strong baselines across multiple mathematical reasoning benchmarks. Source code is available at <https://github.com/DawnLIU35/ERPO>

1 INTRODUCTION

Large language models (LLMs) have become the foundation of modern artificial intelligence, exhibiting strong performance across domains such as mathematics, programming, and scientific problem solving (Team et al., 2023; Guo et al., 2025; Yang et al., 2025a; Zheng et al., 2025d; Chen et al., 2024b; 2023; 2024c; Shirkavand et al., 2025a;b; Chen et al., 2025b). A central factor behind these advancements is their capacity for extended reasoning, where models construct coherent, multi-step chains of thought to address complex tasks (Wei et al., 2022; Yao et al., 2023; Muennighoff et al., 2025). Reinforcement learning (RL) has emerged as a key paradigm for strengthening this capability, enabling LLMs to refine their responses through interaction-driven feedback and alignment with verifiable signals or human preferences (Schulman et al., 2017; Ouyang et al., 2022; Rafailov et al., 2023; Chen et al., 2024a). In particular, reinforcement learning with verifiable rewards (RLVR) has proven especially effective, as it leverages tasks with automatically checkable outcomes to provide reliable supervision for scaling reasoning abilities (Shao et al., 2024; Guo et al., 2025; Yang et al., 2025a; Liu et al., 2025).

Among recent advances in reinforcement learning for LLMs, Group Relative Policy Optimization (GRPO) has emerged as a widely adopted RLVR framework (Shao et al., 2024; Guo et al., 2025). Building on this foundation, subsequent research has sought to address key issues of GRPO, including entropy collapse, reward noise, and training instability (Yu et al., 2025; Cui et al., 2025; Zheng et al., 2025a). Furthermore, as an on-policy algorithm, GRPO has motivated efforts to develop more effective sampling strategies beyond basic random decoding (Xu et al., 2025; Zheng et al., 2025c; Hou et al., 2025).

In this work, we identify a limitation shared by the GRPO family of algorithms: as training steps and model size increase, more training prompts become residual prompts that no longer provide training

*Work done during an internship at ByteDance.

Table 1: Proportion of prompts with all-correct responses under different sampling temperatures and model scales. The proportion increases with RL training process and larger model sizes, leaving more residual prompts, thereby reducing diversity and wasting valuable training signals.

	$T = 1.0$	$T = 1.1$	$T = 1.2$
Qwen2.5-3B	0%	–	–
Qwen2.5-3B + DAPO	8.7%	6.2%	2.8%
Qwen2.5-7B	0%	–	–
Qwen2.5-7B + DAPO	21.3%	15.5%	5.5%
Qwen2.5-32B	0%	–	–
Qwen2.5-32B + DAPO	74.8%	62.1%	34.8%

signals yet still contain valuable information that can benefit model performance. Residual prompts are those that initially provide effective training signals at the beginning of training but eventually provide zero training signal or are filtered out by the RL algorithms as the well-trained policy generates all-correct responses for them. This reduces training diversity over time and ultimately hinders further improvement through RL. Furthermore, residual prompts retain learning potential that can be leveraged to further improve model performance, as they help the model retain acquired abilities and may yield novel reasoning traces. Moreover, residual prompts are not necessarily robust—small perturbations, such as increasing the sampling temperature, can easily induce errors. Table 1 reports the proportion of residual prompts with all-correct responses in the training data under different sampling temperatures and model scales.

To better exploit the residual prompts left behind during training, we propose the **Explore Residual Prompts in Policy Optimization (ERPO)** framework. ERPO introduces a novel sampling strategy that maintains a history tracker for each prompt and adaptively increases the sampling temperature for residual prompts that have previously produced all-correct responses. Specifically, ERPO records how many times a model generates all-correct responses for each prompt, and the sampling temperature is determined by this count. The more frequently a prompt yields all-correct responses, the higher the sampling temperature assigned to it, thereby encouraging greater exploration. As shown in Table 1, increasing the sampling temperature enables the model to explore more diverse reasoning traces and generate incorrect responses, which reactivates the training signal and alleviates the collapse of prompt diversity.

Overall, **our contributions** can be summarized as follows:

- We identify a key limitation of the GRPO family: residual prompts accumulate as training progresses and models scale, leading to reduced training diversity and the loss of valuable training signals from residual prompts.
- We propose the ERPO framework, which encourages models to adaptively explore residual data and recover their learning potential. ERPO maintains a history tracker for each prompt and adaptively increases the sampling temperature for residual prompts.
- Extensive experiments on several math reasoning benchmarks demonstrate the effectiveness of ERPO in both average and majority-vote evaluations, with particularly strong improvements on data that are likely not contaminated, such as AIME2025.

2 RELATED WORK

Reinforcement learning for LLM reasoning. Reinforcement learning (RL) has become a central approach for enhancing the reasoning abilities of large language models (LLMs) in domains such as mathematics, programming, and problem solving (Dubey et al., 2024; Zhou et al., 2025). Early general-purpose algorithms like Proximal Policy Optimization (PPO) provided a practical framework for fine-tuning LLMs through sampled rollouts and reward feedback (Schulman et al., 2015; 2017). More recently, RLVR methods such as Group Relative Policy Optimization (GRPO) have emerged as effective alternatives to PPO, removing the critic model while maintaining strong performance on reasoning benchmarks (Guo et al., 2025; Shao et al., 2024). Several extensions have been

proposed to address the limitations of GRPO: Cui et al. (2025); Wang et al. (2025); Cheng et al. (2025); Zheng et al. (2025c) mitigates the entropy collapse problem during training; Zheng et al. (2025a); Yang et al. (2025b) aims to stabilize the optimization process, and DAPO (Yu et al., 2025) tackles both issues while filtering noisy rewards for training data. However, all these methods obtain no training signal from residual prompts, thereby missing valuable information during training. To address this limitation, ERPO reactivates the training signal of residual prompts and learns useful information from them.

Data Sampling Strategies. The outputs of LLMs rely heavily on data sampling strategies to balance diversity and quality. Common strategies include greedy search, beam search, and various random sampling techniques such as top-k and top-p (Zhao et al., 2023; Minaee et al., 2024). In RLVR, the model generates on-policy responses and assigns them verifiable rewards during training. Basic random decoding is widely used in RLVR algorithms such as GRPO and DAPO (Guo et al., 2025; Yu et al., 2025). Beyond this, several works explore alternative sampling strategies. Hou et al. (2025) leverages tree search to find correct responses with higher probability. Zheng et al. (2025c) forks responses at high-entropy tokens. Shrivastava et al. (2025) dynamically allocates additional training resources to harder problems based on real-time difficulty estimates. Xu et al. (2025) selects a subset of responses to maximize reward variation. Zheng et al. (2025b) predicts and skips uninformative prompts using reward training dynamics. Zhang et al. (2025) progressively exposes the model to increasingly challenging samples. Nevertheless, none of these methods are specifically designed to leverage information from residual prompts.

3 PRELIMINARIES

Notation We define an autoregressive language model parameterized by θ as a policy π_θ . Let q denote a query and $\mathcal{D}$ the query set. For a response o to query q , its likelihood under π_θ is expressed as

$$\pi_\theta(o | q) = \prod_{t=1}^{|o|} \pi_\theta(o_{i,t} | q, o_{i,<t}), \quad (1)$$

where $|o|$ is the number of tokens in o .

Group Relative Policy Optimization (GRPO) (Shao et al., 2024; Guo et al., 2025) has shown strong effectiveness for fine-tuning LLMs. Unlike traditional approaches that rely on a critic network of comparable size to the policy, GRPO estimates the baseline directly from group-level rewards. For a specific question-answer pair (q, a) , the behavior policy $\pi_{\theta_{\text{old}}}$ samples a group of G individual responses $\{o_i\}_{i=1}^G$. Then, the advantage of the i -th response is calculated by normalizing the group-level rewards $\{R_i\}_{i=1}^G$:

$$\hat{A}_{i,t} = \frac{r_i - \text{mean}(\{R_i\}_{i=1}^G)}{\text{std}(\{R_i\}_{i=1}^G)}. \quad (2)$$

Building on the group-normalized advantages, GRPO optimizes the policy with a clipped objective that stabilizes updates and a KL regularization term that constrains divergence from the reference model. The objective is defined as:

$$\begin{aligned} \mathcal{J}_{\text{GRPO}}(\theta) = & \mathbb{E}_{(q,a) \sim \mathcal{D}, \{o_i\}_{i=1}^G \sim \pi_{\theta_{\text{old}}}(\cdot | q)} \\ & \left[\frac{1}{G} \sum_{i=1}^G \frac{1}{|o_i|} \sum_{t=1}^{|o_i|} \left(\min \left(r_{i,t}(\theta) \hat{A}_{i,t}, \text{clip} \left(r_{i,t}(\theta), 1 - \varepsilon, 1 + \varepsilon \right) \hat{A}_{i,t} \right) \right. \right. \\ & \left. \left. - \beta D_{\text{KL}}(\pi_\theta || \pi_{\text{ref}}) \right) \right], \end{aligned} \quad (3)$$

where $r_{i,t}(\theta)$ is the importance ratio between the old and new policy:

$$r_{i,t}(\theta) = \frac{\pi_\theta(o_{i,t} | q, o_{i,<t})}{\pi_{\theta_{\text{old}}}(o_{i,t} | q, o_{i,<t})}. \quad (4)$$

Decoupled Clip and Dynamic sAmpling Policy Optimization (DAPO) (Yu et al., 2025) introduces four key improvements: Clip-Higher promotes output diversity and mitigates entropy col-

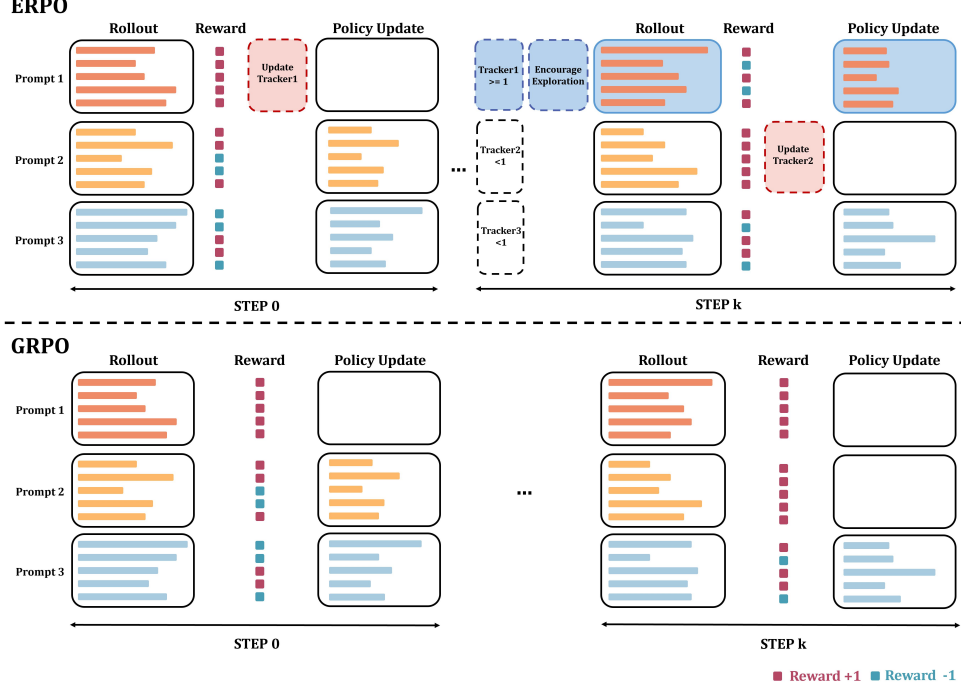


Figure 1: Comparison between ERPO and GRPO. During RL training, the policy gradually learns from the training data, resulting in more residual prompts with all-correct responses. In GRPO, residual prompts yield zero-variance rewards and thus provide no training signal for policy updates, reducing the effectiveness of training data. In contrast, ERPO maintains a tracker for each prompt to record the number of times it produces all-correct responses, and adaptively encourages exploration on residual prompts to trigger incorrect responses and reactivate the training signal.

lapse; Dynamic Sampling is designed to enhance training efficiency and stability; Token-Level Policy Gradient Loss plays a critical role in handling long chain-of-thought reasoning; and Overlong Reward Shaping reduces reward noise while stabilizing optimization. Building on these components, DAPO optimizes the policy with the following objective:

$$\begin{aligned}
 \mathcal{J}_{\text{DAPO}}(\theta) = & \mathbb{E}_{(q,a) \sim \mathcal{D}, \{o_i\}_{i=1}^G \sim \pi_{\theta_{\text{old}}}(\cdot|q)} \\
 & \left[\frac{1}{\sum_{i=1}^G |o_i|} \sum_{i=1}^G \sum_{t=1}^{|o_i|} \min \left(r_{i,t}(\theta) \hat{A}_{i,t}, \text{clip} \left(r_{i,t}(\theta), 1 - \varepsilon_{\text{low}}, 1 + \varepsilon_{\text{high}} \right) \hat{A}_{i,t} \right) \right] \quad (5) \\
 \text{s.t. } & 0 < \left| \{o_i \mid \text{is_equivalent}(a, o_i)\} \right| < G,
 \end{aligned}$$

where a is the ground-truth answer of query q .

4 METHODOLOGY

ERPO is proposed to leverage the training information contained in residual prompts to improve reinforcement learning for reasoning language models. Section 4.1 introduces the idea of reactivating the training signal from residual prompts that are otherwise discarded during RL training. Section 4.2 describes how ERPO predicts whether a prompt is residual and adaptively modifies the sampling strategy. It further explains how ERPO adjusts the sampling temperature to encourage different levels of exploration based on a history tracker.

4.1 REACTIVATE TRAINING SIGNAL

Current RLVR algorithms usually discard residual prompts that contain all-correct responses by assigning them zero advantage (Guo et al., 2025) or by directly filtering them out from the training batch (Yu et al., 2025). Consequently, available training prompts continually decreases with ongoing RL training and increasing model size, leading to a gradual reduction in both the size and diversity of the training dataset. As shown in Table 1, larger models with longer training generally produce more residual prompts that yield all-correct responses. Furthermore, as training progresses and the policy evolves, new reasoning traces and directions may be generated for the residual prompts, which can help the model learn more diverse reasoning patterns. In addition, training on residual prompts can reinforce the reasoning abilities the model has already acquired.

To leverage the training signal from residual prompts that are left behind by current RL algorithms, we propose a simple method to reactivate them. For residual prompts with all-correct responses, we replace the zero advantage with a small positive advantage by introducing a pseudo-negative reward into the advantage computation. The new Reactivated Advantage (RA) for prompts with all-correct responses is:

$$\hat{R}A_{i,t} = \frac{r_i - \text{mean}(\{R_i^+\}_{i=1}^G \cup \{R^-\})}{\text{std}(\{R_i^+\}_{i=1}^G \cup \{R^-\})}. \quad (6)$$

where R^+ is the reward for correct responses and R^- is the reward for incorrect responses. Using RA, residual prompts with all-correct responses still retain a small positive advantage, providing a valid training signal instead of being discarded by the RL algorithm.

4.2 EXPLORE RESIDUAL PROMPTS IN POLICY OPTIMIZATION

Although using the reactivated advantage can force the model to learn information from residual prompts, residual prompts will dominate the training along the training process and model scales up, leaving less negative feedback and impede the training effectiveness. (Chen et al., 2025a; Xiong et al., 2025). Furthermore, the model may suffer from the imbalance between exploration and exploitation, overfitting to a narrow exploration space.

To address this limitation, we propose to adaptively encourage exploration on residual prompts by controlling their sampling temperature. As shown in Table 1, a higher sampling temperature can trigger incorrect responses, thereby reactivating the training signals of residual prompts. Note that training data typically exhibit strong temporal correlations across epochs (Zheng et al., 2022), meaning that a prompt producing all-correct responses in the current epoch is likely to do so again in the following epoch (Zheng et al., 2025b). Thus, we can maintain a history tracker H_i to track how many times the policy generates all-correct responses for a prompt q_i :

$$H_i^{(0)} = 0, \quad H_i^{(t)} = H_i^{(t-1)} + \mathbf{1}_{q_i \text{ has all-correct responses at step } t} \quad (7)$$

Then H_i is used to determine whether we should assign a larger sampling temperature to prompt q_i . If H_i is greater than 0, it means that prompt q_i is already easy for the policy to generate all-correct responses, and it is very likely to provide no training signals the next time the policy samples it. Therefore, we assign a larger sampling temperature to prompts with $H_i > 0$ to encourage more exploration of their reasoning traces and to reactivate the training signal by triggering incorrect responses.

Since the robustness of prompts varies, some residual prompts require only a marginal increase in sampling temperature to induce incorrect responses, whereas others necessitate substantially larger adjustments. At the same time, it is essential to preserve the benefits of on-policy learning by constraining distributional shifts within a reasonable range to ensure stable and effective training. Assigning excessively large sampling temperatures is particularly detrimental for prompts with lower robustness. This trade-off highlights the difficulty of selecting a single, unified sampling temperature that can consistently induce incorrect responses, enhance exploration, and maintain a manageable distribution shift across all residual prompts. Therefore, ERPO introduces a prompt-adaptive adjustment of the sampling temperature:

$$T_i^{(t)} = \min(T_0 + T_s \cdot H_i^{(t)}, T_{max}) \quad (8)$$

where T_0 , T_{max} , and T_s are hyperparameters representing the initial temperature, maximum temperature, and temperature step size, respectively. In this way, ERPO gradually increases the sampling

Algorithm 1 ERPO framework**Input:** Policy π_θ , reward model R ,Prompt set $\mathcal{D} = \{q_i\}_{i=1}^N$, history tracker $\{H_i\}_{i=1}^N$ (init. $H_i^{(0)}=0$),Rollouts per prompt n , temperatures $(T_0, T_{\max})$, step size T_s , steps K **Output:** Updated policy $\pi_{\theta_{\text{updated}}}$

```

1: for  $t = 1, 2, \dots, K$  do
2:   Sample a mini-batch  $\mathbf{q} \subseteq \mathcal{D}$ 
3:   for each  $q_i \in \mathbf{q}$  do
4:      $T_i^{(t)} \leftarrow \min(T_0 + T_s \cdot H_i^{(t-1)}, T_{\max})$ 
5:     Sample  $n$  rollouts  $\mathbf{o}_i = (o_1, \dots, o_n)$  for  $q_i$  using  $\pi_\theta$  at temperature  $T_i^{(t)}$ 
6:     Compute rewards  $\mathbf{r}_i = (r_1, \dots, r_n)$  with  $R$ ;  $\text{acc} \leftarrow \mathbf{1}_{\text{all } o_j \text{ correct}}$ 
7:      $H_i^{(t)} \leftarrow H_i^{(t-1)} + \mathbf{1}_{\text{acc}=1}$ 
8:   end for
9:   Update  $\pi_\theta$  using an RL algorithm with data  $\mathcal{B} = \{(\mathbf{q}, \mathbf{o}, \mathbf{r})\}$ 
10: end for
11: return  $\pi_{\theta_{\text{updated}}} \leftarrow \pi_\theta$ 

```

temperature of residual prompts until the policy generates incorrect responses. This enables ERPO to strike a balance between reactivating training signals, encouraging exploration, and maintaining a reasonable distribution shift. In general, our ERPO framework can be summarized in Algorithm 1:

5 EXPERIMENTS

In this section, we first outline the implementation details, including training details and evaluation. We then present the main results, comparing ERPO against baseline approaches across several math reasoning benchmarks. Finally, we provide additional experimental results to support further analysis.

5.1 IMPLEMENTATION DETAILS

Training details: Following recent studies (Zheng et al., 2025c; Cheng et al., 2025; Shao et al., 2025) that apply RLVR to train LMMs for math reasoning tasks, we adopt Qwen2.5-3B and Qwen2.5-7B (Qwen et al., 2025) as our backbone models. Consistent with prior work (Yu et al., 2025; Cheng et al., 2025; Cui et al., 2025), we use the DAPO-Math-17K dataset (Yu et al., 2025) for training. To achieve strong performance, we adopt the DAPO algorithm (Yu et al., 2025). Prior works (Yu et al., 2025; Cheng et al., 2025) has demonstrated its superior effectiveness and stability over vanilla GRPO, and we employ it both as the baseline and as the optimization method for ERPO. The learning rate is set to 1×10^{-6} with a linear warm-up over 10 rollout steps. For rollout, we use a prompt batch size of 512, sampling 16 responses per prompt. During training, the mini-batch size is set to 512, resulting in 16 gradient updates per rollout step. The initial rollout temperature T_0 is set to 1.0. The temperature increment step T_s is set to 0.02 for Qwen2.5-3B and 0.05 for Qwen2.5-7B, while the maximum rollout temperature $T_{\max}$ is set to 1.2 for Qwen2.5-3B and 1.4 for Qwen2.5-7B, respectively. Rewards are assigned as 1 for correct responses and -1 otherwise. All experiments are conducted using the verl framework (Sheng et al., 2024). More details can be found in the Appendix.

Evaluation: We evaluate our models on AIME 2025/2024, AMC 2023, and MATH500 (Hendrycks et al., 2021), using a rollout temperature of 1.0 and top- p sampling with $p = 0.7$. For AIME and AMC, we sample 32 independent responses for each prompt and report the average accuracy as mean@32 . In addition, we provide the majority-vote (Zhao et al., 2023) accuracy maj@32 as a complementary metric. For the larger and less challenging MATH500 benchmark, we sample 4 responses per prompt and report the average accuracy mean@4 and majority-vote accuracy maj@4 . All evaluations are conducted using the verl framework (Sheng et al., 2024) and follow the same evaluation protocol as DAPO (Yu et al., 2025). More details can be found in the Appendix.

Table 2: Performance comparison of the Qwen2.5-3B and Qwen2.5-7B models trained with DAPO, DAPO+RA, and DAPO+ERPO. Evaluations use mean@32 and maj@32 for AIME25, AIME24 and AMC23; MATH500 is reported with mean@4 and maj@4. The Avg. columns average the mean and maj across datasets.

Method	AIME25		AIME24		AMC23		MATH500		Avg.	
	mean@32	maj@32	mean@32	maj@32	mean@32	maj@32	mean@4	maj@4	mean	maj
Qwen2.5-3B										
DAPO	4.5	8.9	9.5	15.2	58.4	68.0	50.4	50.8	30.7	35.7
+RA	7.7	11.4	9.6	14.1	59.1	67.1	52.0	52.5	32.1	36.3
+ERPO	5.5	8.8	10.3	16.0	60.8	70.0	51.1	51.6	31.9	36.6
Qwen2.5-7B										
DAPO	12.6	16.9	17.5	20.1	76.7	81.4	60.3	60.9	41.8	44.8
+RA	13.5	16.4	16.1	18.1	75.8	80.2	61.2	61.6	41.7	44.1
+ERPO	14.2	19.4	19.0	21.2	76.4	81.5	61.7	62.0	42.8	46.0

5.2 BENCHMARK COMPARISONS

In this section, we compare the performance of DAPO, Reactivated Advantage (RA), and ERPO on the AIME25, AIME24, AMC23, and MATH500 benchmarks. The detailed results are shown in Table 2. On Qwen-3B, both RA and ERPO achieve higher mean and majority-vote accuracy than the baseline DAPO. The improvement of RA demonstrates that residual prompts still contain valuable training information and should not be totally excluded from RL training. On AIME2025, RA achieves a remarkable performance gain compared to the baseline: around a 70% improvement on *mean@32* and a 28% improvement on *maj@32*. Since AIME2025 is shown to suffer less from data contamination during model pretraining than the other math benchmarks (Wu et al., 2025), these results confirm that learning on residual prompts is particularly helpful for tasks that are novel and challenging for the model.

On Qwen2.5-7B, ERPO achieves the best overall performance in both mean and majority-vote accuracy compared to DAPO and RA, indicating the scalability of our algorithm. On AIME2025, ERPO achieves the largest improvement over the baseline, with an increase of approximately 12% and 16% on *mean@32* and *maj@32*. However, unlike the results on the 3B model, RA performs worse than ERPO. A possible reason is that reactivating all residual prompts may lead to overfitting when the proportion of residual prompts is high during training. As shown in Table 1, Qwen2.5-7B has more than 20% residual prompts, making this issue more pronounced as the model scales up. In contrast, ERPO avoids this problem by setting $T_{\max}$, which prevents unbounded increases in sampling temperature. Once a residual prompt is fully learned and robust to higher temperatures, it no longer provides a training signal.

In summary, RA verifies that residual prompts contain information that can still benefit model training and should not be totally discarded. ERPO further provides an effective sampling strategy that leverages training signals from residual prompts and scales effectively to larger models.

5.3 EXPLORATION ON RESIDUAL PROMPTS

To investigate the effect of sampling temperature on residual prompts, we conduct experiments to measure the proportion of residual prompts under different temperature settings. Specifically, we select a 2k subset from our training dataset DAPO-Math-17K, sample each prompt 16 times following the same training configuration, and calculate the proportion of residual prompts within this subset. We evaluate Qwen2.5-3B, Qwen2.5-7B, and Qwen2.5-32B trained with DAPO. For Qwen2.5-32B+DAPO, we use the publicly released checkpoints from DAPO (Yu et al., 2025). The detailed results are presented in Table 1. Our findings highlight three key observations: (1) the proportion of residual prompts increases after training; (2) larger models tend to produce more residual prompts, revealing the challenge of scaling RLVR with model size; and (3) higher sampling temperatures encourage greater exploration and can elicit more incorrect responses from residual prompts.

Table 3: Proportion of prompts with **all-incorrect** responses under different sampling temperatures and model scales.

	$T = 1.0$	$T = 1.1$	$T = 1.2$
Qwen2.5-3B	73.0%	–	–
Qwen2.5-3B + DAPO	39.6%	40.6%	44.0%
Qwen2.5-7B	68.9%	–	–
Qwen2.5-7B + DAPO	25.1%	26.9%	33.2%
Qwen2.5-32B	48.9%	–	–
Qwen2.5-32B + DAPO	7.0%	7.9%	15.2%

Table 4: Sensitivity analysis on temperature increment step T_s and maximum rollout temperature $T_{\max}$ using Qwen2.5-3B. Evaluations use mean@32/maj@32 for AIME25, AIME24, AMC23; MATH uses mean@4 and maj@4. The Avg. columns average mean and maj across datasets (using mean@4/maj@4 for MATH).

Setting	AIME25		AIME24		AMC23		MATH		Avg.	
	mean@32	maj@32	mean@32	maj@32	mean@32	maj@32	mean@4	maj@4	mean	maj
$T_s = 0.02, T_{\max} = 1.2$	5.5	8.8	10.3	16.0	60.8	70.0	51.1	51.6	31.9	36.6
$T_s = 0.05, T_{\max} = 1.2$	5.6	9.4	8.8	14.2	59.2	69.5	51.2	51.9	31.2	36.3
$T_s = 0.05, T_{\max} = 1.4$	3.6	5.5	9.9	14.6	59.6	67.7	52.2	52.7	31.3	35.1

On the other hand, we also examine the effect of sampling temperature on prompts with all-incorrect responses. The experimental settings are kept the same, and the results are presented in Table 3. The findings indicate that RL training and model scaling reduce the proportion of prompts with all-incorrect responses. Moreover, sampling temperature has a much smaller impact on this proportion than on residual prompts. Therefore, ERPO is applied only to residual prompts that are more likely to yield all-correct responses.

5.4 SENSITIVE ANALYSIS

We conduct a sensitivity analysis on the temperature increment step (T_s) and the maximum rollout temperature ($T_{\max}$) to evaluate their impact on the Qwen2.5-3B model. Specifically, we experiment with three parameter settings: $T_s = 0.02, T_{\max} = 1.2$; $T_s = 0.05, T_{\max} = 1.2$; and $T_s = 0.05, T_{\max} = 1.4$. The performance under these settings is reported in Table 4. The results show that $T_s = 0.02, T_{\max} = 1.2$ yields the best performance, while increasing either T_s or $T_{\max}$ leads to performance degradation. Notably, the setting $T_s = 0.05, T_{\max} = 1.2$ still outperforms the vanilla DAPO algorithm, whereas $T_s = 0.05, T_{\max} = 1.4$ achieves performance comparable to the DAPO baseline. Since smaller models are generally less robust than larger ones, these results on the 3B model further validate the effectiveness of ERPO and demonstrate that it is not highly sensitive to the choice of hyperparameters.

5.5 FURTHER ANALYSIS

Sampling Temperature The average and maximum sampling temperatures during the ERPO training process are shown in Figure 2. The maximum temperature increases linearly, whereas the average temperature increases exponentially, indicating that more prompts become residual prompts whose sampling temperatures are raised by ERPO. Setting an upper bound on the temperature, $T_{\max}$, is necessary to prevent uncontrolled growth of the sampling temperature.

Residual Prompts Figure 3 shows the number of residual prompts with all-correct responses and the number of prompts with all-incorrect responses in a training batch of size 512. During training, the number of prompts with all-incorrect responses continues to decrease, while the number of residual prompts steadily increases. Moreover, the growth rate of residual prompts is higher than the decay rate of all-incorrect prompts, underscoring the importance of leveraging residual prompts. In addition, the history tracker for Qwen2.5-3B and Qwen2.5-7B indicates that 15.3% and 40.2% of

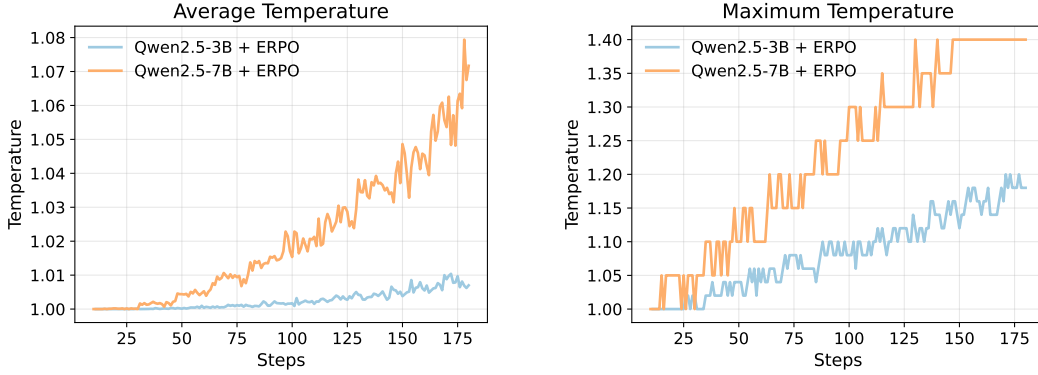


Figure 2: The average and maximum sampling temperatures during the ERPO training process. The steps shown here are the prompt generation steps.

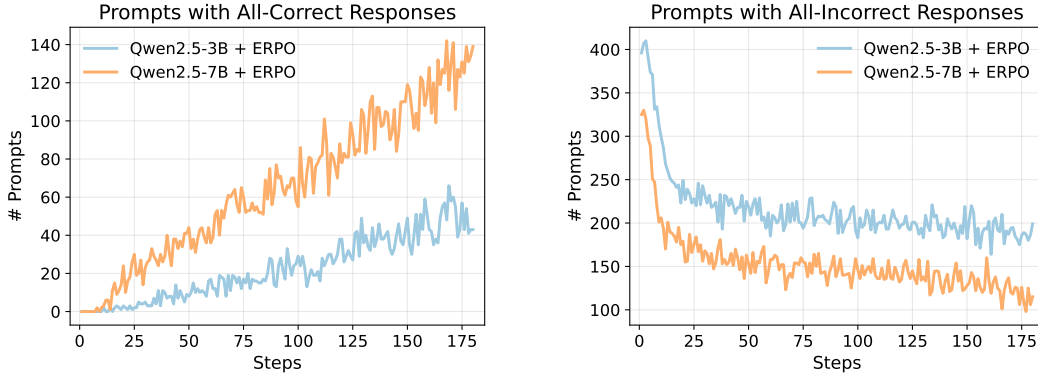


Figure 3: The number of residual prompts with all-correct responses and prompts with all-incorrect responses during the ERPO training process. The steps shown here are the prompt generation steps.

the prompts in the training dataset have a record $H_i > 0$, further demonstrating the critical role of ERPO in the training process.

6 CONCLUSION

In this work, we address a key limitation of GRPO-based reinforcement learning for LLMs: the accumulation of residual prompts that diminish training diversity and leave valuable signals underutilized. To tackle this, we introduce the ERPO framework, which adaptively adjusts sampling temperature based on prompt history to reactivate training signals and encourage broader exploration. Our experiments across multiple math reasoning benchmarks demonstrate that ERPO not only mitigates prompt collapse but also improves both average and majority-vote performance, with especially strong gains on tasks less affected by data contamination. These results highlight the potential of exploiting residual prompts as a promising direction for advancing reinforcement learning with verifiable rewards.

ETHICS STATEMENT

This work uses only publicly available mathematical datasets without personal or sensitive information. The study does not involve human subjects or animals. Our method focuses on improving reasoning in math tasks, with minimal risk of societal harm, and is intended solely for research purposes.

REFERENCES

- Huayu Chen, Kaiwen Zheng, Qingsheng Zhang, Ganqu Cui, Yin Cui, Haotian Ye, Tsung-Yi Lin, Ming-Yu Liu, Jun Zhu, and Haoxiang Wang. Bridging supervised learning and reinforcement learning in math reasoning. *arXiv preprint arXiv:2505.18116*, 2025a.
- Lichang Chen, Jiuhai Chen, Chenxi Liu, John Kirchenbauer, Davit Sotolia, Chen Zhu, Tom Goldstein, Tianyi Zhou, and Heng Huang. Optune: Efficient online preference tuning. *arXiv preprint arXiv:2406.07657*, 2024a.
- Ruibo Chen, Tianyi Xiong, Yihan Wu, Guodong Liu, Zhengmian Hu, Lichang Chen, Yanshuo Chen, Chenxi Liu, and Heng Huang. Gpt-4 vision on medical image classification—a case study on covid-19 dataset. *arXiv preprint arXiv:2310.18498*, 2023.
- Ruibo Chen, Yihan Wu, Lichang Chen, Guodong Liu, Qi He, Tianyi Xiong, Chenxi Liu, Junfeng Guo, and Heng Huang. Your vision-language model itself is a strong filter: Towards high-quality instruction tuning with data selection. *arXiv preprint arXiv:2402.12501*, 2024b.
- Ruibo Chen, Yihan Wu, Yanshuo Chen, Chenxi Liu, Junfeng Guo, and Heng Huang. A watermark for order-agnostic language models. *arXiv preprint arXiv:2410.13805*, 2024c.
- Yanshuo Chen, Zhengmian Hu, Yihan Wu, Ruibo Chen, Yongrui Jin, Marcus Zhan, Chengjin Xie, Wei Chen, and Heng Huang. Enhancing privacy in biosecurity with watermarked protein design. *Bioinformatics*, pp. btaf141, 2025b.
- Daixuan Cheng, Shaohan Huang, Xuekai Zhu, Bo Dai, Wayne Xin Zhao, Zhenliang Zhang, and Furu Wei. Reasoning with exploration: An entropy perspective. *arXiv preprint arXiv:2506.14758*, 2025.
- Ganqu Cui, Yuchen Zhang, Jiacheng Chen, Lifan Yuan, Zhi Wang, Yuxin Zuo, Haozhan Li, Yuchen Fan, Huayu Chen, Weize Chen, et al. The entropy mechanism of reinforcement learning for reasoning language models. *arXiv preprint arXiv:2505.22617*, 2025.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. The llama 3 herd of models. *arXiv e-prints*, pp. arXiv–2407, 2024.
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- 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. *arXiv preprint arXiv:2103.03874*, 2021.
- Zhenyu Hou, Ziniu Hu, Yujiang Li, Rui Lu, Jie Tang, and Yuxiao Dong. Treerl: Llm reinforcement learning with on-policy tree search. *arXiv preprint arXiv:2506.11902*, 2025.
- Chenxi Liu, Tianyi Xiong, Ruibo Chen, Yihan Wu, Junfeng Guo, Tianyi Zhou, and Heng Huang. Modality-balancing preference optimization of large multimodal models by adversarial negative mining. *arXiv preprint arXiv:2506.08022*, 2025.
- Shervin Minaee, Tomas Mikolov, Narjes Nikzad, Meysam Chenaghlu, Richard Socher, Xavier Amatriain, and Jianfeng Gao. Large language models: A survey. *arXiv preprint arXiv:2402.06196*, 2024.
- Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke Zettlemoyer, Percy Liang, Emmanuel Candès, and Tatsunori Hashimoto. s1: Simple test-time scaling. *arXiv preprint arXiv:2501.19393*, 2025.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35: 27730–27744, 2022.

- Qwen, :, An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. Qwen2.5 technical report, 2025. URL <https://arxiv.org/abs/2412.15115>.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *Advances in neural information processing systems*, 36:53728–53741, 2023.
- John Schulman, Sergey Levine, Pieter Abbeel, Michael Jordan, and Philipp Moritz. Trust region policy optimization. In *International conference on machine learning*, pp. 1889–1897. PMLR, 2015.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- Rulin Shao, Shuyue Stella Li, Rui Xin, Scott Geng, Yiping Wang, Sewoong Oh, Simon Shaolei Du, Nathan Lambert, Sewon Min, Ranjay Krishna, et al. Spurious rewards: Rethinking training signals in rlvr. *arXiv preprint arXiv:2506.10947*, 2025.
- Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, YK Li, Yang Wu, et al. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*, 2024.
- Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. Hybridflow: A flexible and efficient rlhf framework. *arXiv preprint arXiv: 2409.19256*, 2024.
- Reza Shirkavand, Shangqian Gao, Peiran Yu, and Heng Huang. Cost-aware contrastive routing for llms. *arXiv preprint arXiv:2508.12491*, 2025a.
- Reza Shirkavand, Qi He, Peiran Yu, and Heng Huang. Bilevel zofo: Bridging parameter-efficient and zeroth-order techniques for efficient llm fine-tuning and meta-training. *arXiv preprint arXiv:2502.03604*, 2025b.
- Vaishnavi Shrivastava, Ahmed Awadallah, Vidhisha Balachandran, Shivam Garg, Harkirat Behl, and Dimitris Papailiopoulos. Sample more to think less: Group filtered policy optimization for concise reasoning. *arXiv preprint arXiv:2508.09726*, 2025.
- Gemini Team, Rohan Anil, Sebastian Borgeaud, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soriccut, Johan Schalkwyk, Andrew M Dai, Anja Hauth, Katie Millican, et al. Gemini: a family of highly capable multimodal models. *arXiv preprint arXiv:2312.11805*, 2023.
- Shenzhi Wang, Le Yu, Chang Gao, Chujie Zheng, Shixuan Liu, Rui Lu, Kai Dang, Xionghui Chen, Jianxin Yang, Zhenru Zhang, et al. Beyond the 80/20 rule: High-entropy minority tokens drive effective reinforcement learning for llm reasoning. *arXiv preprint arXiv:2506.01939*, 2025.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35:24824–24837, 2022.
- Mingqi Wu, Zhihao Zhang, Qiaole Dong, Zhiheng Xi, Jun Zhao, Senjie Jin, Xiaoran Fan, Yuhao Zhou, Huijie Lv, Ming Zhang, et al. Reasoning or memorization? unreliable results of reinforcement learning due to data contamination. *arXiv preprint arXiv:2507.10532*, 2025.
- Wei Xiong, Jiarui Yao, Yuhui Xu, Bo Pang, Lei Wang, Doyen Sahoo, Junnan Li, Nan Jiang, Tong Zhang, Caiming Xiong, et al. A minimalist approach to llm reasoning: from rejection sampling to reinforce. *arXiv preprint arXiv:2504.11343*, 2025.
- Yixuan Even Xu, Yash Savani, Fei Fang, and Zico Kolter. Not all rollouts are useful: Down-sampling rollouts in llm reinforcement learning. *arXiv preprint arXiv:2504.13818*, 2025.

- An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, et al. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*, 2025a.
- Shihui Yang, Chengfeng Dou, Peidong Guo, Kai Lu, Qiang Ju, Fei Deng, and Rihui Xin. Dcpo: Dynamic clipping policy optimization. *arXiv preprint arXiv:2509.02333*, 2025b.
- Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Tom Griffiths, Yuan Cao, and Karthik Narasimhan. Tree of thoughts: Deliberate problem solving with large language models. *Advances in neural information processing systems*, 36:11809–11822, 2023.
- Qiyang Yu, Zheng Zhang, Ruofei Zhu, Yufeng Yuan, Xiaochen Zuo, Yu Yue, Weinan Dai, Tiantian Fan, Gaohong Liu, Lingjun Liu, et al. Dapo: An open-source llm reinforcement learning system at scale. *arXiv preprint arXiv:2503.14476*, 2025.
- Xiaojiang Zhang, Jinghui Wang, Zifei Cheng, Wenhao Zhuang, Zheng Lin, Minglei Zhang, Shaojie Wang, Yinghan Cui, Chao Wang, Junyi Peng, et al. Srpo: A cross-domain implementation of large-scale reinforcement learning on llm. *arXiv preprint arXiv:2504.14286*, 2025.
- Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, et al. A survey of large language models. *arXiv preprint arXiv:2303.18223*, 1(2), 2023.
- Chujie Zheng, Shixuan Liu, Mingze Li, Xiong-Hui Chen, Bowen Yu, Chang Gao, Kai Dang, Yuqiong Liu, Rui Men, An Yang, et al. Group sequence policy optimization. *arXiv preprint arXiv:2507.18071*, 2025a.
- Haizhong Zheng, Rui Liu, Fan Lai, and Atul Prakash. Coverage-centric coresets for high pruning rates. *arXiv preprint arXiv:2210.15809*, 2022.
- Haizhong Zheng, Yang Zhou, Brian R Bartoldson, Bhavya Kailkhura, Fan Lai, Jiawei Zhao, and Beidi Chen. Act only when it pays: Efficient reinforcement learning for llm reasoning via selective rollouts. *arXiv preprint arXiv:2506.02177*, 2025b.
- Tianyu Zheng, Tianshun Xing, Qingshui Gu, Taoran Liang, Xingwei Qu, Xin Zhou, Yizhi Li, Zhoufutu Wen, Chenghua Lin, Wenhao Huang, et al. First return, entropy-eliciting explore. *arXiv preprint arXiv:2507.07017*, 2025c.
- Tong Zheng, Hongming Zhang, Wenhao Yu, Xiaoyang Wang, Xinyu Yang, Runpeng Dai, Rui Liu, Huiwen Bao, Chengsong Huang, Heng Huang, et al. Parallel-rl: Towards parallel thinking via reinforcement learning. *arXiv preprint arXiv:2509.07980*, 2025d.
- Xiangxin Zhou, Zichen Liu, Anya Sims, Haonan Wang, Tianyu Pang, Chongxuan Li, Liang Wang, Min Lin, and Chao Du. Reinforcing general reasoning without verifiers. *arXiv preprint arXiv:2505.21493*, 2025.

Appendix

A EXPERIMENTAL DETAILS

A.1 TRAINING DETAILS

We provide detailed settings of various parameters in the DAPO algorithm, which serves as both the baseline and the optimization method for Reactivated Advantage (RA) and ERPO. The KL coefficient is fixed at 0 across all experiments. The clip ratio is set to $\epsilon_{low} = 0.2$ and $\epsilon_{high} = 0.28$. The maximum response length is set to 10,240 for experiments on the Qwen2.5-3B model and the RA algorithm, and to 20,480 for experiments with DAPO and ERPO on the Qwen2.5-7B model. The overlong buffer is set to 4,096, with an overlong penalty factor of 1. 180 prompt generation steps/2880 policy update steps is used for all experiments.

A.2 EVALUATION DETAILS

We follow the same evaluation protocol as DAPO (Yu et al., 2025), using the verl framework () to assess all benchmarks. Specifically, each question from the benchmark is prepended with the prompt `Solve the following math problem step by step. The last line of your response should be of the form Answer: $Answer (without quotes), where $Answer is the solution to the problem.` and appended with the prompt `\n\nRemember to put your answer on its own line after "Answer:". This structure is identical to that used in the training data. We then follow the same workflow as DAPO to extract the final answer from the model responses.`

B ADDITIONAL EXPERIMENTAL RESULTS

Comparison with additional baselines. We further compare our model performance with another baseline Entropy (Cheng et al., 2025), which use the same backbone model Qwen2.5-7B, training dataset DAPO-Math-17K, optimization method DAPO and very similar hyperparameters. The results are shown in Table 5. Results show that ERPO outperforms Entropy on every benchmark by a large margin, demonstrating the effectiveness of ERPO.

Table 5: Performance comparison of the Qwen2.5-7B model trained with the Entropy baseline and ERPO. For the Entropy baseline, we report the values provided in their paper. For ERPO, we report the *mean@32* scores on AIME25, AIME24, and AMC23, and the *mean@4* score on MATH500.

Method	AIME25	AIME24	AMC23	MATH500	Avg.
Qwen2.5-7B					
Entropy	11.8	12.6	57.8	58.5	35.2
ERPO	14.2	19.0	76.4	61.7	42.8

Model Performance. We further present the model performance on AIME25 throughout training in Figure 4. On both Qwen2.5-3B and Qwen2.5-7B, ERPO consistently outperforms the DAPO baseline for most of the training process, demonstrating its effectiveness on novel and challenging math tasks that are less affected by data contamination. RA achieves the best performance on Qwen2.5-3B but only the second-best on Qwen2.5-7B, suggesting that training on residual prompts can provide notable benefits, but its advantages diminish as model size scales up.

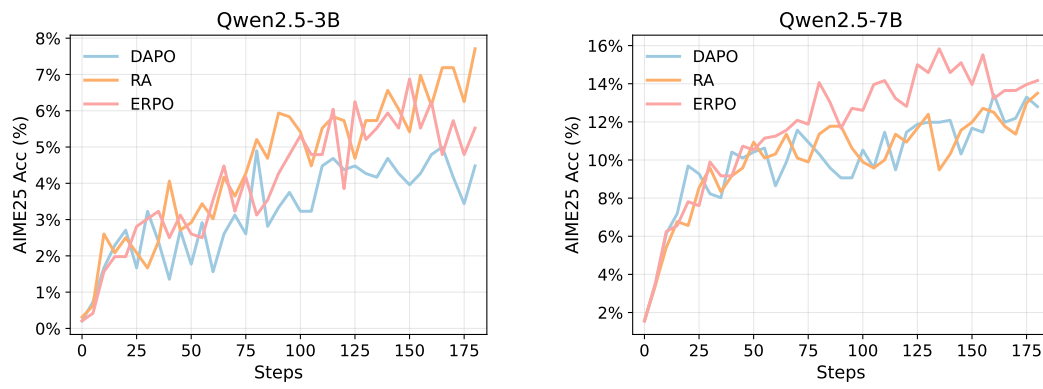


Figure 4: Performance of $mean@32$ on AIME2025.