
Diversity-Aware Policy Optimization for Large Language Model Reasoning

Jian Yao¹, Ran Cheng^{1,2,3}*, Xingyu Wu¹, Jibin Wu^{1,2}, Kay Chen Tan¹

¹ Department of Data Science and Artificial Intelligence, The Hong Kong Polytechnic University

² Department of Computing, The Hong Kong Polytechnic University

³ The Hong Kong Polytechnic University Shenzhen Research Institute, Shenzhen, China
nigel97.yao@connect.polyu.hk, {ran-peter.cheng, xingyu.wu, jibin.wu, kctan}@polyu.edu.hk

Abstract

The reasoning capabilities of large language models (LLMs) have advanced rapidly, particularly following the release of DeepSeek-R1, which has inspired a surge of research into data quality and reinforcement learning (RL) algorithms. Despite the pivotal role diversity plays in RL, its influence on LLM reasoning remains largely underexplored. To bridge this gap, this work presents a systematic investigation into the impact of diversity in RL-based training for LLM reasoning, and proposes a novel diversity-aware policy optimization method. Across evaluations on 12 LLMs, we observe a strong positive correlation between the solution diversity and Potential@k (a novel metric quantifying an LLM’s reasoning potential) in high-performing models. This finding motivates our method to explicitly promote diversity during RL training. Specifically, we design a token-level diversity and reformulate it into a practical objective, then we selectively apply it to positive samples. Integrated into the R1-zero training framework, our method achieves a 3.5% average improvement across four mathematical reasoning benchmarks, while generating more diverse and robust solutions. The code is available at https://github.com/nigelyaoj/R1_zero_Div.

1 Introduction

Recently, the reasoning capabilities of large language models (LLMs) have made remarkable progress, with significant improvements showcased by OpenAI-o1 [35], DeepSeek-R1 [13], and Kimi-k1.5 [43]. Among these advancements, two key innovations have contributed significantly: First, the adoption of a rule-based reward system significantly streamlines the training process by focusing exclusively on rewarding correct final answers and proper output formats, thereby eliminating the complexity associated with process-based reward models [30, 47]. Second, the introduction of a lightweight reinforcement learning (RL) algorithm [13, 43] removes the need for a separate critic model, substantially reducing computational overhead and accelerating the training process. The success of DeepSeek-R1 has attracted numerous follow-up studies [62], which broadly fall into two categories. The first category focuses on improving the quality of training data [33, 15, 21, 2], emphasizing rigorous data set curation through filtering, deduplication, and verification. The second category refines RL algorithms, including detailed optimizations for PPO-based methods such as VCPPO [56] and VAPO [55]; enhancements to GRPO for stability and speed, such as DAPO [54], Dr.GRPO [31] and SRPO [63]; as well as alternative approaches such as REINFORCE++ [20].

While RL has been extensively applied to LLM reasoning, the role of diversity remains largely unexplored in this context, even though it plays a crucial role in RL research [19, 11, 36, 8, 37, 32, 64,

*Correspondence author

12, 58, 66, 6, 52]. In traditional RL tasks, incorporating diversity is widely recognized to facilitate exploration by promoting the selection of more stochastic policies, which helps the policy escape local optima and accelerate the convergence of training. This hypothesis has been experimentally validated in previous work [19, 11, 36]. Beyond empirical evidence, theoretical analyses suggest that policies with higher entropy (a measure of diversity) can smooth the optimization landscape [1]. These findings naturally lead us to ask the following question: **Is promoting diversity essential during RL training for LLM reasoning?**

Intuitively, an LLM capable of generating diverse responses could broaden the exploration of reasoning paths, enabling the model to avoid overfitting to narrow solution patterns in mathematical or logical tasks. To formally address this question, we conduct an evaluation of diversity in LLM reasoning, with a specific focus on mathematical problem-solving. We introduce a novel metric, Potential@k, to quantify an LLM’s reasoning potential (the possible performance gain after RL training). We empirically analyze 12 representative LLMs, examining both their solution diversity and Potential@k scores. Notably, our results reveal a strong positive correlation between solution diversity and Potential@k scores among high-performing models, which suggests that diversity directly contributes to improved final performance after RL training.

The empirical findings motivate us to promote diversity during RL training for LLM reasoning. A commonly used approach for this goal is entropy regularization. However, directly increasing the average entropy of LLM outputs can introduce length bias, as longer responses inherently exhibit higher entropy. To address this, we introduce a token-level diversity metric and reformulate the diversity objective into a practical form. Moreover, promoting diversity often entails a quality-diversity trade-off. To mitigate this, we strategically apply diversity enhancement only to positive samples, thereby enriching solution diversity while preserving training stability. This design is akin to fostering diversity in high-quality policies in population-based RL training, ensuring that exploration is guided by task-relevant performance criteria [48]. Finally, we integrate our diversity objective into the R1-zero training method and evaluate the enhanced approach across 4 mathematical reasoning benchmarks. Experimental results demonstrate a 3.5% average performance gain over standard R1-zero training, while our method can generate more diverse solutions.

To summarize, our key contributions are:

- We present the first formal investigation into the role of diversity in LLM reasoning. Through experiments on mathematical benchmarks, we identify a positive correlation between solution diversity and an LLM’s reasoning potential, as measured by our proposed Potential@k metric. This finding provides empirical motivation for incorporating diversity into policy optimization.
- We propose a novel token-level diversity objective, which is reformulated into a practical metric and selectively applied to positive samples. This design is further supported through gradient behavior analysis, offering an insight for balancing quality and diversity during optimization.
- We evaluate our method on four mathematical reasoning benchmarks, each comprising at least 500 problems with stable evaluation metrics. Our method achieves a 3.5% average improvement over standard R1-zero training and consistently produces more diverse solutions.

2 Preliminary

2.1 RL for LLMs

In the context of RL for LLMs, we frame the LLM generation process as an RL problem. Here, the LLM is modeled as a policy that produces outputs (actions) conditioned on input prompts (states) and receives evaluative feedback (rewards) for its generated responses. This formulation aligns the sequential decision-making nature of language generation with RL’s state-action-reward framework, enabling systematic optimization of the model’s behavior through reward signals.

Formally, in the context of LLM generation for mathematical problem-solving, where each prompt is a question, we define the prompt as $q \in \mathcal{Q}$, where $\mathcal{Q}$ represents the set of all possible questions. The set of all potential text outputs o forms an action space $\mathcal{O}$. Each output o consists of tokens, denoted

as $o := (o^1, o^2, \dots, o^t, \dots)$. To generate an output, a policy $\pi_\theta(\cdot|q)$ parameterized by θ is employed, which generates the output according to the distribution:

$$\pi_\theta(o|q) := \prod_t \pi_\theta(o^t|q, o^{<t}), \quad (1)$$

where $o^{<t} = (o^1, o^2, \dots, o^{t-1})$.

2.2 Reinforcement Learning algorithm

The R1-zero training method proposed by DeepSeek-R1 [13] has attracted significant research attention due to its computational efficiency and effectiveness. In our work, we adopt this training method as our backbone. R1-zero incorporates two key innovations: the GRPO algorithm [40] and a rule-based reward function. In this section, we introduce both components.

Group Relative Policy Optimization (GRPO) GRPO streamlines the process by eliminating the need for a separate critic model, which is usually as large as the policy model, and instead estimates baselines using group scores. Specifically, for each question q , GRPO samples a group of outputs $\{o_1, o_2, \dots, o_G\}$ from the old policy π_{old} and optimizes the policy π_θ by maximizing the following objective:

$$J_{GRPO}(\pi_\theta) = \mathbb{E}_{q \sim \mathcal{Q}, \{o_i\}_{i=1}^G \sim \pi_{old}(\cdot|q)} \left(\frac{1}{G} \sum_{i=1}^G \left(\min \left(\frac{\pi_\theta(o_i|q)}{\pi_{old}(o_i|q)} A_i, \text{clip} \left(\frac{\pi_\theta(o_i|q)}{\pi_{old}(o_i|q)}, 1 - \epsilon, 1 + \epsilon \right) A_i \right) - \beta \mathbb{D}_{KL}(\pi_\theta || \pi_{ref}) \right), \quad (2)$$

where ϵ and β are hyperparameters, the KL term is defined as

$$\mathbb{D}_{KL}(\pi_\theta || \pi_{ref}) = \frac{\pi_{ref}(o_i|q)}{\pi_\theta(o_i|q)} - \log \frac{\pi_{ref}(o_i|q)}{\pi_\theta(o_i|q)} - 1, \quad (3)$$

and the advantage A_i is computed using a group of rewards $\{r_1, r_2, \dots, r_G\}$:

$$A_i = \frac{r_i - \text{mean}(\{r_1, r_2, \dots, r_G\})}{\text{std}(\{r_1, r_2, \dots, r_G\})}. \quad (4)$$

Reward functions In line with DeepSeek-R1 [13], we implement two types of rule-based rewards: accuracy rewards and format rewards. The accuracy reward model assesses whether the response is correct by comparing the predicted answer to the golden reference answer, while the format reward model ensures that the final answer is presented in a `\boxed{\}` format for reliable verification.

3 Correlation between LLMs' reasoning potential and solution diversity

The role of diversity has long been established as critical in traditional RL tasks. Numerous studies [19, 11, 36, 8, 37] have shown that promoting diversity can enhance the final quality of the policy. However, its impact in the realm of RL for LLM reasoning still remains under-explored. In this section, we investigate the relationship between solution diversity and the reasoning abilities of LLMs on mathematical benchmarks. We adopt the equation diversity in prior work [49] to quantify the variety of solutions generated for mathematical problem-solving. For reasoning ability, we introduce a novel metric to evaluate an LLM's training potential (related to the performance gain achieved after RL training).

Experimental setup We evaluate 12 LLMs on the MATH benchmark [16]. For each question, we calculate: (1) **Pass@1 accuracy** using greedy decoding, and (2) **Diversity** with (3) **Potential@k**, both evaluated from 16 sampled responses (temperature=0.9).

For diversity, we adopt the metric (denoted as Div-Equ) from prior work [49], which measures the ratio of distinct equations among the responses:

$$\text{Div-Equ} := \frac{1}{N} \sum_{i=1}^N \frac{|\mathcal{U}_i|}{|\mathcal{A}_i|}, \quad (5)$$

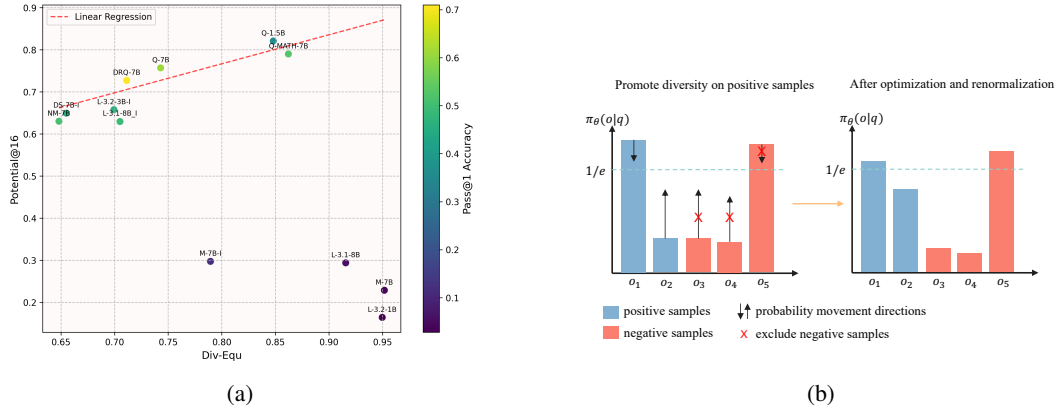


Figure 1: **(a)** Evaluation of Pass@1 accuracy, Div-Equ diversity, and Potential@16 across 12 LLMs on the MATH benchmark. Model naming conventions: Prefixes denote base architectures (Q: Qwen2.5-Math, DS: DeepSeekMath, M: Mistral, L: Llama, DRQ: DeepSeek-R1-Distill-Qwen, NM: NuminaMath); suffix '-I' indicates '-Instruct'. **(b)** Illustration of probability movement during diversity optimization on positive samples.

where $\mathcal{U}_i$ and $\mathcal{A}_i$ are the sets of unique equations and all equations extracted from the k sampled responses (with $k = 16$ in our experiments) of question i , respectively. And $N = 500$ is the amount of the data.

For Potential, we define a metric termed Potential@k to quantify the model's capability to correct answers within k trials (with $k = 16$ in our experiments) on its Pass@1 failure samples. Formally:

$$\text{Potential@k} := \frac{\sum_{i=1}^N \text{Pass@k}(q_i) \cdot (1 - \text{Pass@1}(q_i))}{\sum_{i=1}^N (1 - \text{Pass@1}(q_i))}, \quad (6)$$

where q_i denotes the i -th question.

Empirical findings The results are shown in Figure 1a. The results show a bifurcated pattern: For LLMs with limited reasoning ability ($\text{Pass@1} < 0.4$), we observe no significant relationship between solution diversity and model potential. For stronger performers ($\text{Pass@1} > 0.4$), a clear positive correlation emerges between these metrics. Linear regression on this high-performing subset yields $R^2 = 0.81$, confirming a strong predictive relationship where increased diversity corresponds to higher model potential.

Through an investigation of the Objective 2 in the GRPO algorithm, we observe that for each question in the training set, if all samples within a group are either entirely positive or entirely negative, the advantage score becomes 0, resulting in no gradient update. Crucially, the training signal originates from the reward discrepancy between positive and negative samples within the group, which is inherently linked to our definition of potential (to some extent, the algorithm's improvement can be characterized by the dynamics of this potential metric, as discussed in Appendix B). This indicates that promoting diversity for LLM may result in higher performance after RL training.

Takeaways A positive correlation between the LLM's reasoning potential and solution diversity is observed in our experiment. As illustrated in Section 2.2, the optimization direction is guided by correct answers in multiple sampled responses. This directly links our Potential@k metric to RL training improvements. Hence, the observation strongly motivates us to enhance diversity during the RL training process.

4 Diversity-aware policy optimization

Building on the insights from Section 3, in this section, we introduce an entropy-based diversity and propose its targeted application to positive samples during policy optimization for LLM reasoning. We incorporate this diversity objective into the **R1-zero** training method [13], which employs the

GRPO algorithm with the reward function defined in Section 2.2. We refer to this enhanced approach as **R1-zero-Div**.

4.1 Entropy-based diversity

A straightforward approach is to define diversity as the average entropy of the LLM’s outputs per question i.e., $E_{q \sim Q}[\mathcal{H}(\pi_\theta(\cdot|q))]$. However, this formulation introduces length bias: longer responses inherently exhibit higher entropy (due to more token-level uncertainties), causing the metric to artificially favor longer outputs regardless of actual solution diversity. To address this issue, we introduce token-level entropy, which calculates the entropy for each token sampled from the old policy π_{old} . Formally, we define:

$$\hat{J}_{Div}(\pi_\theta) := \mathbb{E}_{q \sim Q, o \sim \pi_{old}(\cdot|q)} \left[\frac{1}{T} \sum_{t=1}^T \mathcal{H}(\pi_\theta(\cdot|q, o^{<t})) \right], \quad (7)$$

where T is the length of the output.

During training, the gradient of diversity with respect to the policy π_θ in the $\mathcal{H}(\pi_\theta(\cdot|q, o^{<t}))$ is intractable. We therefore reformulate the diversity objective to enable effective backpropagation:

$$\begin{aligned} \hat{J}_{Div}(\pi_\theta) &= \mathbb{E}_{q \sim Q, o \sim \pi_{old}(\cdot|q)} \left[-\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\tilde{o}^t \sim \pi_\theta(\cdot|q, o^{<t})} [\log \pi_\theta(\tilde{o}^t|q, o^{<t})] \right] \\ &= \mathbb{E}_{q \sim Q, o \sim \pi_{old}(\cdot|q)} \left[-\frac{1}{T} \sum_{t=1}^T \frac{\pi_\theta(o^t|q, o^{<t})}{\pi_{old}(o^t|q, o^{<t})} \log \pi_\theta(o^t|q, o^{<t}) \right]. \end{aligned} \quad (8)$$

A proof for the last equation can be found in Appendix A.1. In practice, building on the R1-zero training method, we can use the samples within the group to calculate Objective 8.

4.2 Promoting diversity on positive samples

Empirical evidence indicates that the direct application of Objective 8 inadvertently increases diversity in incorrect solutions. Intuitively, negative samples offer more room for diversity enhancement, which can skew the model’s optimization process. To address this issue, we concentrate on promoting diversity exclusively within positive samples:

$$J_{Div}(\pi_\theta) = \mathbb{E}_{q \sim Q, o \sim \pi_{old}(\cdot|q)} \left[-\mathbb{I}(r=1) \cdot \frac{1}{T} \sum_{t=1}^T \frac{\pi_\theta(o^t|q, o^{<t})}{\pi_{old}(o^t|q, o^{<t})} \log \pi_\theta(o^t|q, o^{<t}) \right], \quad (9)$$

where $\mathbb{I}(\cdot)$ denotes the indicator function and r is the accuracy reward for output o .

This is akin to fostering diversity in high-quality policies in population-based RL training [48], while we focus on positive samples rather than policies here. Beyond intuitive justification, we further justify this design by analyzing the gradient on each token.

According to Equation 8, we have:

$$\nabla_{\pi_\theta} \hat{J}_{Div}(\pi_\theta) = \mathbb{E}_{q \sim Q, o \sim \pi_{old}(\cdot|q)} \left[-\frac{1}{T} \sum_{t=1}^T \frac{\nabla_{\pi_\theta} [\pi_\theta(o^t|q, o^{<t}) \log \pi_\theta(o^t|q, o^{<t})]}{\pi_{old}(o^t|q, o^{<t})} \right]. \quad (10)$$

Thus, the gradient can be decomposed into per-token contributions (each term in the summation contributes a component). Up to a constant scaling factor, the gradient from each token is:

$$-\nabla_{\pi_\theta} \pi_\theta(o^t|q, o^{<t}) \log \pi_\theta(o^t|q, o^{<t}) = -[1 + \log \pi_\theta(o^t|q, o^{<t})] \cdot \nabla_{\pi_\theta} \pi_\theta(o^t|q, o^{<t}). \quad (11)$$

Hence, for tokens with small probabilities (in that case $\pi_\theta(o^t|q, o^{<t}) < e^{-1}$, and this holds for most of tokens since the sum of probability is equal to 1), the gradient aligns with $\nabla_{\pi_\theta} \pi_\theta(o^t|q, o^{<t})$. This suggests that the diversity component’s gradient actively promotes increasing the probability of

low-probability tokens, which inherently offer substantial growth potential. However, this tendency is undesirable for negative samples. Thus, excluding diversity enhancement for negative samples mitigates conflicts between solution quality and diversity. A visual illustration is provided in Figure 1b. Moreover, the experimental results in Section 5.3 and Appendix E.1 further support our design.

Finally, we incorporate the diversity optimization into the standard R1-zero training, and use the samples in the group to calculate the diversity, yielding the final training objective:

$$\begin{aligned}
J(\pi_\theta) &= J_{GRPO}(\pi_\theta) + \lambda \cdot J_{Div}(\pi_\theta) \\
&= \mathbb{E}_{q \sim \mathcal{Q}, \{o_i\}_{i=1}^G \sim \pi_{old}(\cdot|q)} \frac{1}{G} \sum_{i=1}^G \left[\min \left(\frac{\pi_\theta(o_i|q)}{\pi_{old}(o_i|q)} A_i, \text{clip} \left(\frac{\pi_\theta(o_i|q)}{\pi_{old}(o_i|q)}, 1 - \epsilon, 1 + \epsilon \right) A_i \right) \right. \\
&\quad \left. - \beta \mathbb{D}_{KL}(\pi_\theta || \pi_{ref}) - \lambda \mathbb{I}(r_i = 1) \cdot \frac{1}{T_i} \sum_{t=1}^{T_i} \frac{\pi_\theta(o_i^t|q, o_i^{<t})}{\pi_{old}(o_i^t|q, o_i^{<t})} \log \pi_\theta(o_i^t|q, o_i^{<t}) \right], \quad (12)
\end{aligned}$$

where λ is the diversity weight and i denotes the i -th sample in the group. In practice, we choose $\lambda = 0.01$. Other implementation details are provided in Section 5.1 and Appendix D.

5 Experiments

In this experimental section, we aim to address the following questions:

- Q1. Can our method effectively enhance reasoning abilities and provide diverse solutions?
- Q2. Does the design of the diversity coefficient λ influence the results?
- Q3. Does our method demonstrate consistent performance across different model sizes?

5.1 Experimental setup

Base models We choose Qwen2.5-Math-7B (Qwen7B) [51] as our base model, which is commonly used for mathematical reasoning benchmarks [59, 67, 24]. Additionally, we conduct an ablation study using Qwen2.5-Math-1.5B (Qwen1.5B)[51] to assess the effectiveness of our approach in smaller LLMs.

Benchmarks We selected 4 mathematical benchmarks to evaluate the models' reasoning abilities: GSM8K [7], MATH500 [16], Olympiad Bench [14], and College Math [42]. Each contains at least 500 data points for testing. We excluded some commonly used mathematical benchmarks that provide limited data, e.g, AIME24² with 30 items, as they can lead to unstable and biased evaluation outcomes. We train the base model on the GSM8K training set and then evaluate on the 4 benchmarks.

Baselines The most pertinent baselines for comparison are the base model itself and the base model trained via R1-zero. Additionally, we incorporate the latest prominent "R1-zero-Like" models with similar backbones **for reference**: SimpleRL-Zoo [59], PRIME-Zero-7B [9]. It is important to note that these methods are trained with different computational resources and datasets, making direct comparisons challenging. Our approach is designed to enhance diversity rather than compete directly with these methods. In fact, our method is compatible with and can be integrated into these existing approaches.

Implementation details For R1-zero-Div, we train the base model on the GSM8K training set using the loss function in Equation 12, with a learning rate of 3×10^{-6} and the AdamW optimizer. During rollout, we sample 6 responses with a temperature of 0.9 and train for 2 epochs. Our implementation is built on TRL [46] and runs on 8×A6000 GPUs. For R1-zero, we maintain identical settings to R1-zero-Div but exclude the diversity objective. For other baselines, we evaluate open-sourced models downloaded from Hugging Face³, following the settings recommended in their original papers. Additional implementation details are provided in Appendix D.

²https://huggingface.co/datasets/Maxwell-Jia/AIME_2024

³<https://huggingface.co>

Table 1: Pass@1 accuracy on mathematical benchmarks.

Method	GSM8K	MATH500	Olympiad Bench	College Math	Avg
Qwen2.5-Math-7B	57.6	51.8	16.1	21.4	36.7
R1-zero	88.7	74.6	37.3	43.3	61.0
R1-zero-Div (Ours)	91.7	78.2	40.1	47.6	64.4
SimpleRL-Zoo	90.2	80.0	39.0	47.2	64.1
Eurus-2-7B-PRIME	88.0	74.4	39.4	46.6	62.1

Table 2: Avg@8 accuracy on mathematical benchmarks.

Method	GSM8K	MATH500	Olympiad Bench	College Math	Avg
Qwen2.5-Math-7B	53.37 (0.56)	48.10 (0.82)	15.80 (0.22)	19.36 (0.14)	34.16
R1-zero	87.77 (0.86)	72.97 (1.20)	37.26 (0.52)	42.22 (0.31)	60.06
R1-zero-Div (Ours)	90.64 (0.89)	76.92 (1.24)	39.19 (0.55)	47.49 (0.32)	63.56
SimpleRL-Zoo	89.46 (0.87)	77.15 (1.23)	39.43 (0.57)	47.19 (0.34)	63.31
Eurus-2-7B-PRIME	88.31 (0.86)	73.92 (1.18)	36.56 (0.50)	45.27 (0.30)	61.02

5.2 Main results

R1-zero-Div enhances reasoning abilities We evaluate the reasoning performance using Pass@1 accuracy, as shown in Table 1. In our experiment, R1-zero-Div demonstrates superior performance compared to R1-zero, achieving an average improvement of 3.5%. Despite being trained with limited computational resources (discussed in Appendix B), R1-zero-Div achieves comparable results to state-of-the-art methods (SimpleRL-Zoo and Eurus-2-7B-PRIME). These results suggest that promoting diversity on positive samples in training can effectively enhance the model’s reasoning capabilities. Also, following the recommendations in prior work [4, 17], we evaluated 8 samples per question with a temperature of 0.5. We report Avg@8 and its standard error in the Table 2. The conclusion regarding the effectiveness of our approach remains consistent with the pass@1 metric results.

R1-zero-Div generates diverse solutions We empirically demonstrate that R1-zero-Div produces more diverse solutions than other RL-finetuning baselines. Our evaluation on the GSM8K test set generates 5 responses for each of 1,319 questions, measuring diversity through three metrics: Div-Equ, and two additional metrics in prior work [26]: (1) N-gram diversity (proportion of distinct n-grams per response, capturing intra-diversity) and (2) Self-BLEU diversity (100 minus Self-BLEU score, capturing inter-diversity). All metrics range from 0 to 100, with higher values indicating greater diversity. As shown in Table 3, while RL fine-tuning methods significantly reduce diversity (compared to the base model), R1-zero-Div effectively preserves diversity. We further provide concrete examples in Appendix E showing that R1-zero-Div generates distinct solutions for the same question.

5.3 Ablation study

We conduct an ablation study to analyze (1) the impact of different diversity weights and (2) our method’s generalization capability on smaller base models.

Analysis on the choice of diversity weights λ Table 4 presents Pass@1 accuracy when applying different λ values to promote diversity on positive samples (denoted as “pos”). The results demonstrate that small values ($\lambda \leq 0.02$) effectively enhance reasoning performance, with $\lambda = 0.01$ emerging as the optimal choice in our experimental setup. We further compare diversity promotion strategies: positive samples only (“pos”) versus all samples (“pos+neg”). The marginal improvement observed when applying diversity to all samples supports our methodological design choice in Section 4.2.

Table 3: Diversity of different methods on GSM8K test set.

Method	Div-Equ	Div-N-gram	Div-Self-BLEU
Qwen2.5-Math-7B	92.26	29.29	85.98
Eurus-2-7B-PRIME	60.86	24.08	48.20
SimpleRL-Zoo	74.89	25.41	49.32
R1-zero	75.02	27.75	56.00
zero-Div (Ours)	79.29	29.60	58.89

Experiment on 1.5B base model We perform both R1-zero-Div and R1-zero on the Qwen2.5-Math-1.5B base model [51], with results shown in Table 5. The experiments demonstrate that, compared to R1-zero, R1-zero-Div enhances reasoning performance on 3 out of 4 benchmarks, achieving an average improvement of 2.3%, validating the scalability of our approach to a smaller model.

Table 4: Ablation Study on different diversity weights on mathematical benchmarks

Method	GSM8K	MATH500	Olympiad Bench	College Math	Avg
$\lambda = 0$	88.7	74.6	37.3	43.3	61.0
$\lambda = 0.05$, pos	88.1	74.8	38.2	45.8	61.7
$\lambda = 0.02$, pos	<u>90.7</u>	<u>76.0</u>	<u>38.4</u>	<u>45.9</u>	<u>62.8</u>
$\lambda = 0.01$, pos	91.7	78.2	40.1	47.6	64.4
$\lambda = 0.01$, pos+neg	89.8	76.6	39.6	46.9	63.2

Table 5: Ablation Study on Qwen2.5-Math-1.5B base model

Method	GSM8K	MATH500	Olympiad Bench	College Math	Avg
Qwen2.5-Math-1.5B	39.4	36.4	23.0	6.6	26.3
R1-zero	82.9	66.4	32.1	43.1	56.1
R1-zero-Div (Ours)	83.2	70.4	32.0	43.9	57.4

6 Related work

RL for LLM reasoning The reasoning capabilities of LLMs have seen remarkable progress recently, with notable improvements demonstrated by OpenAI-o1 [35], DeepSeek-R1 [13], and Kimi-k1.5 [43]. Our work builds upon the R1-zero training method proposed by DeepSeek-R1 [13], which significantly improves LLM reasoning through two innovations that simplify the training pipeline and accelerate training: the GRPO algorithm [40], which replaces critic models with group score baselines, and a rule-based reward system that focuses solely on final answer correctness and output format.

Subsequent research has advanced this approach in two directions: (1) improving training data quality [28, 33, 15, 54, 21, 18] and (2) refining RL algorithms. Regarding RL algorithm refinement, one category focuses on PPO-like methods. SimpleRL-zero [60] demonstrates that PPO with replacing the reward model by a rule-based reward function can significantly improve the LLM’s reasoning ability. VinePPO [25] leverages the flexibility of language environments to compute unbiased Monte Carlo-based estimates, eliminating the need for large value networks. VCPPO [56] employs a pretrained value model to address value initialization bias and decouples Generalized Advantage Estimation (GAE) computation between the actor and critic to mitigate reward signal decay. VAPO [55] further shows that value-based RL frameworks outperform value-free methods in long Chain-of-Thought

reasoning. The second category focuses on GRPO enhancements for stability and speed. DAPO [54] identifies the critical shortcomings (entropy collapse, training instability, and biased loss) in the original GRPO algorithm and addresses them via decoupled clipping and dynamic sampling. Dr.GRPO [31] reveals two biases in GRPO: response-level length bias and question-level difficulty bias. SRPO [63] introduces a two-stage history-resampling method to improve training efficiency. Alternative approaches explore algorithms like REINFORCE: Kimi-k1.5 [43] demonstrates stable training with REINFORCE-like policy gradients, while REINFORCE++ [20] and GPG [5] aim to enhance REINFORCE’s stability and scalability, respectively.

Diversity in RL Research on policy diversity in deep reinforcement learning can be categorized into three groups based on how diversity is utilized [48]. The first category uses diversity primarily to improve exploration efficiency, where diversity emerges as a byproduct of maximizing final task performance [19, 11, 36, 8, 37, 53]. The second category treats diversity either as a constraint (optimizing quality subject to diversity constraints) or as an objective (optimizing diversity under quality constraints) [32, 64, 12, 58, 66]. The third category optimizes quality and diversity simultaneously, known as Quality-Diversity RL methods [6, 38, 45, 3]. Our work extends the first paradigm to RL for LLM reasoning. While existing research in this category has proposed various diversity metrics, such as distance regularization between the current policy and a previous policy [19], reward randomization [41], we develop our approach based on a simple yet effective entropy-based diversity metric.

Diversity in LLMs Prior work has explored diversity in LLMs across several domains. GEM [29] proposes methods to preserve diversity during supervised fine-tuning, while Bstar [61] investigates the exploration-exploitation tradeoff in self-improvement settings. Additional studies have examined diversity in reinforcement learning from human feedback [34, 26] and LLM ensembles [44]. However, diversity remains understudied in RL for LLM reasoning scenarios. To our knowledge, we are the first to formally analyze diversity and propose a principled diversity-aware training method for this setting.

7 Conclusion, limitations, and discussion

In this research, we investigate the role of diversity in RL for LLM reasoning. Through comprehensive evaluations across 12 LLMs, we empirically establish a strong positive correlation between a model’s reasoning potential and the diversity of its generated solutions, underscoring the necessity of fostering diversity during RL training. To this end, we introduce a novel diversity-aware policy optimization method that optimizes the token-level diversity in positive samples. Experimentally, we demonstrate that our method not only enhances LLMs’ reasoning ability but also generates more diverse solutions. By bridging the gap between diversity promotion and policy optimization, we aim to provide new insights for advancing the robustness and creativity of LLMs in complex reasoning scenarios.

Due to computational constraints, our experiments were conducted on $8 \times$ NVIDIA A6000 GPUs, which restricted our analysis to 1.5B and 7B parameter-scale models. This naturally introduces a limitation: the generalizability of our diversity-aware policy optimization method to larger-scale LLMs remains to be explored. While our method demonstrates significant improvements on mid-sized models, extrapolating these findings to larger architectures may require adjustments to the entropy regularization scheme or training dynamics, given the known differences in optimization landscapes across model scales. We urge future research to investigate these scalability challenges and hope our work will inspire the community to explore diversity-enhanced RL strategies for both small and large LLMs, fostering more robust reasoning capabilities across the spectrum of model architectures. Beyond scaling, the diversity-aware optimization mechanism could also be extended to other LLM tasks [65, 50, 23, 22], suggesting its broader applicability beyond reasoning tasks.

Another promising future direction lies in the semantic definition of diversity. In this work, we employ entropy-based regularization to implicitly promote diverse behaviors during LLM generation, which captures statistical variance in output distributions. However, many real-world applications demand user-intended diversity (e.g., requiring both algebraic and arithmetic solutions to a math problem, or generating code with distinct algorithmic approaches). Such scenario-specific diversity requires explicit modeling of user-defined diversity, a challenge well-studied in RL [10, 48]. By bridging LLM reasoning with explicit diversity optimization from RL, future work could unlock

more controllable and context-aware generative capabilities, addressing the gap between statistical diversity and human-intentional variety in complex tasks.

Acknowledgments and Disclosure of Funding

This work was supported in part by the National Natural Science Foundation of China under Grant U21A20512 and the Research Grants Council of the Hong Kong SAR under Grant No. C5052-23G, Grant PolyU 15229824, Grant PolyU 15218622, and Grant PolyU 15215623. This work was also partially supported by Guangdong Basic and Applied Basic Research Foundation (No. 2024B1515020019). This work was also partially supported by the Research Grants Council of the Hong Kong SAR (Grant No. PolyU15217424, PolyU25216423) and the Hong Kong Polytechnic University (Project IDs: P0043563).

References

- [1] Zafarali Ahmed, NicolasLe Roux, Mohammad Norouzi, and Dale Schuurmans. Understanding the impact of entropy on policy optimization. *International Conference on Machine Learning*, 2019.
- [2] Alon Albalak, Duy Phung, Nathan Lile, Rafael Rafailov, Kanishk Gandhi, Louis Castricato, Anikait Singh, Chase Blagden, Violet Xiang, Dakota Mahan, et al. Big-math: A large-scale, high-quality math dataset for reinforcement learning in language models. *arXiv preprint arXiv:2502.17387*, 2025.
- [3] Sumeet Batra, Bryon Tjanaka, Matthew C Fontaine, Aleksei Petrenko, Stefanos Nikolaidis, and Gaurav Sukhatme. Proximal policy gradient arborescence for quality diversity reinforcement learning. *arXiv preprint arXiv:2305.13795*, 2023.
- [4] Nikhil Chandak, Shashwat Goel, and Ameya Prabhu. Incorrect baseline evaluations call into question recent llm-rl claims, 2025. *Notion Blog*, 2025.
- [5] Xiangxiang Chu, Hailang Huang, Xiao Zhang, Fei Wei, and Yong Wang. Gpg: A simple and strong reinforcement learning baseline for model reasoning. *arXiv preprint arXiv:2504.02546*, 2025.
- [6] Geoffrey Cideron, Thomas Pierrot, Nicolas Perrin, Karim Beguir, and Olivier Sigaud. Qd-rl: Efficient mixing of quality and diversity in reinforcement learning. corr abs/2006.08505 (2020). *arXiv preprint arXiv:2006.08505*, 2020.
- [7] 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.
- [8] Edoardo Conti, Vashisht Madhavan, Felipe Petroski Such, Joel Lehman, Kenneth Stanley, and Jeff Clune. Improving exploration in evolution strategies for deep reinforcement learning via a population of novelty-seeking agents. *Advances in neural information processing systems*, 31, 2018.
- [9] Ganqu Cui, Lifan Yuan, Zefan Wang, Hanbin Wang, Wendi Li, Bingxiang He, Yuchen Fan, Tianyu Yu, Qixin Xu, Weize Chen, et al. Process reinforcement through implicit rewards. *arXiv preprint arXiv:2502.01456*, 2025.
- [10] Li Ding, Jenny Zhang, Jeff Clune, Lee Spector, and Joel Lehman. Quality diversity through human feedback: Towards open-ended diversity-driven optimization. *arXiv preprint arXiv:2310.12103*, 2023.
- [11] Benjamin Eysenbach, Abhishek Gupta, Julian Ibarz, and Sergey Levine. Diversity is all you need: Learning skills without a reward function. *arXiv preprint arXiv:1802.06070*, 2018.

- [12] Mahsa Ghasemi, Evan Scope Crafts, Bo Zhao, and Ufuk Topcu. Multiple plans are better than one: Diverse stochastic planning. In *Proceedings of the International Conference on Automated Planning and Scheduling*, volume 31, pages 140–148, 2021.
- [13] 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.
- [14] Chaoqun He, Renjie Luo, Yuzhuo Bai, Shengding Hu, Zhen Leng Thai, Junhao Shen, Jinyi Hu, Xu Han, Yujie Huang, Yuxiang Zhang, et al. Olympiadbench: A challenging benchmark for promoting agi with olympiad-level bilingual multimodal scientific problems. *arXiv preprint arXiv:2402.14008*, 2024.
- [15] Zhiwei He, Tian Liang, Jiahao Xu, Qiuzhi Liu, Xingyu Chen, Yue Wang, Linfeng Song, Dian Yu, Zhenwen Liang, Wenxuan Wang, et al. Deepmath-103k: A large-scale, challenging, decontaminated, and verifiable mathematical dataset for advancing reasoning. *arXiv preprint arXiv:2504.11456*, 2025.
- [16] 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.
- [17] Andreas Hochlehnert, Hardik Bhatnagar, Vishaal Udandara, Samuel Albanie, Ameya Prabhu, and Matthias Bethge. A sober look at progress in language model reasoning: Pitfalls and paths to reproducibility. *arXiv preprint arXiv:2504.07086*, 2025.
- [18] Andreas Hochlehnert, Hardik Bhatnagar, Vishaal Udandara, Ameya Prabhu, and Matthias Bethge. Curatedthoughts: Data curation for rl training datasets, 2025.
- [19] Zhang-Wei Hong, Tzu-Yun Shann, Shih-Yang Su, Yi-Hsiang Chang, Tsu-Jui Fu, and Chun-Yi Lee. Diversity-driven exploration strategy for deep reinforcement learning. *Advances in neural information processing systems*, 31, 2018.
- [20] Jian Hu. Reinforce++: A simple and efficient approach for aligning large language models. *arXiv preprint arXiv:2501.03262*, 2025.
- [21] Jingcheng Hu, Yinmin Zhang, Qi Han, Daxin Jiang, Xiangyu Zhang, and Heung-Yeung Shum. Open-reasoner-zero: An open source approach to scaling up reinforcement learning on the base model, 2025.
- [22] Beichen Huang, Xingyu Wu, Yu Zhou, Jibin Wu, Liang Feng, Ran Cheng, and Kay Chen Tan. Evaluation of large language models as solution generators in complex optimization. *IEEE Computational Intelligence Magazine*, 20(4):56–70, 2025.
- [23] Yuxiao Huang, Wenjie Zhang, Liang Feng, Xingyu Wu, and Kay Chen Tan. How multimodal integration boost the performance of llm for optimization: Case study on capacitated vehicle routing problems. In *2025 IEEE Symposium for Multidisciplinary Computational Intelligence Incubators (MCII)*, pages 1–7. IEEE, 2025.
- [24] Hugging Face. Open r1: A fully open reproduction of deepseek-r1, January 2025.
- [25] Amirhossein Kazemnejad, Milad Aghajohari, Eva Portelance, Alessandro Sordoni, Siva Reddy, Aaron Courville, and Nicolas Le Roux. Vineppo: Unlocking rl potential for llm reasoning through refined credit assignment. *arXiv preprint arXiv:2410.01679*, 2024.
- [26] Robert Kirk, Ishita Mediratta, Christoforos Nalmpantis, Jelena Luketina, Eric Hambro, Edward Grefenstette, and Roberta Raileanu. Understanding the effects of rlhf on llm generalisation and diversity. *arXiv preprint arXiv:2310.06452*, 2023.
- [27] Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the ACM SIGOPS 29th Symposium on Operating Systems Principles*, 2023.

- [28] Xuefeng Li, Haoyang Zou, and Pengfei Liu. Limr: Less is more for rl scaling. *arXiv preprint arXiv:2502.11886*, 2025.
- [29] Ziniu Li, Congliang Chen, Tian Xu, Zeyu Qin, Jiancong Xiao, Zhi-Quan Luo, and Ruoyu Sun. Preserving diversity in supervised fine-tuning of large language models. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [30] Hunter Lightman, Vineet Kosaraju, Yuri Burda, Harrison Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step. In *The Twelfth International Conference on Learning Representations*, 2023.
- [31] Zichen Liu, Changyu Chen, Wenjun Li, Penghui Qi, Tianyu Pang, Chao Du, Wee Sun Lee, and Min Lin. Understanding rl-zero-like training: A critical perspective. *arXiv preprint arXiv:2503.20783*, 2025.
- [32] Muhammad A Masood and Finale Doshi-Velez. Diversity-inducing policy gradient: Using maximum mean discrepancy to find a set of diverse policies. *arXiv preprint arXiv:1906.00088*, 2019.
- [33] Chunyang Meng, Shijie Song, Haogang Tong, Maolin Pan, and Yang Yu. Deepscaler: Holistic autoscaling for microservices based on spatiotemporal gnn with adaptive graph learning. In *2023 38th IEEE/ACM International Conference on Automated Software Engineering (ASE)*, pages 53–65. IEEE, 2023.
- [34] Sonia K Murthy, Tomer Ullman, and Jennifer Hu. One fish, two fish, but not the whole sea: Alignment reduces language models’ conceptual diversity. *arXiv preprint arXiv:2411.04427*, 2024.
- [35] OpenAI. Learning to reason with llms. <https://openai.com/index/learning-to-reason-with-llms>, 2024.
- [36] Jack Parker-Holder, Aldo Pacchiano, Krzysztof M Choromanski, and Stephen J Roberts. Effective diversity in population based reinforcement learning. *Advances in Neural Information Processing Systems*, 33:18050–18062, 2020.
- [37] Zhenghao Peng, Hao Sun, and Bolei Zhou. Non-local policy optimization via diversity-regularized collaborative exploration. *arXiv preprint arXiv:2006.07781*, 2020.
- [38] Thomas Pierrot, Valentin Macé, Felix Chalumeau, Arthur Flajolet, Geoffrey Cideron, Karim Beguir, Antoine Cully, Olivier Sigaud, and Nicolas Perrin-Gilbert. Diversity policy gradient for sample efficient quality-diversity optimization. In *Proceedings of the Genetic and Evolutionary Computation Conference*, pages 1075–1083, 2022.
- [39] Jeff Rasley, Samyam Rajbhandari, Olatunji Ruwase, and Yuxiong He. Deepspeed: System optimizations enable training deep learning models with over 100 billion parameters. In *Proceedings of the 26th ACM SIGKDD international conference on knowledge discovery & data mining*, pages 3505–3506, 2020.
- [40] Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, YK Li, Y Wu, et al. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*, 2024.
- [41] Zhenggang Tang, Chao Yu, Boyuan Chen, Huazhe Xu, Xiaolong Wang, Fei Fang, Simon Du, Yu Wang, and Yi Wu. Discovering diverse multi-agent strategic behavior via reward randomization. *arXiv preprint arXiv:2103.04564*, 2021.
- [42] Zhengyang Tang, Xingxing Zhang, Benyou Wang, and Furu Wei. Mathscale: Scaling instruction tuning for mathematical reasoning. *arXiv preprint arXiv:2403.02884*, 2024.
- [43] Kimi Team, Angang Du, Bofei Gao, Bowei Xing, Changjiu Jiang, Cheng Chen, Cheng Li, Chenjun Xiao, Chenzhuang Du, Chonghua Liao, et al. Kimi k1.5: Scaling reinforcement learning with llms. *arXiv preprint arXiv:2501.12599*, 2025.

- [44] Selim Furkan Tekin, Fatih Ilhan, Tiansheng Huang, Sihao Hu, and Ling Liu. Llm-topla: Efficient llm ensemble by maximising diversity. *arXiv preprint arXiv:2410.03953*, 2024.
- [45] Bryon Tjanaka, Matthew C Fontaine, Julian Togelius, and Stefanos Nikolaidis. Approximating gradients for differentiable quality diversity in reinforcement learning. In *Proceedings of the Genetic and Evolutionary Computation Conference*, pages 1102–1111, 2022.
- [46] Leandro von Werra, Younes Belkada, Lewis Tunstall, Edward Beeching, Tristan Thrush, Nathan Lambert, Shengyi Huang, Kashif Rasul, and Quentin Gallouédec. Trl: Transformer reinforcement learning. <https://github.com/huggingface/trl>, 2020.
- [47] Peiyi Wang, Lei Li, Zhihong Shao, RX Xu, Damai Dai, Yifei Li, Deli Chen, Yu Wu, and Zhifang Sui. Math-shepherd: Verify and reinforce llms step-by-step without human annotations. *arXiv preprint arXiv:2312.08935*, 2023.
- [48] Shuang Wu, Jian Yao, Haobo Fu, Ye Tian, Chao Qian, Yaodong Yang, Qiang Fu, and Yang Wei. Quality-similar diversity via population based reinforcement learning. In *The eleventh international conference on learning representations*, 2023.
- [49] Ting Wu, Xuefeng Li, and Pengfei Liu. Progress or regress? self-improvement reversal in post-training. *arXiv preprint arXiv:2407.05013*, 2024.
- [50] Xingyu Wu, Sheng-hao Wu, Jibin Wu, Liang Feng, and Kay Chen Tan. Evolutionary computation in the era of large language model: Survey and roadmap. *IEEE Transactions on Evolutionary Computation*, 2024.
- [51] An Yang, Beichen Zhang, Binyuan Hui, Bofei Gao, Bowen Yu, Chengpeng Li, Dayiheng Liu, Jianhong Tu, Jingren Zhou, Junyang Lin, Keming Lu, Mingfeng Xue, Runji Lin, Tianyu Liu, Xingzhang Ren, and Zhenru Zhang. Qwen2.5-math technical report: Toward mathematical expert model via self-improvement. *arXiv preprint arXiv:2409.12122*, 2024.
- [52] Hanlin Yang, Jian Yao, Weiming Liu, Qing Wang, Hanmin Qin, Kirk Tang, Jiechao Xiong, Chao Yu, Kai Li, Junliang Xing, et al. Diverse policies recovering via pointwise mutual information weighted imitation learning. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [53] Jian Yao, Weiming Liu, Haobo Fu, Yaodong Yang, Stephen McAleer, Qiang Fu, and Wei Yang. Policy space diversity for non-transitive games. *Advances in Neural Information Processing Systems*, 36:67771–67793, 2023.
- [54] Qiyang Yu, Zheng Zhang, Ruofei Zhu, Yufeng Yuan, Xiaochen Zuo, Yu Yue, Tiantian Fan, Gaohong Liu, Lingjun Liu, Xin Liu, et al. Dapo: An open-source llm reinforcement learning system at scale, 2025. URL <https://arxiv.org/abs/2503.14476>, 2025.
- [55] Yufeng Yuan, Qiyang Yu, Xiaochen Zuo, Ruofei Zhu, Wenyuan Xu, Jiaze Chen, Chengyi Wang, TianTian Fan, Zhengyin Du, Xiangpeng Wei, et al. Vapo: Efficient and reliable reinforcement learning for advanced reasoning tasks. *arXiv preprint arXiv:2504.05118*, 2025.
- [56] Yufeng Yuan, Yu Yue, Ruofei Zhu, Tiantian Fan, and Lin Yan. What’s behind ppo’s collapse in long-cot? value optimization holds the secret. *arXiv preprint arXiv:2503.01491*, 2025.
- [57] Yang Yue, Zhiqi Chen, Rui Lu, Andrew Zhao, Zhaokai Wang, Shiji Song, and Gao Huang. Does reinforcement learning really incentivize reasoning capacity in llms beyond the base model? *arXiv preprint arXiv:2504.13837*, 2025.
- [58] Tom Zahavy, Brendan O’Donoghue, Andre Barreto, Volodymyr Mnih, Sebastian Flennerhag, and Satinder Singh. Discovering diverse nearly optimal policies with successor features. *arXiv preprint arXiv:2106.00669*, 2021.
- [59] Weihao Zeng, Yuzhen Huang, Qian Liu, Wei Liu, Keqing He, Zejun Ma, and Junxian He. Simplerl-zoo: Investigating and taming zero reinforcement learning for open base models in the wild, 2025.

- [60] Weihao Zeng, Yuzhen Huang, Wei Liu, Keqing He, Qian Liu, Zejun Ma, and Junxian He. 7b model and 8k examples: Emerging reasoning with reinforcement learning is both effective and efficient. <https://hkust-nlp.notion.site/simpler1-reason>, 2025. Notion Blog.
- [61] Weihao Zeng, Yuzhen Huang, Lulu Zhao, Yijun Wang, Zifei Shan, and Junxian He. B-star: Monitoring and balancing exploration and exploitation in self-taught reasoners. *arXiv preprint arXiv:2412.17256*, 2024.
- [62] Chong Zhang, Yue Deng, Xiang Lin, Bin Wang, Dianwen Ng, Hai Ye, Xingxuan Li, Yao Xiao, Zhanfeng Mo, Qi Zhang, et al. 100 days after deepseek-r1: A survey on replication studies and more directions for reasoning language models. *arXiv preprint arXiv:2505.00551*, 2025.
- [63] 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.
- [64] Yunbo Zhang, Wenhao Yu, and Greg Turk. Learning novel policies for tasks. In *International Conference on Machine Learning*, pages 7483–7492. PMLR, 2019.
- [65] Yu Zhou, Xingyu Wu, Jibin Wu, Liang Feng, and Kay Chen Tan. HM3: Hierarchical multi-objective model merging for pretrained models. *arXiv preprint arXiv:2409.18893*, 2024.
- [66] Zihan Zhou, Wei Fu, Bingliang Zhang, and Yi Wu. Continuously discovering novel strategies via reward-switching policy optimization. *arXiv preprint arXiv:2204.02246*, 2022.
- [67] Yuxin Zuo, Kaiyan Zhang, Shang Qu, Li Sheng, Xuekai Zhu, Biqing Qi, Youbang Sun, Ganqu Cui, Ning Ding, and Bowen Zhou. Ttrl: Test-time reinforcement learning. *arXiv preprint arXiv:2504.16084*, 2025.

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: Yes, the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope.

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: Please refer to the Limitation Section.

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: The paper has some theoretical analysis. We provide a complete and correct analysis.

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: Please refer to Experiment Section and Appendix.

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: [Yes]

Justification: Our code and instructions are included in the supplementary material. The data we use for the experiments are all from open-access datasets.

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, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: Please refer to Experiment Section and Appendix.

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: [Yes]

Justification: Please refer to Experiment Section.

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: Please refer to the implementation details in the Appendix.

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: The research conducted in the paper conform with the NeurIPS 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 list potential positive societal impacts in the Appendix.

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: The paper poses no such risks.

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 properly credit data, paper, and ideas that we used in this paper.

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 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: [\[Yes\]](#)

Justification: We document well about the asset.

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: This paper does not involve crowdsourcing nor research with human subjects.

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: This paper does not involve crowdsourcing nor research 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.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

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

Justification: We describe the pipeline to fine-tune the LLM in the Experiment Section.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Theoretical analysis

A.1 Proof for Equation 8

The equation we want to prove is:

$$\begin{aligned} & \mathbb{E}_{q \sim \mathcal{Q}, o \sim \pi_{old}(\cdot|q)} \left[-\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\tilde{o}^t \sim \pi_{\theta}(\cdot|q, o^{<t})} [\log \pi_{\theta}(\tilde{o}^t|q, o^{<t})] \right] \\ &= \mathbb{E}_{q \sim \mathcal{Q}, o \sim \pi_{old}(\cdot|q)} \left[-\frac{1}{T} \sum_{t=1}^T \frac{\pi_{\theta}(o^t|q, o^{<t})}{\pi_{old}(o^t|q, o^{<t})} \log \pi_{\theta}(o^t|q, o^{<t}) \right]. \end{aligned} \quad (13)$$

Since T is a random variable that depends on π_{old} , the proof is not straightforward. We prove it in two stages.

(1). When T is fixed, the proof proceeds straightforwardly by examining each term in the summation. Note that $o^{<t}$ is sampled from π_{old} while $\tilde{o}^t$ is sampled from π_{θ} , hence:

$$\begin{aligned} & \mathbb{E}_{q \sim \mathcal{Q}, o \sim \pi_{old}(\cdot|q)} \left[-\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\tilde{o}^t \sim \pi_{\theta}(\cdot|q, o^{<t})} [\log \pi_{\theta}(\tilde{o}^t|q, o^{<t})] \right] \\ &= -\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{q \sim \mathcal{Q}, o \sim \pi_{old}(\cdot|q)} [\mathbb{E}_{\tilde{o}^t \sim \pi_{\theta}(\cdot|q, o^{<t})} [\log \pi_{\theta}(\tilde{o}^t|q, o^{<t})]] \\ &= -\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{q \sim \mathcal{Q}, o^{<t} \sim \pi_{old}(\cdot|q)} [\mathbb{E}_{\tilde{o}^t \sim \pi_{\theta}(\cdot|q, o^{<t})} [\log \pi_{\theta}(\tilde{o}^t|q, o^{<t})]] \\ &= -\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{q \sim \mathcal{Q}, o^{<t} \sim \pi_{old}(\cdot|q)} \left[\mathbb{E}_{o^t \sim \pi_{old}(\cdot|q, o^{<t})} \left[\frac{\pi_{\theta}(o^t|q, o^{<t})}{\pi_{old}(o^t|q, o^{<t})} \log \pi_{\theta}(o^t|q, o^{<t}) \right] \right] \\ &= -\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{q \sim \mathcal{Q}, o^{<t} \sim \pi_{old}(\cdot|q)} \left[\frac{\pi_{\theta}(o^t|q, o^{<t})}{\pi_{old}(o^t|q, o^{<t})} \log \pi_{\theta}(o^t|q, o^{<t}) \right] \\ &= -\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{q \sim \mathcal{Q}, o \sim \pi_{old}(\cdot|q)} \left[\frac{\pi_{\theta}(o^t|q, o^{<t})}{\pi_{old}(o^t|q, o^{<t})} \log \pi_{\theta}(o^t|q, o^{<t}) \right] \\ &= \mathbb{E}_{q \sim \mathcal{Q}, o \sim \pi_{old}(\cdot|q)} \left[-\frac{1}{T} \sum_{t=1}^T \frac{\pi_{\theta}(o^t|q, o^{<t})}{\pi_{old}(o^t|q, o^{<t})} \log \pi_{\theta}(o^t|q, o^{<t}) \right]. \end{aligned} \quad (14)$$

The second and fifth equations hold because we may add or remove any random variables that are not in the target expectation. The third equation results from applying importance sampling to reweight probabilities.

(2). For the case that T is a random variable, roughly, the idea is to apply the law of total probability:

$$\begin{aligned} & \mathbb{E}_{q \sim \mathcal{Q}, o \sim \pi_{old}(\cdot|q)} \left[-\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\tilde{o}^t \sim \pi_{\theta}(\cdot|q, o^{<t})} [\log \pi_{\theta}(\tilde{o}^t|q, o^{<t})] \right] \\ &= \mathbb{E}_{q \sim \mathcal{Q}} \left[\sum_{T_0=0}^{\infty} P(T = T_0) \mathbb{E}_{o \sim \pi_{old}(\cdot|q, T=T_0)} \left[-\frac{1}{T_0} \sum_{t=1}^{T_0} \mathbb{E}_{\tilde{o}^t \sim \pi_{\theta}(\cdot|q, o^{<t})} [\log \pi_{\theta}(\tilde{o}^t|q, o^{<t})] \right] \right] \end{aligned} \quad (15)$$

And apply case (1) to finish the proof.

B More discussion

More discussion about Potential@k

The definition of Potential@k aims to quantify the performance improvement achievable through RL training for LLMs. By examining its formulation, we derive:

$$\text{Potential@k} := \frac{\sum_{i=1}^N \text{Pass@k}(q_i) \cdot (1 - \text{Pass@1}(q_i))}{\sum_{i=1}^N (1 - \text{Pass@1}(q_i))} \approx \sum_{i=1}^N [\text{Pass@k}(q_i) - \text{Pass@1}(q_i)]. \quad (16)$$

This metric essentially captures the discrepancy between Pass@k and Pass@1. While Pass@k is often treated as the performance boundary for RL training on LLM [57], our Potential@k specifically measures the performance gain from RL training, approximated by subtracting Pass@1 (a measure for initial performance) from Pass@k.

For each question q_i before training begins, if $\text{Pass@1}(q_i) = 1$, the question is already mastered with no improvement potential. When $\text{Pass@1}(q_i) = 0$ but $\text{Pass@k}(q_i) = 1$, GRPO training uses positive samples from k trials to teach the correct response. If both $\text{Pass@1}(q_i) = 0$ and $\text{Pass@k}(q_i) = 0$, the question provides no training signal as it remains unsolved. Hence, our definition of Potential@k focuses training on questions with partial capability, excluding both mastered and unsolvable questions, thereby capturing the true learning potential through the Pass@k to Pass@1 performance gap.

Why the performance of our reproduction of R1-zero is worse than the state-of-the-art methods reproductions (e.g. SimpleRL-Zoo)?

We believe the performance gap between our R1-zero reproduction and SimpleRL-Zoo’s implementation stems primarily from resource constraints. Our experiments were conducted on a modest $8 \times \text{A6000}$ GPUs setup, necessitating several efficiency optimizations: we employed the simpler GSM8K dataset, constrained generation lengths (appropriate for GSM8K’s short responses). In contrast, SimpleRL-Zoo utilized significantly more powerful $2 \times 8 \text{ H100-80G}$ GPUs, trained on more complex datasets with longer response lengths. Importantly, our study’s primary objective was not to surpass SimpleRL-Zoo’s results, but rather to demonstrate that our diversity-enhanced method outperforms standard R1-zero. Our method of independence can be applied to enhance the SimpleRL-Zoo and other state-of-the-art methods.

Why are the 4 mathematical datasets chosen?

We require the number of data points in the test dataset to be at least 500. Since we find that the results are unstable when we test on a small dataset. In some recent work, they report by sampling many times and calculate the average. However, we think the distribution shift issue still exists (i.e., the small test data may biasly represent the hard/medium/easy-level benchmark)

C Broader impacts

Our diversity-aware RL approach for LLM reasoning offers valuable benefits for AI applications. First, in education, generating multiple valid reasoning paths could enhance AI tutoring systems by providing alternative solution strategies to students. Second, for scientific research, the improved ability to explore diverse reasoning approaches may aid in hypothesis generation and problem-solving where multiple perspectives are valuable.

D Implementation details

We provide more details for experiments in Section 5.

D.1 Experiment environment

For training R1-zero and R1-zero-Div, the codebase runs on Python 3.11, utilizing TRL 0.16.0 [46] with PyTorch 2.5.1. We employ DeepSpeed [39] for distributed training and incorporate vLLM 0.7.2 [27] for efficient rollout, all deployed on $8 \times \text{NVIDIA A6000}$ GPUs. Each experiment runs for 3 days. For other baselines, we evaluate open-sourced models downloaded from Hugging Face⁴

⁴<https://huggingface.co>

A conversation between User and Assistant. The user asks a question, and the Assistant solves it. The assistant first thinks about the reasoning process in the mind and then provides the user with the answer. The assistant should 1) Identify core concepts and required formulas. 2) Break down solutions into logical, numbered steps. 3) Verify results using alternative methods or substitutions. Put your final answer within `\boxed{}`

Figure 2: System prompt

For evaluation, we utilize the code from Qwen2.5-Math.⁵ To calculate Pass@1, we use greedy decoding for our models and baselines, except for SimpleRL-Zoo [59], which we evaluate using temperature=1 and top-p=0.95 as suggested in their paper.

D.2 Hyperparameter settings

We provide the system prompt in Figure 2 and other detailed hyperparameter settings in Table 6. Due to computational resource constraints, we train on the simpler dataset (GSM8K), which allows for a shorter maximum response length, and use a well-designed prompt to obtain a stronger initial checkpoint. The experiment settings for R1-zero and R1-zero-Div are the same except for $\lambda = 0$ in R1-zero and $\lambda = 0.01$ in R1-zero-Div.

Table 6: Hyperparameter settings

Hyperparameter	Value
<i>General settings</i>	
dataset	GSM8K
max prompt length	256
max completion length	756
num generations	6
use vllm	true
vllm gpu memory utilization	0.5
torch dtype	bfloat16
learning rate	3.0e-06
lr scheduler type	cosine
beta	0.0001
zero stage	2
offload optimizer device	CPU
offload param device	none
distributed type	DEEPSPEED
<i>Base model: Qwen/Qwen2.5-Math-7B</i>	
num train epochs	2
per device train batch size	1
gradient accumulation steps	64
<i>Base model: Qwen/Qwen2.5-Math-1.5B</i>	
num train epochs	3
per device train batch size	6
gradient accumulation steps	16

⁵<https://github.com/QwenLM/Qwen2.5-Math>



Figure 3: Entropy during the RL training

Table 7: Avg@8 on GPQA Diamond dataset

Method	GPQA Diamond
Qwen2.5-Math-7B	28.5 (0.64)
R1-zero	32.7 (0.70)
R1-zero-Div (Ours)	34.3 (0.79)
SimpleRL-Zoo	33.9 (0.56)
Eurus-2-7B-PRIME	34.0 (0.74)

E More experiment results

E.1 Entropy during the RL training

We analyze the entropy dynamics during training. As shown in the Figure 3, when $\lambda = 0$ (i.e., the baseline R1-zero method), entropy collapses rapidly, indicating a loss of exploration. When applying the diversity objective with $\lambda = 0.01$ to all samples (both positive and negative), entropy exhibits late-stage exploration during training; however, this tends to degrade model quality, as reflected by the final performance in Table 4. In contrast, applying the diversity objective with $\lambda = 0.01$ exclusively to positive samples strikes a better balance between quality and diversity, yielding the optimal final performance.

E.2 Evaluation on GPQA

We present the avg@8 results on the GPQA Diamond dataset in Table 7, where we directly evaluated the models listed in Table 1. These results demonstrate the generalizability of our method.

E.3 Solution diversity in R1-zero-Div

We show the case where our method generates diverse solutions in Figure 4. The text outputs are generated by our model (R1-zero-Div training on Qwen2.5-Math-7B), and we render LaTeX expressions for better visualization. As illustrated in Figure 4, the model generates two distinct yet mathematically valid solutions to the problem. The first solution demonstrates elegant mathematical

Question: There exist real numbers x and y , both greater than 1, such that $\log_x(y^x) = \log_y(x^{4y}) = 10$. Find xy .

Solution 1:

Given the equations:

$$\log_x(y^x) = 10 \quad \text{and} \quad \log_y(x^{4y}) = 10$$

1. First, we simplify the equations using logarithmic properties:

$$\begin{aligned} \log_x(y^x) = 10 &\implies x \cdot \log_x(y) = 10 \implies \log_x(y) = \frac{10}{x}, \\ \log_y(x^{4y}) = 10 &\implies 4y \cdot \log_y(x) = 10 \implies \log_y(x) = \frac{10}{4y} = \frac{5}{2y} \end{aligned}$$

2. Using the change of base formula, we know that $\log_x(y) = \frac{1}{\log_y(x)}$. Therefore:

$$\begin{aligned} \log_x(y) = \frac{10}{x} \quad \text{and} \quad \log_y(x) = \frac{5}{2y} \\ \frac{10}{x} = \frac{1}{\frac{5}{2y}} \implies \frac{10}{x} = \frac{2y}{5} \implies 10 \cdot 5 = x \cdot 2y \implies 50 = 2xy \implies xy = 25 \end{aligned}$$

Thus, the final answer is: 25

Solution 2:

Given the equations:

$$\log_x(y^x) = 10 \quad \text{and} \quad \log_y(x^{4y}) = 10$$

we can use the properties of logarithms to simplify these equations.

1. Simplifying the first equation:

$$\begin{aligned} \log_x(y^x) = 10 &\implies x \log_x(y) = 10 \implies \log_x(y) = \frac{10}{x} \\ \text{Using the change of base formula } \log_x(y) &= \frac{\log(y)}{\log(x)}, \text{ we get: } \frac{\log(y)}{\log(x)} = \frac{10}{x} \\ \text{This implies: } \log(y) &= \frac{10 \log(x)}{x} \end{aligned}$$

2. Simplifying the second equation:

$$\begin{aligned} \log_y(x^{4y}) = 10 &\implies 4y \log_y(x) = 10 \implies \log_y(x) = \frac{10}{4y} = \frac{5}{2y} \\ \text{Using the change of base formula } \log_y(x) &= \frac{\log(x)}{\log(y)}, \text{ we get: } \frac{\log(x)}{\log(y)} = \frac{5}{2y} \\ \text{This implies: } \log(x) &= \frac{5 \log(y)}{2y} \end{aligned}$$

Now we have two equations:

$$\log(y) = \frac{10 \log(x)}{x}, \quad \log(x) = \frac{5 \log(y)}{2y}$$

Substituting the second equation into the first:

$$\log(y) = \frac{10 \left(\frac{5 \log(y)}{2y} \right)}{x} = \frac{25 \log(y)}{xy}$$

Dividing both sides by $\log(y)$ (assuming $\log(y) \neq 0$):

$$1 = \frac{25}{xy} \implies xy = 25$$

Thus, the value of xy is: 25

Figure 4: Solution generated by R1-zero-Div

insight by utilizing the reciprocal property of logarithmic functions. The second solution adopts a more systematic approach, first establishing two equations for $\log(y)$ and $\log(x)$ (more precisely, $\lg y$ and $\lg x$). And then solve these two equations. Both methods ultimately arrive at the correct solution, showcasing the model's ability to discover multiple valid reasoning paths.