

Group Adaptive Clipping Policy Optimization

Sheng Jia^{1,2}, Xiao Wang², Shiva Prasad Kasiviswanathan^{2†}, Rein Houthoofd^{2‡}

¹University of Toronto, ²Amazon

{shengama, waxao, kasivisw, reinrh}@amazon.com

Abstract

Group relative policy optimization for reinforcement learning with verifiable rewards (RLVR) typically uses a fixed importance-sampling (IS) ratio clipping boundary across all rollouts. We identify a key limitation: rare correct rollouts on harder problems and abundant correct rollouts on easier problems are clipped at comparable rates, despite contributing very different learning signals. Rollouts with low group success exhibit larger IS ratios and carry stronger gradient signal for exploration and solving new problems, yet are disproportionately suppressed by fixed clipping.

To address this, we propose *Group Adaptive Clipping Policy Optimization (GAPO)*, a plugin modification to GRPO methods that adapts the clipping boundary to the rollout advantage. GAPO is motivated by a reverse-KL trust-region perspective, which suggests that rollouts with larger learning signal should receive proportionally greater update headroom. GAPO requires no reward shaping and preserves the standard PPO/GSPO surrogate while adapting only the clipping threshold. Across Qwen and Llama models, GAPO consistently improves both Pass@1 and Pass@k over fixed clipping and advantage-shaping baselines on math reasoning and coding benchmarks where the pass rates by the base model are relatively low¹.

1 Introduction

Reinforcement learning with verifiable rewards (RLVR) trains language models using PPO-style clipped policy optimization (Schulman et al., 2017), where a fixed importance-sampling (IS) ratio boundary approximates a trust region by suppressing updates once the policy moves too far from its previous iterate. In group-relative RLVR methods

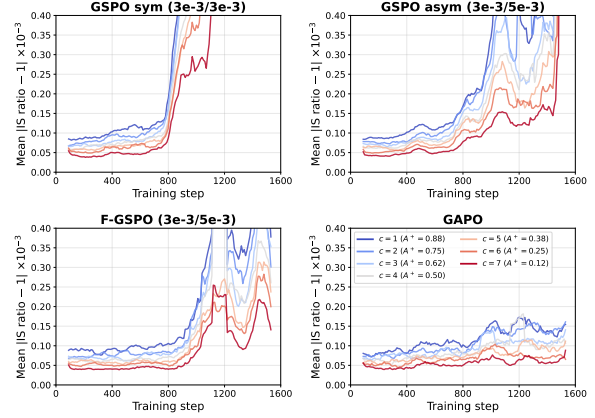


Figure 1: **Mean deviation of the importance-sampling (IS) ratio over RLVR training, binned by correctness count c within a group.** Two patterns hold across algorithms: (1) the IS ratio is consistently higher for scarcely correct rollouts (small c , fewer correct solutions per group), and (2) for these rollouts the ratio continues to grow at a rate equal to or exceeding that of rollouts on easier problems (large c). Together these patterns motivate adapting the clip boundary to question difficulty: under a fixed boundary, the rare high-advantage rollouts on hard problems hit the clip first and are suppressed disproportionately early.

such as GRPO (Shao et al., 2024), GSPO (Zheng et al., 2025), and DAPO (Yu et al., 2025), rollouts are assigned advantages relative to other samples from the same prompt. Specifically, in a group of k rollouts with c correct solutions, correct rollouts receive advantage $A_i = (k - c)/k$, creating a structured spectrum ranging from scarce, high-advantage correct rollouts (low c) to abundant, low-advantage correct rollouts (high c).

A fixed clipping boundary treats this spectrum uniformly, despite substantial differences in learning signal. We observe that before clipping activates, IS ratios grow proportionally to advantage: scarce correct rollouts on harder problems move away from the reference policy substantially faster than abundant correct rollouts on easier prob-

[†] Equal advising.

¹Code available: <https://github.com/Sheng-J/GAPO>. Correspondence to: shengama@amazon.com

lems (Figure 1). Yet once clipping begins, both are truncated at comparable rates under a uniform boundary. This creates a mismatch: scarce correct rollouts, which provide valuable signal for exploration and solving new problems, are clipped too aggressively, while abundant redundant rollouts are granted equal update headroom despite contributing less to improvement.

A natural response might be to widen the clipping boundary uniformly, as explored in DAPO (Yu et al., 2025). However, this fails to address the core issue: the problem is not the absolute width of the clipping boundary, but its uniformity across rollouts with fundamentally different advantages. What must change is the *relative* clip width across different levels of group success.

Under a reverse-KL trust-region perspective, the optimal policy update at a single prompt allocates probability to responses proportionally to their advantage, yielding a trust-region-optimal importance-sampling (IS) ratio that scales exponentially with advantage (Chen et al., 2018). This suggests a simple principle for clipping: rollouts with larger learning signal should receive proportionally greater update headroom. Pushing an IS ratio substantially beyond this optimum either violates the trust region or over-allocates probability mass to a single response at the expense of alternatives, making the optimal ratio a natural guide for adaptive clipping.

We leverage this observation to introduce *Group Adaptive Policy Optimization (GAPO)*, a simple plug-in modification to group-relative policy optimization that adapts the per-rollout clip threshold to rollout advantage. In RLVR, binary rewards induce only a small number of discrete positive advantage levels, allowing the adaptive clipping threshold to be computed directly from the group success statistic c . GAPO preserves the standard PPO/GSPO surrogate objective and modifies only the clipping boundary. Unlike reward- or advantage-shaping approaches, GAPO continues to optimize pass@1 directly and does not exhibit the pass@1 drift that can arise from shaping objectives based on difficulty (Plyusov et al., 2026), inference-time pass@ k (Walder and Karkhanis, 2025; Chen et al., 2025b; Tang et al., 2025), or diversity (Li et al., 2025). We make the following contributions:

- We identify a *clipping asymmetry* in RLVR: under uniform clip boundaries, scarce correct rollouts (low c , high advantage) are clipped

at rates comparable to abundant correct rollouts, despite carrying substantially stronger learning signal (Figure 3). We show that a reverse-KL trust-region perspective (Equation 5, 6) naturally motivates *advantage-dependent clipping*, yielding a trust-region-optimal IS ratio that scales exponentially with advantage and reduces to a practical closed-form clipping rule in RLVR through the group statistic c (Equation 8).

- We introduce Group Adaptive Clipping Policy Optimization (GAPO), a simple plug-in modification to PPO/GSPO that adapts only the clipping boundary while preserving the standard surrogate objective and direct optimization of pass@1 (Equation 11).
- We demonstrate consistent improvements in pass@1 and pass@ k across Qwen2.5-Math-1.5B, Llama-3.2-3B-Instruct, and DeepSeek-R1-Distill-Qwen-1.5B on mathematical reasoning and code generation benchmarks (Table 1, 2, 3, Figure 4), while maintaining strong correlation between IS ratio and advantage throughout training (Figure 5). Figure 7 shows a checkpoint intervention experiment to rule out common confounds.

2 Preliminaries

2.1 RL with Verifiable Rewards

We consider reinforcement learning with verifiable rewards (RLVR) for language model post-training. Given a prompt x drawn from a dataset $\mathcal{D}$, the policy π_θ generates a complete response $y \sim \pi_\theta(\cdot | x)$, which receives a binary verifiable reward $R(x, y) \in \{0, 1\}$ from an automated verifier. The training objective is the expected reward

$$J(\theta) = \mathbb{E}[R(x, y)], \quad x \sim \mathcal{D}, y \sim \pi_\theta(\cdot | x). \quad (1)$$

In group-relative policy optimization with k rollouts per prompt $\{y_i\}_{i=1}^k \sim \pi_{\theta_{\text{old}}}(\cdot | x)$, a correct rollout i in a group with c correct has advantage

$$A_i = r_i - \bar{r} = 1 - \frac{c}{k} = \frac{k - c}{k}, \quad (2)$$

where $r_i \in \{0, 1\}$. Scarce correct rollouts ($c=1$) receive large advantage $(k-1)/k$, while abundant correct rollouts ($c=k-1$) receive small advantage $1/k$. In this work, we do not perform advantage

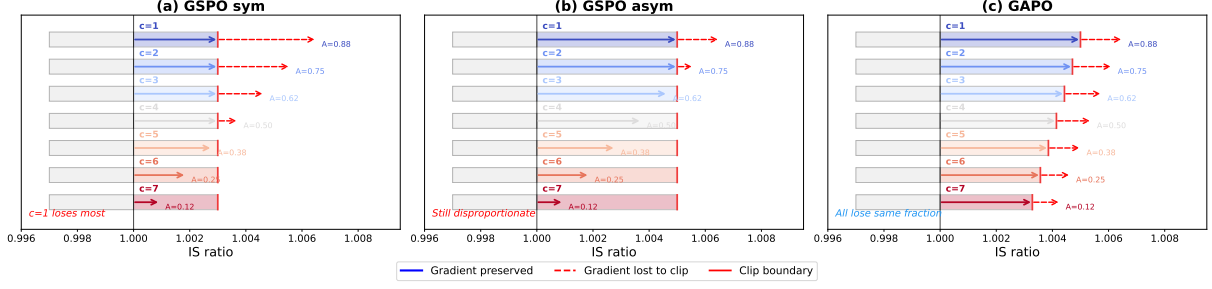


Figure 2: **(a) Symmetric clipping (b) Asymmetric Clipping (c) Adaptive Clipping.** Each row represents a correct rollout with group statistic c and advantage A_i . Solid arrows show preserved gradient; dashed red arrows show gradient lost to IS surpassing clipping thresholds. (a) GSPO symmetric: a uniform boundary clips scarce correct rollouts disproportionately. (b) GSPO asymmetric: wider upper boundary but still uniform across c . (c) GAPO: per- c adaptive boundary equation 11 keeps gradient updates proportional to advantage equation 8 even as clipping starts firing.

normalization for unbiased advantage estimate (Liu et al., 2025).

Standard PPO-style clipping applies a fixed trust region $[1-\epsilon, 1+\epsilon]$ uniformly to all rollouts. We observe that this disproportionately suppresses high-advantage (scarce) rollouts, which push their importance ratio past the clip boundary fastest.

2.2 Trust Region Policy Optimization

Trust regions and reverse KL. We build on the policy-improvement guarantee of Schulman et al. (2015). In the RLVR setting the bound takes the form

$$J(\theta) \geq L_{\pi_{\theta_{\text{old}}}}(\pi_{\theta}) - C \cdot \max_x D_{\text{KL}}(\pi_{\theta_{\text{old}}}(\cdot | x) \| \pi_{\theta}(\cdot | x)), \quad (3)$$

where the surrogate, written in importance-sampling form, is

$$\begin{aligned} L_{\pi_{\theta_{\text{old}}}}(\pi_{\theta}) &= \widehat{\mathbb{E}}_i \left[\frac{\pi_{\theta}(y_i | x)}{\pi_{\theta_{\text{old}}}(y_i | x)} A_i \right] \\ &= \widehat{\mathbb{E}}_i [\rho_i A_i], \end{aligned} \quad (4)$$

with $\widehat{\mathbb{E}}_i[\cdot]$ denoting the empirical expectation over a batch of rollouts $\{(x, y_i)\}$ drawn from $\pi_{\theta_{\text{old}}}(\cdot | x)$, and C depending only on the reward range. The bound follows from Pinsker’s inequality applied to the worst-case total-variation distance between policies; maximising its right-hand side guarantees monotonic improvement of $J(\theta)$ at each update.

Equation 3 is stated with the forward direction $D_{\text{KL}}(\pi_{\theta_{\text{old}}} \| \pi_{\theta})$, but the same guarantee holds with the reverse direction (Chen et al., 2018). The reason is symmetry: because $D_{\text{TV}}(P, Q) = D_{\text{TV}}(Q, P)$,

Pinsker’s inequality bounds the same quantity from either side, so $D_{\text{KL}}(\pi_{\theta} \| \pi_{\theta_{\text{old}}})$ may replace $D_{\text{KL}}(\pi_{\theta_{\text{old}}} \| \pi_{\theta})$ in equation 3 without weakening the guarantee. The constrained optimization we solve is therefore

$$\begin{aligned} \max_{\theta} \quad & \widehat{\mathbb{E}}_i[\rho_i A_i] \\ \text{s.t.} \quad & \widehat{\mathbb{E}}_i[D_{\text{KL}}(\pi_{\theta}(\cdot | x) \| \pi_{\theta_{\text{old}}}(\cdot | x))] \leq \delta, \end{aligned} \quad (5)$$

where the empirical expectation over the batch replaces the worst-case maximum for tractability, as in PPO (Schulman et al., 2017).

3 Group Adaptive Policy Optimization

3.1 Per-Prompt Trust Region Optimization

Following Chen et al. (2018), we consider local policy optimization at a state s , which is a prompt x in our case. This makes the batch-level constraint in equation 5 per-prompt, exact only for single-prompt batches (see Limitations). The Lagrangian of maximizing equation 5 given x is

$$\begin{aligned} \mathcal{L}_x &= \sum_y \pi_{\theta}(y|x) A(x, y) \\ &\quad - \lambda D_{\text{KL}}(\pi_{\theta}(\cdot | x) \| \pi_{\theta_{\text{old}}}(\cdot | x)). \end{aligned} \quad (6)$$

By solving the following Euler–Lagrange equation (Gelfand and Fomin, 2000),

$$\begin{aligned} \frac{\partial \mathcal{L}_x}{\partial \pi_{\theta}(y|x)} &= A(x, y) - \lambda \left(\log \frac{\pi_{\theta}(y|x)}{\pi_{\theta_{\text{old}}}(y|x)} + 1 \right) \\ &= 0. \end{aligned} \quad (7)$$

we obtain the stationary point for the target policy $\pi_{\theta}^*(y|x) \propto \pi_{\theta_{\text{old}}}(y|x) \exp\left(\frac{A(x, y)}{\lambda}\right)$. At each

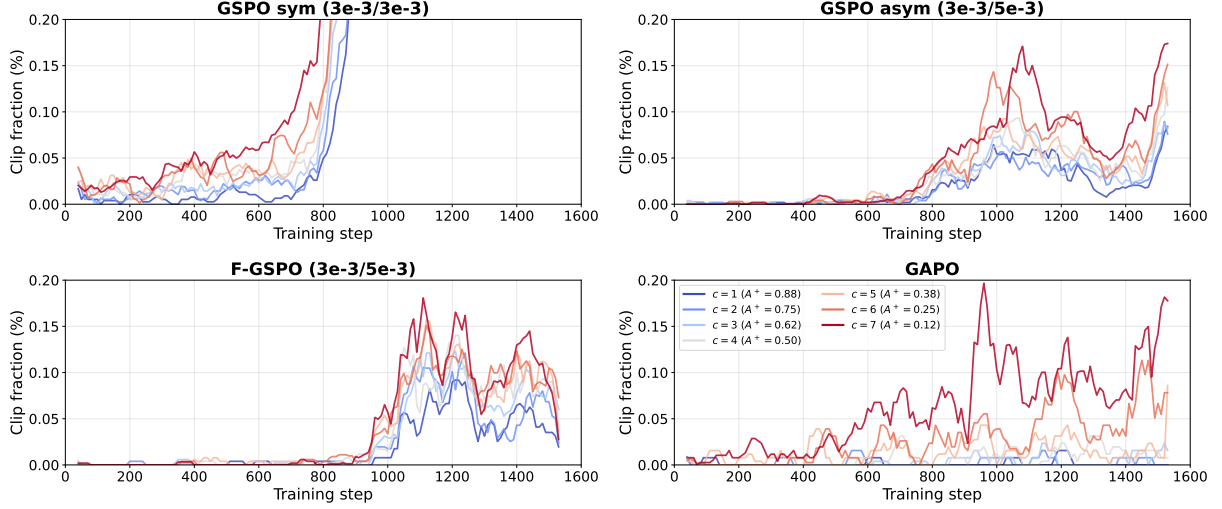


Figure 3: **Token clip fraction over training of Qwen2.5-Math-1.5B, by group correctness c .** Fraction of tokens clipped *within* rollouts at each c . Under fixed clipping (GSPO sym / asym), scarce correct rollouts on harder questions (low c , larger advantage A_i equation 2) are clipped at rates comparable to abundant correct rollouts (high c). GAPO clips high- c rollouts more and low- c rollouts less, retaining gradient on the rare correct rollouts most valuable for exploration.

rollout i for prompt x , the optimal IS ratio is

$$\rho_i^* = \frac{\pi_{\theta}^*(y_i|x)}{\pi_{\theta_{\text{old}}}(y_i|x)} \propto \exp(A_i/\lambda). \quad (8)$$

Interpretation on the optimal IS ratio. ρ_i^* describes how much the probability of generating response y_i should be pushed up for prompt x to maximize the expected return on prompt x . Pushing ρ_i past ρ_i^* means either the trust region constraint is violated, or probability of generating some other response y_j is not pushed according to equation 8, lowering the expected return under this formulation. The clip should therefore fire at ρ_i^* , giving

$$\epsilon_{\text{hi}}(i) \geq \rho_i^* - 1 = \exp\left(\frac{A_i}{\lambda}\right) - 1. \quad (9)$$

GSPO uses clip widths $\epsilon \sim 10^{-3}$, and GAPO inherits this scale, so $\epsilon_{\text{hi}} \sim 10^{-3}$. Inverting $\epsilon_{\text{hi}} = \exp(A_i/\lambda) - 1$ gives $|A_i/\lambda| = \log(1 + \epsilon_{\text{hi}}) \approx 10^{-3}$, hence $\lambda \sim A_{\text{max}}/10^{-3} \sim 10^3 \gg A_{\text{max}}$ (with $A_{\text{max}} = (k-1)/k \leq 1$). The linearization $\exp(A_i/\lambda) \approx 1 + A_i/\lambda$ then gives

$$\epsilon_{\text{hi}}(i) \approx \frac{A_i}{\lambda} = \frac{k - c_i}{k\lambda}. \quad (10)$$

The linearization is only for presentation cleanliness. At the sequence-IS level it is essentially exact. At the token-IS scale $\epsilon_{\text{hi}}^{\text{max}}=0.28$, it is conservative (smaller ϵ_{hi} than exact, so still within the trust region), preserves the monotonic ordering in c , and keeps the relative error on A_i/λ bounded (2–13%).

3.2 GAPO Adaptive Clip Formula

Normalising to interpolate between ϵ_{lo} (minimum, for $c = k$) and $\epsilon_{\text{hi}}^{\text{max}}$ (maximum, for $c = 1$), we obtain the GAPO adaptive upper clip:

$$\epsilon_{\text{hi}}(c) = \epsilon_{\text{lo}} + (\epsilon_{\text{hi}}^{\text{max}} - \epsilon_{\text{lo}}) \cdot \frac{k - c}{k - 1} \quad (11)$$

with boundaries $\epsilon_{\text{hi}}(c=1) = \epsilon_{\text{hi}}^{\text{max}}$, $\epsilon_{\text{hi}}(c=k) = \epsilon_{\text{lo}}$. Scarce correct rollouts ($c=1$) receive maximum headroom, while abundant correct rollouts ($c \rightarrow k$) are constrained to the minimum. Incorrect rollouts always use ϵ_{lo} (the upper clip is irrelevant for negative advantage). For rollouts with $A_i < 0$ where ϵ_{lo} can become active, the update scales with $|A| = \frac{c}{k}$, which is small on hard problems (low c), precisely where correct rollouts are scarce. Adaptive ϵ_{lo} matters more for the trust-region of penalizing incorrect rollouts on easy problems (high c). It affects neither the positive signal nor the clipping bias that prevents exploration (Yu et al., 2025), so we fix ϵ_{lo} .

Combined with GSPO’s sequence-level importance ratio $s_i(\theta)$, the full GAPO objective is

$$\mathcal{L}^{\text{GAPO}} = \mathbb{E} \left[\max \left(-A_i \cdot s_i(\theta), -A_i \cdot \text{clip}(s_i(\theta), 1 - \epsilon_{\text{lo}}, 1 + \epsilon_{\text{hi}}(c_i)) \right) \right], \quad (12)$$

where the sequence-level ratio is the geometric

mean of per-token ratios:

$$s_i(\theta) = \left(\frac{\pi_\theta(y_i | x)}{\pi_{\theta_{\text{old}}}(y_i | x)} \right)^{1/|y_i|}. \quad (13)$$

Both GAPO’s adaptive clip and GSPO’s sequence-level ratio operate at the rollout level, and the clip threshold directly gates the sequence-level ratio derived from the same per-rollout trust region argument.

4 Experiments & Results

We first compare the training dynamics of GAPO on Qwen2.5-Math-1.5B with three fixed clipping baselines. Specifically, we are interested in (i) how clip fraction within rollouts at each c differs between fixed and adaptive clipping (Figure 3), (ii) which clipping schedule maintains a high correlation between the empirical IS ratio and advantage (Figure 5), and (iii) whether this correlation translates into higher pass@1 and pass@ k (Figure 4). We then benchmark GAPO against representative RLVR methods on two base models without SFT (Qwen2.5-Math-1.5B, Llama-3.2-3B-Instruct) in Table 3 and one reasoning-distilled base model (DeepSeek-R1-Distill-Qwen-1.5B) in Table 1, 2.

4.1 Experiment Setup

Models and datasets. We train and evaluate three base models spanning two model families and two pretraining regimes: Qwen2.5-Math-1.5B (math-pretrained) (Yang et al., 2024), Llama-3.2-3B-Instruct (general-purpose instruction-tuned) (Grattafiori et al., 2024), and DeepSeek-R1-Distill-Qwen-1.5B (reasoning-distilled). We include this post-SFT model, since their RLVR usually has a different training dynamic, and we aim to show GAPO works for all settings. Math RL runs train on the DeepScaleR dataset with 39,202 samples after filtering duplicates (Tan et al., 2026). For the RL runs with code generation data in Table 2, we use 24,269 samples from DeepCoder (Luo et al., 2025) after filtering overlong prompts: 16,238 from SYNTHETIC-1 (Mattern et al., 2025), 7,432 from TACO (Li, 2024), and 599 from LiveCodeBench 2023-5-1 to 2024-7-31 (Jain et al., 2025).

Training details. Our main method builds on GSPO’s sequence-level importance ratio $s_i(\theta)$ equation 13, with clip boundaries $(\epsilon_{\text{lo}}, \epsilon_{\text{hi}}^{\text{max}}) = (3e-3, 5e-3)$ for Qwen2.5-Math-1.5B experiments and $(\epsilon_{\text{lo}}, \epsilon_{\text{hi}}^{\text{max}}) = (7e-5, 3e-4)$ for DS-R1-distilled models, as distilled models

have small IS drifts. For token-IS, we use clip boundaries $\epsilon_{\text{lo}} = 0.2$ and $\epsilon_{\text{hi}}^{\text{max}} = 0.28$. For RLVR with code generation data, we follow DeepCoder recipe (Luo et al., 2025) and also use token-IS. We use group size $k = 8$. Loss is aggregated at the token level (sum of per-token losses divided by total token count). We do not normalise the advantage by $\text{std}(\{R(q, o_1), \dots, R(q, o_k)\})$, following Dr.GRPO (Liu et al., 2025); this std normalisation introduces a question-level difficulty bias, upweighting groups whose rewards are nearly all 1 or 0. Up-weighting groups with many rewards=1 can counteract updates with rare correct ones. Full hyperparameters are listed in Table 4, 5.

Baselines. We compare against fixed-clip RLVR methods at matched clip widths: GRPO (Shao et al., 2024), Dr.GRPO (Liu et al., 2025), GSPO with symmetric and asymmetric clip ranges (Zheng et al., 2025), and the focal-shaping variants F-GRPO (Plyusov et al., 2026) and F-GSPO. These apply advantage shaping $\tilde{A}_i = (1 - c/k)^\gamma A_i$ with $\gamma \geq 0$. This selection spans the standard group-relative baselines and the two main alternative directions for adapting the surrogate (advantage shaping and asymmetric clipping) for exploration. The symmetric clipping $(\epsilon_{\text{lo}}=3e-3, \epsilon_{\text{hi}}^{\text{max}}=3e-3)$ baseline is also chosen to show that GAPO’s improvement comes not from clipping tighter, but from the relative difference in clip width across rollouts with different group success rates. We also evaluate a token-IS variant of GAPO against token-IS baselines. Among non-adaptive methods, GSPO’s sequence-level IS is the strongest baseline in our setting, so we adopt it as the base for both GAPO and the advantage-shaping baseline F-GSPO.

Evaluation. We evaluate on standard mathematical reasoning benchmarks (AIME24, AIME25 (Ye et al., 2025), AMC, MATH500 (Hendrycks et al., 2021), Minerva (Lewkowycz et al., 2022), OlympiadBench (He et al., 2024)) and code generation benchmarks (LiveCodeBench (Jain et al., 2025), HumanEval+ (Liu et al., 2023)), and we report IFEval (Zhou et al., 2023) as an out-of-domain probe for instruction-following retention. For Qwen2.5-Math-1.5B (Yang et al., 2024) and Llama-3.2-3B-Instruct (Grattafiori et al., 2024), we report pass@1 and pass@256 with temperature 1.0, top- $p = 1$, $n = 256$ samples per prompt, and T_{max} of 3092 and 8192 respectively. For DeepSeek-R1-Distill-Qwen-1.5B we report pass@1 and pass@16 with temperature 0.6, top- $p = 1$, $n = 16$, and $T_{\text{max}} =$

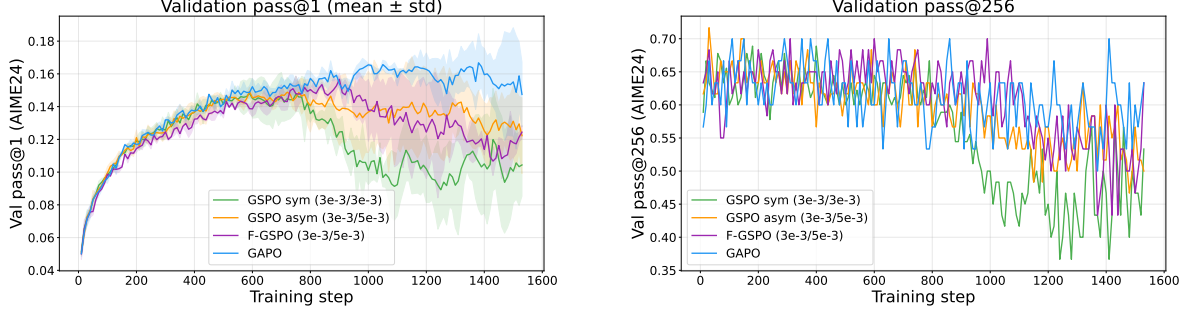


Figure 4: **Validation pass@1 and pass@256 on AIME24 over training.** Qwen2.5-Math-1.5B under four RLVR algorithms (mean $\pm$ std across seeds). GAPO sustains higher pass@1 and pass@256 in late training, while fixed-clip baselines (especially symmetric GSPO) show pass@256 collapse after step ~ 1000 . The diversity loss adaptive clipping is designed to avoid by retaining gradient on low- c rollouts (Figures 3, 5).

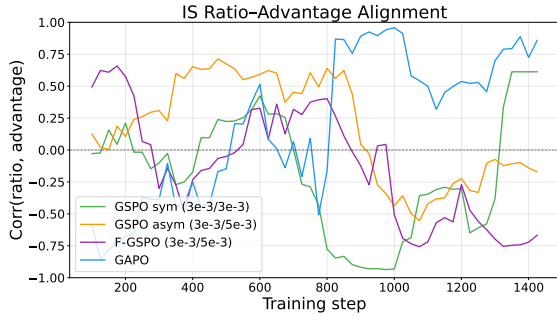


Figure 5: **IS-advantage correlation over training.** Windowed Pearson correlation between per- c IS-ratio deviation and advantage ($W = 200$ steps). Positive correlation reflects the trust-region-optimal relationship $\rho_i^* \propto \exp(A_i/\lambda)$ equation 8. GAPO maintains $r > 0.8$ in late training; fixed-clip baselines degrade toward negative correlation.

24576 to accommodate longer rollouts.

4.2 Adaptive clipping preserves the IS-advantage correlation

We compare GAPO against three fixed-clip baselines on Qwen2.5-Math-1.5B across clip fraction (Fig. 3), IS-advantage correlation (Fig. 5), and validation pass@1 / pass@256 (Fig. 4).

Under the reverse-KL trust region at a single prompt, the optimal target IS ratio equation 8 is proportional to advantage (to first order). We do not directly optimise toward this target; we instead clip proportional to advantage. Even so, Figure 5 shows GAPO retains a high correlation between the importance-sampling ratio and advantage throughout training, whereas fixed clipping breaks the correlation by clipping scarce correct rollouts at the same fraction as abundant correct ones.

Step0 – Step400. Qwen2.5-Math-1.5B has no rea-

soning capability, so the correlation at step 0 reflects sampling noise from RL and datasampler initialization rather than RL dynamics.

Step400 – Step600. This is a clip-free regime for scarce correct rollouts from low c groups. This results in a consistent rise in correlation across all baseline methods. Figure 3 shows clipping fires infrequently in this window, so the advantage-ratio coupling is preserved.

Step600+. The baselines only diverge once the clip fraction rises, at which point GAPO sustains the correlation while fixed clipping collapses it. Figure 4 shows GAPO also sustains higher pass@1 and pass@256 over the same window. Appendix B provides the intervention evidence that clipping itself, rather than a common cause, drives the divergence. Beyond aggregated metrics, Figure 8 shows GAPO maintains problem coverage. While baselines progressively lose solve rates on medium-difficulty AIME24 problems, GAPO retains 10.9 problems with $> 5\%$ solve rate in late training compared to 6.7-8.6 for baselines. Appendix E shows an example of reasoning paths.

4.2.1 Ruling out confounds before clipping

Since the GSPO-asym uniform clipping baseline began firing aggressively around step 600, we took the step-600 checkpoint and switched to adaptive clipping to continue RLVR. Figure 7 in Appendix B shows GAPO continue-finetuned (ckpt600) maintains higher adv-IS correlation and pass@k, ruling out factors before aggressive clipping (e.g. reduced advantage spectrum) as the cause of the decline in advantage-IS correlation.

Base model: DeepSeek-R1-Distill-Qwen-1.5B. Data: DeepScaleR Pass@1 / Pass@16 (Temperature=0.6, Top-p=1, n=16, Tmax=24576)				
In-domain	Base	GSPO	F-GSPO	GAPO
AIME24	28.5/60.0	41.3/73.3	40.2/73.3	44.0/76.7
AIME25	22.3/43.3	29.4/50.0	29.6/46.7	30.8/56.7
AMC	71.4/ 95.0	79.2/ 95.0	82.2/ 95.0	83.44/95.0
MATH500	70.4/82.4	83.0/ 94.7	85.5/93.2	85.5/93.2
Minerva	19.0/47.4	26.6/50.4	27.9/44.3	29.1/48.9
Olympiad	41.5/60.3	57.8/76.2	59.0/72.4	59.4/76.2

Table 1: **Comparison of GAPO with GSPO baselines on DeepSeek-R1-Distill-Qwen-1.5B.** Pass@1 / Pass@16 across six in-domain math benchmarks at group size $N=8$. *Base* is the pretrained model evaluated without RL training.

Base model: DeepSeek-R1-Distill-Qwen-1.5B. Data: DeepCoder Pass@1 (Temperature=0.6, Top-p=1, n=2, Tmax=24576)			
In-domain	Base	DeepCoder-1.5B (reproduce)	GAPO-DeepCoder-1.5B
LCB-v5 (8/1/24-2/1/25)	16.9	22.4	24.8
HumanEval+	58.3	68.2	71.7

Table 2: **GAPO remains effective in the coding domain.** *DeepCoder-1.5B* is our reproduction of the DeepCoder recipe (Luo et al., 2025) with Tmax=24576. *GAPO-DeepCoder-1.5B* applies adaptive clipping on top of that recipe with $\epsilon_{hi}^{max} \in [0.2, 0.28]$ (token-IS).

4.3 Benchmark results

Tables 1, 2, and 3 report pass rates on three base models and the nine benchmarks. Results are from one run per configuration, as is common in RLVR. Seeded runs on Qwen2.5-1.5B-Math appear in Table 6, 7. GAPO is compared against fixed-clip RLVR methods (GRPO, Dr.GRPO, symmetric and asymmetric GSPO) and the focal-shaping variants (F-GRPO, F-GSPO) at matched clip widths.

DeepSeek-R1-Distill-Qwen-1.5B. On a DS-R1 distilled model with longer rollouts ($T_{max} = 24576$), GAPO leads or matches pass@1 on all six in-domain benchmarks (Table 1), with the largest gains on AIME24 (44.0/76.7, +2.7/+3.4 over GSPO) and AIME25 (30.8/56.7, +1.4/+6.7 over GSPO). Table 2 shows GAPO remains effective in RLVR with code generation tasks.

Qwen2.5-1.5B-Math. Table 3 shows GAPO achieves the best average pass@1 / pass@256 (37.9/76.3) among GSPO variants, and the highest AIME24 pass@1 (17.9, vs. 15.4 for F-GSPO and 16.2 for asymmetric GSPO). At the token-IS clip range, GAPO-token-IS outperforms GRPO, F-GRPO, and Dr.GRPO on average and leads pass@1 on five of the six math benchmarks, with the largest gains on AIME24 pass@1 (14.6) and IFEval (20.9 vs. 11.4–19.4 for the baselines). These gains are

significant on at least 3 of 6 benchmarks against every baseline (Table 6, 7).

Llama-3.2-3B-Instruct. GAPO again leads on average pass@1 / pass@256 (24.1/63.4) and improves AIME24 pass@1 over the fixed-clip GSPO variants (14.6 vs. 10.7–14.4). GAPO also achieves the best IFEval (56.4), tying F-GRPO.

Across all three base models, GAPO improves pass@1 without sacrificing pass@ k . Gains concentrate on the harder benchmarks (AIME24, +1.72 over asymmetric GSPO, significance $p < 0.001$ in Table 7) and shrink on easier ones, where AMC regresses against the same baseline (Table 3, 7). This is expected, since adaptive clipping redistributes gradient toward scarce correct rollouts, which are only scarce on problems the model rarely solves.

4.4 Ablation Study

We ablate three key algorithm choices for GAPO. **Seq-IS vs token-IS.** In Table 3, the more principled sequence-level adaptive IS clipping (AIME24 pass@1 17.9%, math average pass@1 37.9%) outperforms token-level adaptive IS clipping (AIME24 pass@1 14.6%, math average 37.6%).

Group size k . k controls how finely advantages can be differentiated within a group, and hence the resolution of the per-prompt clip threshold. Increasing k from 4 to 8 further widens the performance gap between GAPO and GSPO (Figure 6). On the math training data, increasing k beyond 8 shows no noticeable gain, but the training was more stable.

Adaptive upper clipping bound ϵ_{hi}^{max} . Across the range in Table 8, GSPO with fixed $\epsilon_{hi} = \epsilon_{hi}^{max}$ is worse than GAPO on AIME24. Because adaptive clipping constrains abundant correct rollouts to save the trust-region budget, slightly relaxing the value $\epsilon_{hi}^{max} = 5e-3$, which was tuned for fixed-clipping, to $5.25e-3$ yields a small improvement.

5 Related Work

5.1 LLM Post-training

SFT. s1 (Muennighoff et al., 2025) shows that SFT on only 1K high-quality DeepSeek-R1 (Guo et al., 2025) traces can surpass o1-preview (Jaech et al., 2024). SSFT (Jia et al., 2026) uses a matching loss to distill from multiple teachers to capture diverse reasoning modes. On-policy distillation (Agarwal et al., 2024) leverages the access to teacher logits to perform token-level distribution matching between the student and the teacher.

Method	ϵ_{lo}	ϵ_{hi}	In-domain							OOD
			Avg.	AIME24	AIME25	AMC	MATH500	Minerva	Olympiad	IFEval
Qwen2.5-1.5B-Math			Pass@1 / Pass@256 (Temperature=1.0, Top-p=1, n=256, Tmax=3092)							
GRPO	0.2	0.2	36.7/74.4	13.8/61.1	9.9/58.0	53.1/96.2	75.4/95.6	31.9/61.1	36.3/74.3	12.2
F-GRPO	0.2	0.2	36.3/74.5	13.0/60.7	10.5/57.9	51.6/95.9	74.7/96.1	31.0/61.0	37.0/75.5	11.4
DrGRPO	0.2	0.28	33.8/75.1	13.1/60.0	8.85/63.3	51.8/97.5	70.5/94.6	22.3/62.1	36.0/73.2	19.4
GAPO-token-IS	0.2	0.28	37.6/75.6	14.6/63.3	10.2/60.0	54.9/97.5	75.8/96.4	32.3/62.1	37.6/74.2	20.9
GSPO	3e-3	3e-3	36.8/76.2	15.3/60.0	9.28/60.0	53.0/100	75.2/97.0	30.9/62.9	37.3/77.3	22.1
GSPO	3e-3	5e-3	37.7/74.1	16.2/60.0	9.54/56.7	55.3/95.0	75.3/95.6	31.0/63.6	37.6/73.7	20.1
F-GSPO	3e-3	5e-3	36.9/76.1	15.4/63.3	9.44/56.7	53.7/97.5	75.0/97.0	30.7/63.4	37.1/78.5	21.0
GAPO	3e-3	5e-3	37.9/76.3	17.9/63.3	10.6/56.7	53.4/97.5	76.5/97.5	31.2/64.1	38.0/78.5	22.2
Llama3.2-3B-Instruct			Pass@1 / Pass@256 (Temperature=1.0, Top-p=1, n=256, Tmax=8192)							
GRPO	0.2	0.2	23.0/59.9	10.7/40.7	0.7/21.5	30.5/88.2	55.0/90.6	21.8/59.0	19.4/59.3	54.1
F-GRPO	0.2	0.2	23.0/63.4	12.1/46.1	1.0/29.5	29.8/90.6	54.1/92.9	21.0/60.1	20.1/61.3	56.4
GSPO	3e-3	3e-3	22.7/56.5	13.6/36.7	0.4/16.7	27.8/86.5	51.5/87.8	20.3/57.6	22.3/53.9	54.9
GSPO	3e-3	5e-3	23.0/59.2	14.4/46.7	0.7/20.0	27.9/87.5	52.0/89.2	20.5/55.7	22.4/56.1	54.8
F-GSPO	3e-3	5e-3	23.8/59.2	13.2/36.7	0.9/23.3	31.1/92.5	54.0/87.8	21.1/57.6	22.5/57.5	56.3
GAPO	3e-3	5e-3	24.1/63.4	14.6/53.3	1.0/23.3	31.6/91.2	54.6/93.2	20.3/60.5	22.4/58.9	56.4

Table 3: **Comparison of GAPO with fixed-clip RLVR baselines on two base models.** Pass@1 / Pass@256 on six in-domain math benchmarks (Avg., AIME24, AIME25, AMC, MATH500, Minerva, Olympiad) and one OOD benchmark (IFEval) for Qwen2.5-1.5B-Math and Llama-3.2-3B-Instruct, trained at group size $N=8$. Columns ϵ_{lo} and ϵ_{hi} give the lower and upper clip thresholds (equal values denote symmetric clipping; differing values denote asymmetric clipping). Blue shaded rows mark our methods. To assess the statistical significance, 95% confidence intervals and pairwise significance tests over 3 training seeds are reported in Table 6, 7 respectively.

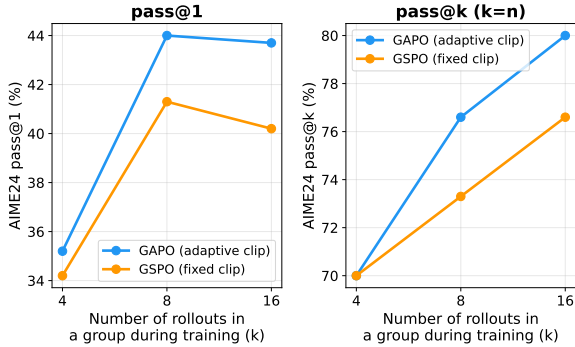


Figure 6: Ablation on the number of rollouts per group k used during training. Scaling from $k=4$ to $k=8$ yields clear improvement in both pass@1 and pass@256 on AIME24 eval, and widens the gap between fixed clipping (GSPO) and adaptive clipping (GAPO).

RLVR. Khatri et al. (2025) surveys group-relative RLVR methods. For token loss aggregation, Dr.GRPO (Liu et al., 2025) normalizes by a constant and DAPO (Yu et al., 2025) by total batch tokens, both to avoid length bias. For advantage normalization, Dr.GRPO removes the within-group std rescaling to preserve unbiasedness, while REINFORCE++ (Hu et al., 2025a) substitutes batch-level std. All use token-level IS with fixed clipping, with DAPO widening the upper threshold asym-

metrically. GSPO (Zheng et al., 2025) instead uses sequence-level IS and clipping. CISPO (Chen et al., 2025a) clips the IS weight directly while preserving every token’s gradient, a weighted REINFORCE. DAPO’s clip-higher and CISPO’s gradient preservation share our motivation of retaining exploration tokens. We arrive at the same goal through an explicit closed-form derivation: the per-prompt reverse-KL trust-region-optimal IS ratio suggests a clip threshold proportional to advantage. Other exploration-oriented approaches scale the number of rollouts (Hu et al., 2025b), introduce separate fixed clipping thresholds for correct and incorrect rollouts (Karaman et al., 2026), or only optimize policy on exploratory forking tokens (Wang et al., 2025b,a).

5.2 Advantage and Reward Shaping for Exploration in RLVR

Yue et al. (2025) found that pass@ k of RLVR-trained LLMs score below the base model at large k , motivating a line of work that directly optimizes pass@ k (Chen et al., 2025b; Walder and Karkhanis, 2025; Thrampoulidis et al., 2025). However, pass@1, the deployment metric, has been observed to degrade under this training (Walder and Karkhanis, 2025; Barakat et al., 2026). Rather than directly optimizing pass@ k , others propose

dedicated advantage-shaping schemes to encourage exploration. F-GRPO (Plyusov et al., 2026) and He et al. (2025) apply advantage shaping based on within-group success count; Zhou et al. (2025) reweights the loss across groups by success rate; and Gai et al. (2025) adds a differential-smoothing reward based on per-rollout correctness. In contrast, GAPO still optimizes the pass@1 objective, encouraging exploration by maintaining high correlation between the empirical IS ratio and advantage, which preserves gradients on rare correct rollouts relative to redundant ones, according to the target IS ratio for the reverse KL trust region formulation.

6 Conclusion

We provide a simple plug-in adaptive clipping method, Group Adaptive Policy Optimization (GAPO), motivated by the reverse-KL trust-region-optimal IS ratio at a single prompt. Under this trust region, the optimal IS ratio scales exponentially with advantage, suggesting that the per-rollout clip threshold should scale with advantage rather than be applied uniformly. In RLVR, binary rewards reduce this to a closed-form schedule indexed by the group correctness count c , requiring no new hyperparameters beyond the existing clip range. Across Qwen2.5-Math-1.5B, Llama-3.2-3B-Instruct, and DeepSeek-R1-Distill-Qwen-1.5B on mathematical reasoning and competitive coding benchmarks, GAPO consistently improves pass@1 while retaining pass@ k , and maintains a high IS–advantage correlation throughout training where fixed clipping collapses it. Unlike reward-shaping and advantage-shaping approaches for exploration, GAPO still optimizes the deployment metric pass@1 directly while implicitly encouraging exploration through optimally allocating trust-region constraint.

Limitations

Fixed ϵ_{lo} across rollouts. We adapt only the upper clip threshold ϵ_{hi} based on A_i , which governs positive-advantage rollouts. The lower clip ϵ_{lo} for negative-advantage (incorrect) rollouts is held fixed throughout. Because ϵ_{lo} acts only on incorrect rollouts, fixing it leaves the positive learning signal GAPO targets untouched. In future work, we plan to study how jointly adapting the headroom for penalizing incorrect rollouts based on task difficulty indirectly affects the model’s exploration behavior.

Single prompt trust-region analysis. Our derivation of the per-rollout optimal IS ratio ρ_i^* is based on the reverse-KL trust region at a single prompt x (Section 3.1). Since equation 5 constrains the KL averaged over the batch, the multiplier λ is in principle shared across prompts; we treat it as a global constant, which is what makes equation 11 depend on c alone. RLVR runs always have > 1 prompts or tasks in a mini-batch, so the gap is real, though the observed IS-advantage correlations (Figure 5) suggest the per-prompt heuristic still provides useful empirical guidance.

Ethical Considerations

Scope and broader impact. GAPO is a methodological contribution to RLVR: it adapts the per-rollout clipping threshold without introducing new model capabilities, training data, or evaluation protocols beyond those standard in the field. The intended impact is to improve the training efficiency and final performance of LLMs on mathematical reasoning with applications in education.

Data and models. All training data (DeepScaleR, DeepCoder-preview-dataset) and evaluation benchmarks (AIME24/25, MATH500, AMC, Olympiad, Minerva, LCB-v5, HumanEval+, IFEval) are publicly available. The base models (Qwen2.5-Math-1.5B, Llama-3.2-3B-Instruct, DeepSeek-R1-Distill-Qwen-1.5B) are openly released. No human subjects, annotators, or proprietary data were involved.

Dual-use considerations. As with any improvement to LLM training, GAPO may indirectly contribute to producing more capable reasoning systems, which carry the general dual-use considerations common to language modeling research. GAPO does not specifically lower the cost of generating harmful content, nor does it target capabilities relevant to harmful applications; we do not anticipate risks beyond those inherent to RLVR training in general.

Acknowledgements

We used Claude Opus 4.7 (Anthropic, 2026) for occasional grammar checks on individual sentences, which were entered manually through the chat interface. No text, code, experimental results, or figures were LLM-generated. All hypotheses, technical content, and conclusions are our own, and we verified them ourselves.

References

- Rishabh Agarwal, Nino Vieillard, Yongchao Zhou, Piotr Stanczyk, Sabela Ramos Garea, Matthieu Geist, and Olivier Bachem. 2024. [On-policy distillation of language models: Learning from self-generated mistakes](#). In *The Twelfth International Conference on Learning Representations*.
- Anthropic. 2026. Claude Opus 4.7. <https://claude.ai>. Generative AI Chatbot.
- Anas Barakat, Souradip Chakraborty, Khushbu Pahwa, and Amrit Singh Bedi. 2026. Why pass@k optimization can degrade pass@1: Prompt interference in llm post-training. *arXiv preprint arXiv:2602.21189*.
- Aili Chen, Aonian Li, Bangwei Gong, Binyang Jiang, Bo Fei, Bo Yang, Boji Shan, Changqing Yu, Chao Wang, Cheng Zhu, and 1 others. 2025a. Minimax-m1: Scaling test-time compute efficiently with lighting attention. *arXiv preprint arXiv:2506.13585*.
- Gang Chen, Yiming Peng, and Mengjie Zhang. 2018. An adaptive clipping approach for proximal policy optimization. *arXiv preprint arXiv:1804.06461*.
- Zhipeng Chen, Xiaobo Qin, Youbin Wu, Yue Ling, Qinghao Ye, Wayne Xin Zhao, and Guang Shi. 2025b. Pass@k training for adaptively balancing exploration and exploitation of large reasoning models. *arXiv preprint arXiv:2508.10751*.
- Jingchu Gai, Guanning Zeng, Huaqing Zhang, and Aditi Raghunathan. 2025. Differential smoothing mitigates sharpening and improves llm reasoning. *arXiv preprint arXiv:2511.19942*.
- Izrail Moiseevitch Gelfand and Sergei Vasilyevich Fomin. 2000. *Calculus of variations*. Dover Publications, Mineola, NY.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, and 1 others. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shitong Ma, Peiyi Wang, Xiao Bi, and 1 others. 2025. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*.
- Andre Wang He, Daniel Fried, and Sean Welleck. 2025. [Rewarding the unlikely: Lifting GRPO beyond distribution sharpening](#). In *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*, pages 25548–25560, Suzhou, China. Association for Computational Linguistics.
- Chaoqun He, Renjie Luo, Yuzhuo Bai, Shengding Hu, Zhen Thai, Junhao Shen, Jinyi Hu, Xu Han, Yujie Huang, Yuxiang Zhang, and 1 others. 2024. Olympiadbench: A challenging benchmark for promoting agi with olympiad-level bilingual multimodal scientific problems. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 3828–3850.
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. 2021. Measuring mathematical problem solving with the math dataset. *arXiv preprint arXiv:2103.03874*.
- Jian Hu, Jason Klein Liu, Haotian Xu, and Wei Shen. 2025a. Reinforce++: Stabilizing critic-free policy optimization with global advantage normalization. *arXiv preprint arXiv:2501.03262*.
- Jian Hu, Mingjie Liu, Ximing Lu, Fang Wu, Zaid Harchaoui, Shizhe Diao, Yejin Choi, Pavlo Molchanov, Jun Yang, Jan Kautz, and 1 others. 2025b. Brorl: Scaling reinforcement learning via broadened exploration. *arXiv preprint arXiv:2510.01180*.
- Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, and 1 others. 2024. Openai o1 system card. *arXiv preprint arXiv:2412.16720*.
- Naman Jain, Alex Gu, Wen-Ding Li, Fanjia Yan, Tianjun Zhang, Sida Wang, Armando Solar-Lezama, Koushik Sen, and Ion Stoica. 2025. Livecodebench: Holistic and contamination free evaluation of large language models for code. In *International Conference on Learning Representations*, volume 2025, pages 58791–58831.
- Sheng Jia, Xiao Wang, and Shiva Kasiviswanathan. 2026. Training large language models to reason in parallel with global forking tokens. In *International Conference on Learning Representations*, volume 2026, pages 84820–84846.
- Batuhan K Karaman, Aditya Rawal, Suhaila Shakiah, Mohammad Ghavamzadeh, Mingyi Hong, Arijit Biswas, and Ruida Zhou. 2026. Dispo: Enhancing training efficiency and stability in reinforcement learning for large language model mathematical reasoning. *arXiv preprint arXiv:2602.00983*.
- Devvrit Khatri, Lovish Madaan, Rishabh Tiwari, Rachit Bansal, Sai Surya Duvvuri, Manzil Zaheer, Inderjit S Dhillon, David Brandfonbrener, and Rishabh Agarwal. 2025. The art of scaling reinforcement learning compute for llms. *arXiv preprint arXiv:2510.13786*.
- Aitor Lewkowycz, Anders Andreassen, David Dohan, Ethan Dyer, Henryk Michalewski, Vinay Ramasesh, Ambrose Slone, Cem Anil, Imanol Schlag, Theo Gutman-Solo, and 1 others. 2022. Solving quantitative reasoning problems with language models. *Advances in neural information processing systems*, 35:3843–3857.
- Kaixin Li. 2024. [Verified taco problems](https://huggingface.co/datasets/likeixin/TACO-verified). <https://huggingface.co/datasets/likeixin/TACO-verified>.

- Tianjian Li, Yiming Zhang, Ping Yu, Swarnadeep Saha, Daniel Khashabi, Jason Weston, Jack Lanchantin, and Tianlu Wang. 2025. Jointly reinforcing diversity and quality in language model generations. *arXiv preprint arXiv:2509.02534*.
- Jiawei Liu, Chunqiu Steven Xia, Yuyao Wang, and Lingming Zhang. 2023. Is your code generated by chatgpt really correct? rigorous evaluation of large language models for code generation. *Advances in neural information processing systems*, 36:21558–21572.
- Zichen Liu, Changyu Chen, Wenjun Li, Penghui Qi, Tianyu Pang, Chao Du, Wee Sun Lee, and Min Lin. 2025. Understanding rl-zero-like training: A critical perspective. *arXiv preprint arXiv:2503.20783*.
- Ilya Loshchilov and Frank Hutter. 2017. Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*.
- Michael Luo, Sijun Tan, Roy Huang, Ameen Patel, Alpav Ariyak, Qingyang Wu, Xiaoxiang Shi, Rachel Xin, Colin Cai, Maurice Weber, and 1 others. 2025. Deepcoder: A fully open-source 14b coder at o3-mini level. *Notion Blog*, 1. <https://www.together.ai/blog/deepcoder>.
- Justus Mattern, Sami Jaghouar, Manveer Basra, Jan-nik Straube, Matthew Di Ferrante, Felix Gabriel, Jack Min Ong, Vincent Weisser, and Johannes Hagemann. 2025. *Synthetic-1: Two million collaboratively generated reasoning traces from deepseek-rl*. <https://www.primeintellect.ai/blog/synthetic-1-release>.
- Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke Zettlemoyer, Percy Liang, Emmanuel Candès, and Tatsunori Hashimoto. 2025. s1: Simple test-time scaling. *arXiv preprint arXiv:2501.19393*.
- Daniil Plyusov, Alexey Gorbatoyski, Boris Shaposhnikov, Viacheslav Sini, Alexey Malakhov, and Daniil Gavrillov. 2026. F-grpo: Don’t let your policy learn the obvious and forget the rare. *arXiv preprint arXiv:2602.06717*.
- John Schulman, Sergey Levine, Pieter Abbeel, Michael Jordan, and Philipp Moritz. 2015. Trust region policy optimization. In *International conference on machine learning*, pages 1889–1897. PMLR.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. 2017. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*.
- Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. 2024. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *Preprint*, arXiv:2402.03300.
- Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. 2025. Hybridflow: A flexible and efficient rlhf framework. In *Proceedings of the Twentieth European Conference on Computer Systems*, pages 1279–1297.
- Sijun Tan, Michael Luo, Justin Wong, Colin Cai, Xiaoxiang Shi, William Yuan Tang, Manan Roongta, Tianjun Zhang, Li Erran Li, Raluca Ada Popa, and Ion Stoica. 2026. Deepscaler: Effective RL scaling of reasoning models via iterative context lengthening. <https://openreview.net/forum?id=I6GzDCne7U>.
- Yunhao Tang, Kunhao Zheng, Gabriel Synnaeve, and Rémi Munos. 2025. Optimizing language models for inference time objectives using reinforcement learning. *arXiv preprint arXiv:2503.19595*.
- Christos Thrampoulidis, Sadegh Mahdavi, and Wenlong Deng. 2025. Advantage shaping as surrogate reward maximization: Unifying pass@k policy gradients. *arXiv preprint arXiv:2510.23049*.
- Christian Walder and Deep Tejas Karkhanis. 2025. Pass@k policy optimization: Solving harder reinforcement learning problems. In *Advances in Neural Information Processing Systems*, volume 38, Main Conference, pages 152416–152445. Curran Associates, Inc.
- Haozhe Wang, Qixin Xu, Che Liu, Junhong Wu, Fangzhen Lin, and Wenhui Chen. 2025a. Emergent hierarchical reasoning in llms through reinforcement learning. *arXiv preprint arXiv:2509.03646*.
- Shenzhi Wang, Le Yu, Chang Gao, Chujie Zheng, Shixuan Liu, Rui Lu, Kai Dang, Xionghui Chen, Jianxin Yang, Zhenru Zhang, and 1 others. 2025b. Beyond the 80/20 rule: High-entropy minority tokens drive effective reinforcement learning for llm reasoning. *arXiv preprint arXiv:2506.01939*.
- An Yang, Beichen Zhang, Binyuan Hui, Bofei Gao, Bowen Yu, Chengpeng Li, Dayiheng Liu, Jianhong Tu, Jingren Zhou, Junyang Lin, and 1 others. 2024. Qwen2. 5-math technical report: Toward mathematical expert model via self-improvement. *arXiv preprint arXiv:2409.12122*.
- Yixin Ye, Yang Xiao, Tiantian Mi, and Pengfei Liu. 2025. Aime-preview: A rigorous and immediate evaluation framework for advanced mathematical reasoning. <https://github.com/GAIR-NLP/AIME-Preview>. GitHub repository.
- Qiyang Yu, Zheng Zhang, Ruofei Zhu, Yufeng Yuan, Xiaochen Zuo, Yu Yue, Weinan Dai, Tiantian Fan, Gao-hong Liu, juncai liu, LingJun Liu, Xin Liu, Haibin Lin, Zhiqi Lin, Bole Ma, Guangming Sheng, Yuxuan Tong, Chi Zhang, Mofan Zhang, and 17 others. 2025. Dapo: An open-source llm reinforcement learning system at scale. In *Advances in Neural Information Processing Systems*, volume 38, Main Conference, pages 113222–113244. Curran Associates, Inc.

Yang Yue, Zhiqi Chen, Rui Lu, Andrew Zhao, Zhaokai Wang, Shiji Song, and Gao Huang. 2025. Does reinforcement learning really incentivize reasoning capacity in llms beyond the base model? *arXiv preprint arXiv:2504.13837*.

Chujie Zheng, Shixuan Liu, Mingze Li, Xiong-Hui Chen, Bowen Yu, Chang Gao, Kai Dang, Yuqiong Liu, Rui Men, An Yang, and 1 others. 2025. Group sequence policy optimization. *arXiv preprint arXiv:2507.18071*.

Jeffrey Zhou, Tianjian Lu, Swaroop Mishra, Siddhartha Brahma, Sujoy Basu, Yi Luan, Denny Zhou, and Le Hou. 2023. Instruction-following evaluation for large language models. *arXiv preprint arXiv:2311.07911*.

Jingyu Zhou, Lu Ma, Hao Liang, Chengyu Shen, Bin Cui, and Wentao Zhang. 2025. Daro: Difficulty-aware reweighting policy optimization. *arXiv preprint arXiv:2510.09001*.

A Training Hyperparameters

B Intervention study

To isolate uniform clipping as the cause of the IS-advantage correlation decline, we branch from the GSPO-asym step-600 ckpt, where uniform clipping starts firing aggressively, and switch to adaptive clipping. The GAPO branch sustains advantage-IS correlation and Pass@256 while the uniform branch declines, ruling out earlier factors (e.g. reduced advantage spectrum). (Figure 7)

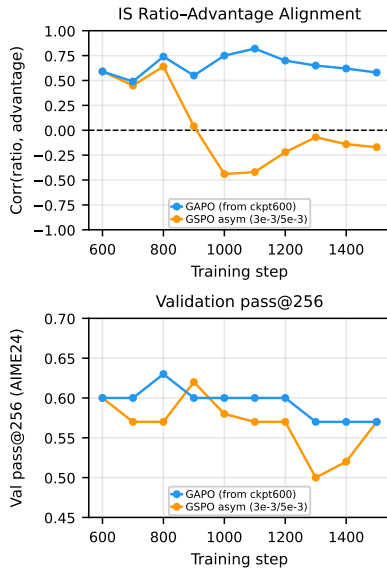


Figure 7: **Clipping, not an earlier factor, drives the correlation decline.** Both branches start from GSPO-asym step 600. Adaptive clipping (GAPO) sustains adv-IS correlation and Pass@256, uniform clipping (GSPO-asym) declines. (Training data: DeepScaleR)

Parameter	Value
<i>Optimization</i>	
Trainer Precision	bfloat16
Optimizer	AdamW (Loshchilov and Hutter, 2017)
Learning rate	1×10^{-6}
LR schedule	Constant
Weight decay	0.01
Gradient clip	1.0
KL coefficient λ	0
Entropy coefficient	0
Reward	Verifiable (0/1)
Training steps	10 epochs
<i>RL Settings</i>	
Inference Precision	bfloat16
Training dataset	DeepScaleR (Tan et al., 2026), 39,202 samples
Rollout batch (prompts)	256
Mini-batch (prompts)	64
Rollouts per prompt k	8 or 16
Token-IS clip ($\epsilon_{lo}, \epsilon_{hi}$)	(0.2, 0.28)
Seq-IS clip ($\epsilon_{lo}, \epsilon_{hi}$)	($3e-3, 5e-3$) ($7e-5, 3e-4$) distilled base
<i>Generation</i>	
Top- p	1.0
Max prompt length	2048
----- Qwen2.5-1.5B-Math/Llama-3.2-3B-Instruct/DS-R1-Distill-Qwen-1.5B -----	
Max response length	3092 / 8192 / 24576
Temperature	1.0 / 1.0 / 1.0

Table 4: **Training hyperparameters used in all math RL experiments.** We use 8xH200 GPUs with the verl (Sheng et al., 2025) framework for both training and evaluation. Training takes about 4–5 days for DS-R1 distilled models and 1 day for Qwen2.5-Math-1.5B and Llama models. These are used for Table 1, 3, Figure 4. For all of our experiments, we use DP=8, TP=1, SP=1 with FSDP.

Parameter	Value
<i>Optimization hyperparameters are the same as Table 4</i>	
<i>RL Settings</i>	
Inference Precision	bfloat16
Training dataset	DeepCoder-Preview-Dataset (Luo et al., 2025), 24,269 samples consisting of 16,238 SYNTHETIC-1 (Mattern et al., 2025), 7,432 from TACO (Li, 2024), and 599 from LiveCodeBench (2023/5/1-2024/7/31) (Jain et al., 2025).
Rollout batch (prompts)	256
Mini-batch (prompts)	64
Rollouts per prompt k	8 or 16
Token-IS clip ($\epsilon_{lo}, \epsilon_{hi}$)	(0.2, 0.28)
<i>Generation parameters are the same as Table 4 under DS-R1-Distill-Qwen-1.5B</i>	

Table 5: **Hyperparams in coding RL runs** for Table 2.

C Robustness Analysis

We stress-test the main-table results (Table 3) along two axes: statistical significance across training seeds, and sensitivity to the clip-width hyperparameter. Table 3 follows the common practice in RLVR. Seeded runs cover the Qwen2.5-1.5B-Math setting, since compute constraints made 3-seed runs infeasible for all configurations. So Table 6 means differ slightly from the single-seed entries in Table 3.

C.1 Significance Tests

Table 6 reports 3-seed means with 95% confidence intervals, and Table 7 the paired per-seed differences with Welch’s t -test significance. GAPO’s Pass@1 gains are significant on at least 3 of 6 benchmarks against every baseline, and strongest on harder reasoning benchmarks like AIME24.

Method	ϵ_{lo}	ϵ_{hi}	Qwen2.5-1.5B-Math Pass@1 (Temperature=1.0, Top-p=1, n=256, Tmax=3092)					
			AIME24	AIME25	AMC	MATH500	Minerva	Olympiad
GRPO	0.2	0.2	13.82±0.81	10.07±0.79	53.16±1.26	75.61±0.99	31.33±0.94	36.90±1.31
F-GRPO	0.2	0.2	12.95±0.75	10.44±0.33	51.53±1.69	74.64±0.64	30.99±1.18	36.95±1.52
DrGRPO	0.2	0.28	13.11±0.69	8.87±0.77	51.65±1.33	70.48±1.01	22.43±1.25	35.92±1.39
GAPO-token-IS	0.2	0.28	14.61±0.83	10.19±0.48	54.87±0.65	75.69±1.03	32.34±0.97	37.53±1.08
GSPO	3e−3	3e−3	15.28±0.74	9.28±0.68	53.12±0.92	75.20±0.60	31.05±1.10	37.28±0.65
GSPO	3e−3	5e−3	16.24±0.61	9.56±0.70	55.23±1.19	75.32±0.52	31.07±1.42	37.55±1.11
F-GSPO	3e−3	5e−3	15.39±0.51	9.34±0.77	53.73±1.17	75.08±1.05	30.63±1.35	37.16±1.41
GAPO	3e−3	5e−3	17.96±0.48	10.56±0.51	53.37±0.83	76.59±1.34	31.70±1.60	38.03±0.58

Table 6: **95% confidence intervals for comparing GAPO with fixed-clip RLVR baselines on Qwen2.5-1.5B-Math.** Each entry is the mean over 3 training seeds $\pm$ a 95% confidence interval (Student- t , $df=2$). Recap on the settings: ϵ_{lo} and ϵ_{hi} are the lower and upper clipping thresholds: equal values denote symmetric clipping, differing values asymmetric clipping. The upper block uses token-level importance sampling and the lower block sequence-level (GSPO-style) importance sampling. Within each block, our method (blue) shares the importance-sampling scheme and the clip range of its baselines, so the only difference is that ϵ_{hi} is adapted per group rather than held fixed.

Comparison (Δ = ours − baseline)	AIME24	AIME25	AMC	MATH500	Minerva	Olympiad
GAPO-token-IS (clip 0.2/0.28) vs. baselines:						
vs. GRPO	+0.79*	+0.12	+1.71*	+0.08	+1.01*	+0.63
vs. F-GRPO	+1.66**	−0.25	+3.34**	+1.05*	+1.35*	+0.58
vs. DrGRPO	+1.50**	+1.32**	+3.22**	+5.21***	+9.91***	+1.61*
GAPO (clip 3e−3/5e−3) vs. baselines:						
vs. GSPO 3/3	+2.68***	+1.28**	+0.25	+1.39*	+0.65	+0.75*
vs. GSPO 3/5	+1.72***	+1.00**	−1.86**	+1.27*	+0.63	+0.48
vs. F-GSPO	+2.57***	+1.22*	−0.36	+1.51*	+1.07	+0.87

Table 7: **Pairwise significance of Pass@1 gains.** Each entry is the mean difference Δ on Pass@1 over 3 training seeds, computed from the per-seed runs summarized in Table 6. Significance is from a Welch’s t -test: * denotes $p < 0.05$, ** denotes $p < 0.01$, and *** denotes $p < 0.001$; unmarked entries are not significant. GAPO-token-IS is compared against the token-IS baselines, GAPO against the GSPO baselines. The gains are statistically significant on at least 3 out of 6 benchmarks for every baseline comparison, and the improvement is noticeable for harder benchmarks like AIME24. There’s only one meaningful regression compared to GSPO 3/5 on AMC.

C.2 Hyperparameter Sweeps

Table 8 sweeps the maximum upper clip $\epsilon_{hi}^{\max}$: GAPO’s adaptive ϵ_{hi} beats uniform clipping at every level tested.

	AIME24 Pass@1	Qwen2.5-1.5B-Math (n=256, Tmax=3092)			
GSPO asym (uniform)	Clip-high (ϵ_{hi}) adaptive range Pass@1	[3e−3] 15.28	[5e−3] 16.24	[5.25e−3] 15.98	[7e−3] 15.91
GAPO (ours)	Clip-high (ϵ_{hi}) adaptive range Pass@1	[3e−3, 5e−3] 17.96	[3e−3, 5e−3] 17.96	[3e−3, 5.25e−3] 18.20	[3e−3, 7e−3] 17.38
Improvement	ours/baseline − 1	+17.54%	+10.59%	+13.89%	+9.24%

Table 8: **GAPO’s adaptive ϵ_{hi} beats uniform clipping at every level tested.** AIME24 Pass@1 on Qwen2.5-1.5B-Math. Each column pairs a uniform fixed clipping GSPO baseline with the GAPO run it is compared against; GAPO’s adaptive range is listed in its own row. ϵ_{lo} is fixed at 3e−3.

D Per-Problem Solve Rate Analysis

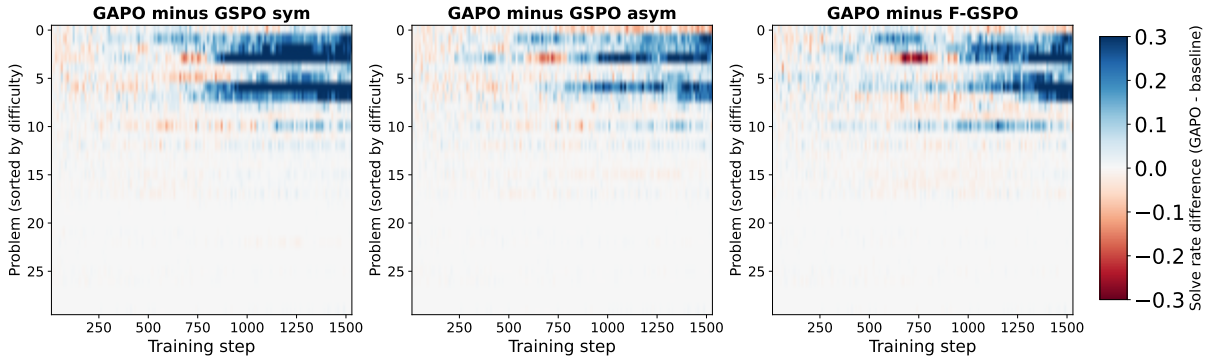


Figure 8: **Per-problem solve rate difference (GAPO minus baseline) on AIME24 validation over training for Qwen2.5-Math-1.5B** (256 samples per problem, evaluated every 10 steps). Problems are sorted by difficulty (easiest at top). Blue indicates higher solve rate for GAPO; red indicates higher solve rate for the baseline. Results are averaged across all seeds per method. GAPO’s advantage is concentrated on medium-difficulty problems (rows 5–12) and grows in late training as baselines progressively lose solve rate on these frontier problems. Averaged over steps 1200–1530, GAPO maintains 10.9 problems with $> 5\%$ solve rate compared to 8.6 (F-GSPO), 8.5 (GSPO asym), and 6.7 (GSPO sym).

E Diverse Correct Solutions Found by GAPO

Problem (AIME24, Chip Placement)

There is a collection of 25 indistinguishable white chips and 25 indistinguishable black chips. Find the number of ways to place some of these chips in the 25 unit cells of a 5×5 grid such that:

- (i) each cell contains at most one chip;
- (ii) all chips in the same row and all chips in the same column have the same colour;
- (iii) any additional chip placed on the grid would violate one or more of the previous two conditions.

<Generation 1> — Color the row/columns, add chips by deduction

```
<|im_start|>assistant
<|im_start|>
“each cell at most one chip, all chips in same row same color, same column same color” ... “rows and columns alternate in color” ... “mixing: choose which rows are white and which are black; similarly for columns” ...  $2^5 - 2 = 30$  ways to choose rows (excluding all-white and all-black) ...  $2^5 - 2 = 30$  ways for columns ... mixed configurations:  $30 \times 30 = 900$  ... monochrome: 2 ... (Comment: each row/column picks white or black independently; cells where row and column colors agree get a chip, the rest stay empty. Excluding the trivial all-white and all-black choices gives  $30 \times 30 = 900$  mixed placements; add the 2 monochrome grids: 902.)
<|im_start|>answer
Answer: ... 902
```

<Generation 2> — Partition the grid into color blocks

```
<|im_start|>assistant
<|im_start|>
“the grid is divided into two rectangles of the same color” ... “all 25 cells same color: 2 configurations” ... “choose a subset of rows and a subset of columns to form the two regions” ...  $2^5 - 2$  ways to divide the rows (any non-empty proper subset) ...  $2^5 - 2$  ways to divide the columns ...  $(2^5 - 2) \times (2^5 - 2) = 30 \times 30 = 900$  partitions ... total  $2 + 900 = 902$  ... (Comment: each valid chip placement corresponds to a partition of the grid into sub-grids; counting non-trivial row/column splits gives  $30 \times 30$ , and the two trivial monochrome cases add 2.)
<|im_start|>answer
Answer: ... 902
```

<Generation 3> — Alternating patterns

```
<|im_start|>assistant
<|im_start|>
“no two chips of different colors are adjacent” ... “each row and each column must be uniform in color” ... “the grid is bipartite in terms of color” ... explicit patterns enumerated: rows {1, 3, 5} white, rows {2, 4} black; or vice versa; ... “choose any subset of the 5 rows to be all white, excluding all-W and all-B:  $2^5 - 2 = 30$ ” ... “similarly 30 choices for columns” ...  $30 \times 30 = 900$  ... plus 2 monochrome ... (Comment: start from a specific pattern, e.g., rows {1, 3, 5} white and rows {2, 4} black, recognise that any way of splitting the 5 rows into white and black works as long as both colors appear, and likewise for the columns; this generalizes to  $30 \times 30$  mixed patterns, plus the 2 single-color grids.)
<|im_start|>answer
Answer: ... 902
```

Figure 9: **Diversity preservation on a single AIME24 problem.** The baseline GSPO produces only 3 correct generations on this problem, while GAPO produces 28 and, near peak validation pass@1, samples three distinct correct reasoning paths. All three converge to the same answer, but enter the problem from different conceptual angles. This suggests that GAPO retains multiple correct modes where the baseline has effectively collapsed.