
On the Design of KL-Regularized Policy Gradient Algorithms for LLM Reasoning

Yifan Zhang^{*1,2} Yifeng Liu^{*1} Huizhuo Yuan¹ Yang Yuan^{3,4}
Quanquan Gu^{1†} Andrew Chi-Chih Yao^{3,4†}

¹University of California, Los Angeles ²Princeton University
³IIS, Tsinghua University ⁴Shanghai Qi Zhi Institute

Abstract

Policy gradient algorithms have been successfully applied to enhance the reasoning capabilities of large language models (LLMs). KL regularization is ubiquitous, yet the design surface, choice of KL direction (forward vs. reverse), normalization (normalized vs. unnormalized), and estimator ($k_1/k_2/k_3$), is scattered across the literature and often intertwined with off-policy estimation. We ask a focused question: under the off-policy setting, what weighting is required for each KL variant so that the surrogate we optimize yields the exact gradient of the intended KL-regularized objective? We answer this with a compact, unified derivation we call the Regularized Policy Gradient (RPG) view. RPG (i) unifies normalized and unnormalized KL variants and shows that the widely-used k_3 penalty is exactly the unnormalized KL; (ii) specifies conditions under which REINFORCE-style losses with stop-gradient are gradient-equivalent to fully differentiable surrogates; (iii) identifies and corrects an off-policy importance-weighting mismatch in GRPO’s KL term; and (iv) introduces RPG-Style Clip, a clipped-importance-sampling step within RPG-REINFORCE that enables stable, off-policy policy-gradient training at scale. On mathematical reasoning benchmarks (AIME24, AIME25), RPG-REINFORCE with RPG-Style Clip improves accuracy by up to +6 absolute percentage points over DAPO. Notably, RPG is a stable and scalable RL algorithm for LLM reasoning, realized via (a) a KL-correct objective, (b) clipped importance sampling, and (c) an iterative reference-policy update scheme.

Project Page: <https://github.com/complex-reasoning/RPG>

1 Introduction

Reinforcement learning (RL), particularly policy gradient (PG) methods, provides a powerful framework for solving sequential decision-making problems in complex environments. These methods have been successfully applied in diverse domains, ranging from robotics to game playing, and have recently become instrumental in fine-tuning large language models (LLMs) to align with human preferences and instructions [Ouyang et al., 2022] and enhancing the reasoning capabilities of LLMs [Shao et al., 2024, Guo et al., 2025]. Classical PG algorithms like REINFORCE [Williams, 1992] optimize policies directly but often suffer from high gradient variance. Advanced methods like Proximal Policy Optimization (PPO) [Schulman et al., 2017] improve stability and sample efficiency, enabling large-scale applications, often by operating in an off-policy manner and employing techniques like training critic models for the estimation of value functions. Our theme in this paper is stability and scalability: which design choices in KL-regularized PG matter for robustness under off-policy sampling, and practical throughput on modern LLM stacks?

A crucial technique for stabilizing policy optimization, especially when deviating from strictly on-policy updates or aiming to control policy complexity, is regularization. Kullback-Leibler (KL)

^{*}Equal contribution; [†]Corresponding authors.

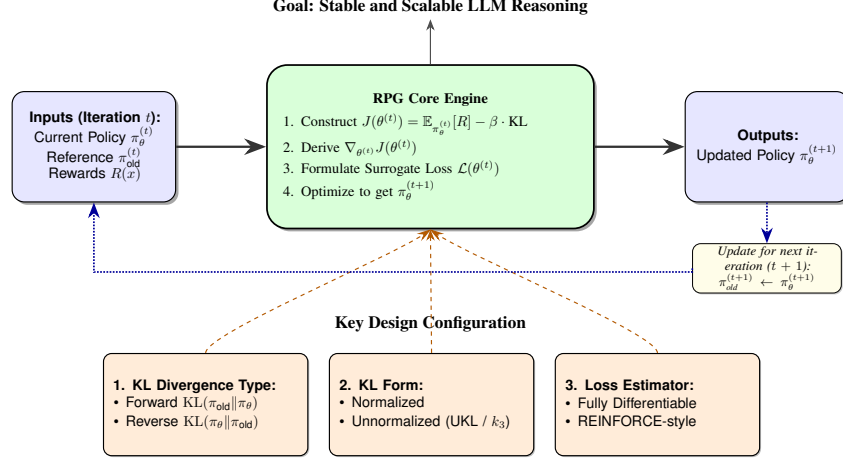


Figure 1: Overview of the iterative Regularized Policy Gradient (RPG) framework proposed in this work. At each iteration t , the central RPG Core Engine processes inputs: the current policy $\pi_{\theta}^{(t)}$, a reference policy $\pi_{\text{old}}^{(t)}$, and associated rewards $R(x)$. The engine’s operation encompasses four main steps: (1) constructing the KL-regularized objective $J(\theta^{(t)})$, which combines the expected reward with a KL divergence term; (2) deriving the off-policy policy gradient $\nabla_{\theta^{(t)}} J(\theta^{(t)})$; (3) formulating a corresponding surrogate loss function $\mathcal{L}(\theta^{(t)})$; and (4) optimizing the policy parameters to yield an updated policy $\pi_{\theta}^{(t+1)}$, aimed at enhancing LLM reasoning capabilities. The specific behavior of the RPG Core Engine is configured by three key design choices: (i) the KL Divergence Type (Forward $\text{KL}(\pi_{\text{old}} \parallel \pi_{\theta})$ or Reverse $\text{KL}(\pi_{\theta} \parallel \pi_{\text{old}})$); (ii) the KL Form (Normalized or Un-normalized, e.g., using UKL / k_3 estimators); and (iii) the Loss Estimator type (Fully Differentiable or REINFORCE-style with Stop-Gradient). The framework operates iteratively, with the updated policy $\pi_{\theta}^{(t+1)}$ from one iteration informing the inputs for the next, including the update of the reference policy $\pi_{\text{old}}^{(t+1)}$, to facilitate continuous learning and performance improvement.

divergence is a commonly used regularizer, penalizing the deviation of the learned policy π_{θ} from a reference policy π_{ref} (e.g., policy from previous iteration $\pi_{\theta_{\text{old}}}$ or a fixed prior policy π^{SFT}). KL regularization helps prevent destructive policy updates, encourages exploration around known good policies, and can prevent catastrophic forgetting or overly confident outputs [Ouyang et al., 2022].

Despite the widespread use of KL regularization in methods such as PPO (often implicitly through reward penalties) and explicit formulations like GRPO [Shao et al., 2024], there exists a considerable variety in how the KL divergence is formulated and estimated. Different choices include Forward KL and Reverse KL, handling potentially unnormalized distributions [Minka et al., 2005] (leading to unnormalized KL (UKL) and unnormalized reverse KL (URKL) formulations), and the use of various estimators like the k_2 and k_3 estimators [Schulman, 2020] designed to potentially reduce variance or offer different properties compared to the standard log-ratio (k_1) estimator. Furthermore, the interplay between the choice of KL formulation, the policy optimization setting (on-policy vs. off-policy), and the derivation of appropriate surrogate loss functions (fully differentiable vs. REINFORCE-style gradient estimators) can lead to subtle differences.

This paper provides systematic derivations and a unifying treatment of KL-regularized policy gradient methods, and revisits classical REINFORCE through the lens of clipped importance sampling. Our main contributions are summarized as follows:

- We derive policy gradients and corresponding surrogate losses for Forward/Reverse KL, in normalized (KL) and unnormalized (UKL) forms, under off-policy sampling with importance weights.
- We give both fully differentiable surrogates and REINFORCE-style losses (with stop-gradient) and prove their gradient-equivalence to the intended regularized objective (Proposition E.1, Appendix N).
- We introduce *RPG-Style Clip*, a clipped-importance REINFORCE estimator (PPO-Clip-like) that substantially improves stability and variance control while preserving the RPG gradients.

- We reveal the equality between the k_3 estimator and unnormalized KL (Appendix F), and show that GRPO’s KL penalty omits an essential importance weight under off-policy sampling. We provide a corrected estimator and loss consistent with the intended objective.
- We present an iterative training framework that periodically updates the reference model to satisfy KL constraints while allowing the policy to depart meaningfully from the initial checkpoint.
- On math reasoning, RPG-REINFORCE (with RPG-Style Clip) yields stable and scalable training and outperforms DAPO by up to +6 absolute points on AIME24/25.

2 Regularized Policy Gradients

We now derive policy gradient estimators for objectives regularized by KL divergence, assuming an online and off-policy setting where expectations are estimated using samples drawn from an old policy π_{old} via importance sampling. In the main text, we focus on the unnormalized objectives (UFKL/URKL) and summarize the corresponding surrogate losses in Table 1. The normalized formulations (FKL/RKL) and their losses are moved to Appendix G (see also Table 4). All proofs are deferred to Appendix M.

2.1 Unnormalized Forward KL Regularization

In scenarios where distributions might not be normalized (i.e., $\int_x \pi(x)dx \neq 1$), the standard KL divergence might not fully capture the dissimilarity. The unnormalized forward KL divergence addresses this by adding a mass correction term. Let $\pi_{\text{old}}(x)$ be a potentially unnormalized reference measure with total mass $Z_{\text{old}} = \int_x \pi_{\text{old}}(x)dx$. Let $\tilde{\pi}_{\text{old}}(x) = \pi_{\text{old}}(x)/Z_{\text{old}}$ be the corresponding normalized probability distribution, such that $\int \tilde{\pi}_{\text{old}}(x)dx = 1$.

Table 1: Summary of fully differentiable surrogate loss functions $\mathcal{L}(\theta)$ for unnormalized KL-regularized objectives (main text). Minimizing $\mathcal{L}(\theta)$ corresponds to maximizing $J(\theta) = \mathbb{E}_{\pi_\theta}[R(x)] - \beta \cdot \text{Divergence}$. Samples x are drawn from $\tilde{\pi}_{\text{old}} = \pi_{\text{old}}/Z_{\text{old}}$. These losses yield $-\nabla_\theta J(\theta)$ via differentiation. Notation: $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$, $R(x)$ is reward, β the regularization strength, and $Z_{\text{old}} = \int \pi_{\text{old}}$. Normalized counterparts are in Appendix G (Table 4).

Regularization (Unnormalized)	Surrogate loss (expectation w.r.t. $\tilde{\pi}_{\text{old}}$)
Forward (UFKL)	$Z_{\text{old}} \mathbb{E}[-w(x)R(x) + \beta(w(x) - \log w(x) - 1)]$
Reverse (URKL)	$Z_{\text{old}} \mathbb{E}[-w(x)R(x) + \beta(w(x) \log w(x) - w(x))]$

Definition 2.1 (Unnormalized Forward KL). The unnormalized forward KL divergence [Minka et al., 2005, Zhu and Rohwer, 1995] between the measure π_{old} and the density π_θ is defined as:

$$\text{UKL}(\pi_{\text{old}} \parallel \pi_\theta) = \underbrace{\int_x \pi_{\text{old}}(x) \log \frac{\pi_{\text{old}}(x)}{\pi_\theta(x)} dx}_{\text{Generalized KL}} + \underbrace{\int_x (\pi_\theta(x) - \pi_{\text{old}}(x)) dx}_{\text{Mass Correction}}.$$

This form is particularly relevant when dealing with reference measures that may not be perfectly normalized or when connecting to certain KL estimators like k_3 (see Remark D.2).

Consider the objective using UKL regularization as follows:

$$J_{\text{UFKL}}(\theta) = \mathbb{E}_{x \sim \pi_\theta}[R(x)] - \beta \text{UKL}(\pi_{\text{old}} \parallel \pi_\theta). \quad (2.1)$$

To estimate this off-policy using samples from the normalized reference $\tilde{\pi}_{\text{old}}(x) = \pi_{\text{old}}(x)/Z_{\text{old}}$, we define the importance weight $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$ (using the unnormalized π_{old}). The gradient and corresponding loss function, incorporating the total mass Z_{old} of the reference measure, are given in Proposition 2.2.

Proposition 2.2 (Policy Gradient and Differentiable Loss for Unnormalized Forward KL). Consider the unnormalized KL regularized objective function in Eq. (2.1). The gradient of $J_{\text{UFKL}}(\theta)$ is:

$$\nabla_\theta J_{\text{UFKL}}(\theta) = Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\left(w(x)R(x) - \beta(w(x) - 1) \right) \nabla_\theta \log \pi_\theta(x) \right].$$

The corresponding surrogate loss for gradient descent optimization, estimated using samples $\{x_i\} \sim \tilde{\pi}_{\text{old}}$, is:

$$\mathcal{L}_{\text{UFKL}}(\theta) = Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-w(x)R(x) + \beta(w(x) - \log w(x) - 1)],$$

satisfying $\nabla_\theta \mathcal{L}_{\text{UFKL}}(\theta) = -\nabla_\theta J_{\text{UFKL}}(\theta)$.

Table 2: Combined performance metrics with 4k context length on the AIME24, and AIME25 mathematical reasoning benchmarks, showing “Last” and “Best” scores. The “Last” score is from the 400th training step, assuming the training process remained stable to that point. The highest score in each column is **bolded**, and the second highest is underlined. RPG and RPG-REINFORCE methods are highlighted with light cyan and light green backgrounds, respectively.

Method	AIME24		AIME25	
	Last	Best	Last	Best
GRPO	0.3458	0.3677	0.2896	0.3042
DAPO	0.4063	<u>0.4479</u>	0.3510	0.3938
RPG-UFKL	0.4031	0.4396	0.3625	0.3979
RPG-URKL	0.3990	0.4219	0.3438	0.3792
RPG-REINFORCE-UFKL	<u>0.4281</u>	0.4375	<u>0.3771</u>	<u>0.4042</u>
RPG-REINFORCE-URKL	0.4458	0.4531	0.4125	0.4313

Remark 2.3 (Interpretation of UFKL Loss and Gradient). The regularization component of the surrogate loss $\mathcal{L}_{\text{UFKL}}(\theta)$, specifically $Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\beta(w(x) - \log w(x) - 1)]$, corresponds to an off-policy estimate of the unnormalized forward KL divergence term $\beta \cdot \text{UKL}(\pi_{\text{old}} \parallel \pi_{\theta})$ present in the objective $J_{\text{UFKL}}(\theta)$. This connection is established via the k_3 estimator (see Remark D.2 and Appendix F). Furthermore, the gradient term $-\beta(w(x) - 1)$ effectively modifies the reward, guiding π_{θ} to match not only the shape of π_{old} but also its overall mass Z_{old} , due to the mass correction component in $\text{UKL}(\pi_{\text{old}} \parallel \pi_{\theta})$.

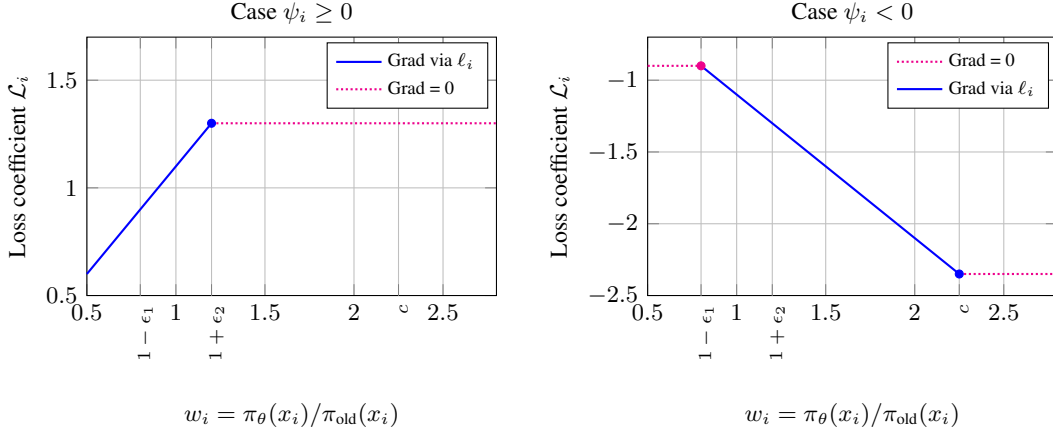


Figure 2: Visualization of the loss coefficient $\mathcal{L}_i$ vs. importance weight w_i based on the specific implementation in Algorithm 2. This version swaps the main branching condition compared to previous versions (branches on $\psi_i > 0$). The plot assumes $\ell_i = -\log \pi_{\theta}(x_i) = 1$ for visualizing the value of $\mathcal{L}_i$. The line styles indicate the nature of the gradient $\nabla_{\theta} \mathcal{L}_i$: **Solid blue:** Gradient exists, flowing only via ℓ_i . The coefficient multiplying $\nabla_{\theta} \ell_i$ depends on $\text{SG}(w_i)$. **Dotted magenta:** Gradient is zero. This occurs when ℓ_i is detached via SG in the loss calculation. Left: Case $\psi_i \geq 0$. Right: Case $\psi_i < 0$.

3 Experiments

In this section, we empirically evaluate our proposed Regularized Policy Gradient (RPG) framework, including both its fully differentiable (RPG) and REINFORCE-style (RPG-REINFORCE) variants. We compare their performance against established baselines on challenging mathematical reasoning tasks using large language models, including GRPO [Shao et al., 2024] and DAPO [Yu et al., 2025]. Our evaluation focuses on task-specific accuracy, training stability, and key training dynamics such as reward, policy entropy, and response length.

Tables 2 and 5 summarize the performance of our RPG algorithms against baselines with 4k and 2k context lengths, reporting both the last and best scores achieved during training on these benchmarks.

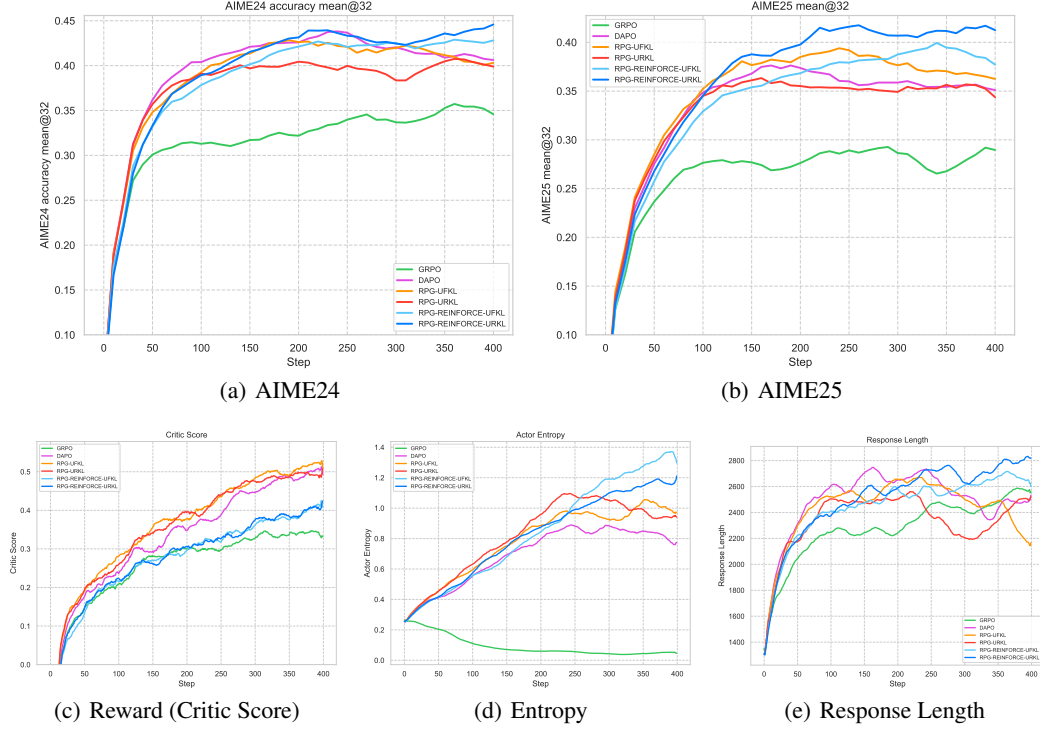


Figure 3: Performance of RPG and REINFORCE-Style Regularized Policy Gradient (RPG-REINFORCE) methods compared to baselines with 4k context length.

Figure 3 and 5 complement these results by illustrating the evaluation scores and training dynamics for the fully differentiable RPG variants and baselines when training the Qwen-3-4B model. These figures display performance on the AIME24 and AIME25 benchmarks, alongside key training metrics: reward (critic score), policy entropy, and average response length. Across settings, the RPG-REINFORCE variants with *RPG-Style Clip* consistently deliver the strongest results, surpassing DAPO and our differentiable RPG by up to +6 absolute points on AIME24/AIME25 at 4k context (see Figure 3).

The quantitative results in Table 2 demonstrate the competitive performance of the proposed RPG and REINFORCE-style RPG frameworks with 4k context length. On AIME24, RPG-REINFORCE variants lead, with RPG-REINFORCE-URKL achieving the best “Best” score (0.4531) and the best “Last” score (0.4458), while RPG-REINFORCE-URKL attain a second best “Last” score (0.4281). For AIME25, RPG-REINFORCE-URKL still achieves the top “Best” score (0.4313) and a strong “Last” score (0.4125) and RPG-REINFORCE-URKL is second only to that. Overall, RPG and RPG-REINFORCE methods rank at or near the top across benchmarks and metrics, while exhibiting stable training dynamics.

4 Conclusion

We introduced RPG, a framework for deriving and organizing KL-regularized policy gradient algorithms for online, off-policy RL. We provided derivations for policy gradients and surrogate loss functions covering forward/reverse KL, normalized/unnormalized distributions, and both fully differentiable and REINFORCE-style estimators. Beyond derivations, we revisited the classical REINFORCE algorithm and made it viable off-policy through *RPG-Style Clip* and iterative reference updates. On LLM reasoning, these design choices deliver stable and scalable training with competitive and superior accuracy relative to strong baselines.

References

- Marah Abdin, Jyoti Aneja, Hany Awadalla, Ahmed Awadallah, Ammar Ahmad Awan, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Jianmin Bao, Harkirat Behl, et al. Phi-3 technical report: A highly capable language model locally on your phone. *arXiv preprint arXiv:2404.14219*, 2024.
- Arash Ahmadian, Chris Cremer, Matthias Gall , Marzieh Fadaee, Julia Kreutzer, Olivier Pietquin, Ahmet  st n, and Sara Hooker. Back to basics: Revisiting reinforce-style optimization for learning from human feedback in llms. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 12248–12267, 2024.
- Mohammad Gheshlaghi Azar, Zhaohan Daniel Guo, Bilal Piot, Remi Munos, Mark Rowland, Michal Valko, and Daniele Calandriello. A general theoretical paradigm to understand learning from human preferences. In *International Conference on Artificial Intelligence and Statistics*, pages 4447–4455. PMLR, 2024.
- Ralph Allan Bradley and Milton E Terry. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952.
- Zixiang Chen, Yihe Deng, Huizhuo Yuan, Kaixuan Ji, and Quanquan Gu. Self-play fine-tuning converts weak language models to strong language models. In *Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024*, 2024.
- Paul F. Christiano, Jan Leike, Tom B. Brown, Miljan Martic, Shane Legg, and Dario Amodei. Deep reinforcement learning from human preferences. In Isabelle Guyon, Ulrike von Luxburg, Samy Bengio, Hanna M. Wallach, Rob Fergus, S. V. N. Vishwanathan, and Roman Garnett, editors, *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, December 4-9, 2017, Long Beach, CA, USA*, pages 4299–4307, 2017.
- Hanze Dong, Wei Xiong, Deepanshu Goyal, Yihan Zhang, Winnie Chow, Rui Pan, Shizhe Diao, Jipeng Zhang, Kashun Shum, and Tong Zhang. RAFT: reward ranked finetuning for generative foundation model alignment. *Trans. Mach. Learn. Res.*, 2023, 2023.
- Kawin Ethayarajh, Winnie Xu, Niklas Muennighoff, Dan Jurafsky, and Douwe Kiela. Kto: Model alignment as prospect theoretic optimization. *arXiv preprint arXiv:2402.01306*, 2024.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, et al. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*, 2024.
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- Jian Hu. Reinforce++: A simple and efficient approach for aligning large language models. *arXiv preprint arXiv:2501.03262*, 2025.
- Kaixuan Ji, Guanlin Liu, Ning Dai, Qingping Yang, Renjie Zheng, Zheng Wu, Chen Dun, Quanquan Gu, and Lin Yan. Enhancing multi-step reasoning abilities of language models through direct q-function optimization. *arXiv preprint arXiv:2410.09302*, 2024.
- Amirhossein Kazemnejad, Milad Aghajohari, Eva Portelance, Alessandro Sordani, 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.
- Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the 29th Symposium on Operating Systems Principles*, pages 611–626, 2023.
- Ziniu Li, Tian Xu, Yushun Zhang, Zhihang Lin, Yang Yu, Ruoyu Sun, and Zhi-Quan Luo. Remax: A simple, effective, and efficient reinforcement learning method for aligning large language models. In *Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024*, 2024.

- Tianqi Liu, Yao Zhao, Rishabh Joshi, Misha Khalman, Mohammad Saleh, Peter J. Liu, and Jialu Liu. Statistical rejection sampling improves preference optimization. In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net, 2024.
- 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.
- Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *7th International Conference on Learning Representations, ICLR 2019, New Orleans, LA, USA, May 6-9, 2019*, 2019.
- Mathematical Association of America’s American Mathematics Competitions MAA. 2024 AIME-I, 2024a. URL https://artofproblemsolving.com/wiki/index.php/2024_AIME_I. Accessed: 2025-05-08.
- Mathematical Association of America’s American Mathematics Competitions MAA. 2024 AIME-II, 2024b. URL https://artofproblemsolving.com/wiki/index.php/2024_AIME_II. Accessed: 2025-05-08.
- Mathematical Association of America’s American Mathematics Competitions MAA. 2025 AIME-I, 2025a. URL https://artofproblemsolving.com/wiki/index.php/2025_AIME_I.
- Mathematical Association of America’s American Mathematics Competitions MAA. 2025 AIME-II, 2025b. URL https://artofproblemsolving.com/wiki/index.php/2025_AIME_II.
- Yu Meng, Mengzhou Xia, and Danqi Chen. Simpo: Simple preference optimization with a reference-free reward. *Advances in Neural Information Processing Systems*, 37:124198–124235, 2024.
- Tom Minka et al. Divergence measures and message passing. *Technical report, Microsoft Research*, 2005.
- Rémi Munos, Michal Valko, Daniele Calandriello, Mohammad Gheshlaghi Azar, Mark Rowland, Zhaohan Daniel Guo, Yunhao Tang, Matthieu Geist, Thomas Mesnard, Côme Fiegel, Andrea Michi, Marco Selvi, Sertan Girgin, Nikola Momchev, Olivier Bachem, Daniel J. Mankowitz, Doina Precup, and Bilal Piot. Nash learning from human feedback. In *Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024*, 2024.
- OpenAI. ChatGPT, 2022. URL <https://chat.openai.com/>.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in Neural Information Processing Systems*, 35: 27730–27744, 2022.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36:53728–53741, 2023.
- John Schulman. Approximating kl divergence. <http://joschu.net/blog/kl-approx.html>, March 2020. Accessed on October 20, 2025.
- John Schulman, Philipp Moritz, Sergey Levine, Michael I. Jordan, and Pieter Abbeel. High-dimensional continuous control using generalized advantage estimation. In Yoshua Bengio and Yann LeCun, editors, *4th International Conference on Learning Representations, ICLR 2016, San Juan, Puerto Rico, May 2-4, 2016, Conference Track Proceedings*, 2016.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- 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.

- Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. Hybridflow: A flexible and efficient RLHF framework. In *Proceedings of the Twentieth European Conference on Computer Systems, EuroSys 2025, Rotterdam, The Netherlands, 30 March 2025 - 3 April 2025*, pages 1279–1297. ACM, 2025.
- Richard S Sutton, Andrew G Barto, et al. *Reinforcement learning: An introduction*, volume 1. MIT press Cambridge, 1998.
- 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.
- Qwen Team. Qwen3 technical report, 2025. URL <https://arxiv.org/abs/2505.09388>.
- Ronald J Williams. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine learning*, 8:229–256, 1992.
- Yue Wu, Zhiqing Sun, Huizhuo Yuan, Kaixuan Ji, Yiming Yang, and Quanquan Gu. Self-play preference optimization for language model alignment. In *The Thirteenth International Conference on Learning Representations, ICLR 2025, Singapore, April 24-28, 2025*, 2025.
- Wei Xiong, Hanze Dong, Chenlu Ye, Ziqi Wang, Han Zhong, Heng Ji, Nan Jiang, and Tong Zhang. Iterative preference learning from human feedback: bridging theory and practice for rlhf under kl-constraint. In *Proceedings of the 41st International Conference on Machine Learning*, pages 54715–54754, 2024.
- Jing Xu, Andrew Lee, Sainbayar Sukhbaatar, and Jason Weston. Some things are more cringe than others: Preference optimization with the pairwise cringe loss. *arXiv preprint arXiv:2312.16682*, 18, 2023.
- An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, et al. Qwen2. 5 technical report. *arXiv preprint arXiv:2412.15115*, 2024.
- Deheng Ye, Zhao Liu, Mingfei Sun, Bei Shi, Peilin Zhao, Hao Wu, Hongsheng Yu, Shaojie Yang, Xipeng Wu, Qingwei Guo, et al. Mastering complex control in moba games with deep reinforcement learning. In *Proceedings of the AAAI conference on artificial intelligence*, volume 34, pages 6672–6679, 2020.
- 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. *arXiv preprint arXiv:2503.14476*, 2025.
- 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.
- Zheng Yuan, Hongyi Yuan, Chengpeng Li, Guanting Dong, Keming Lu, Chuanqi Tan, Chang Zhou, and Jingren Zhou. Scaling relationship on learning mathematical reasoning with large language models. *arXiv preprint arXiv:2308.01825*, 2023.
- Yifan Zhang, Ge Zhang, Yue Wu, Kangping Xu, and Quanquan Gu. Beyond bradley-terry models: A general preference model for language model alignment. In *Proceedings of the 42nd International Conference on Machine Learning*, 2025.
- Yao Zhao, Rishabh Joshi, Tianqi Liu, Misha Khalman, Mohammad Saleh, and Peter J Liu. Slic-hf: Sequence likelihood calibration with human feedback. *arXiv preprint arXiv:2305.10425*, 2023.
- Chujie Zheng, Shixuan Liu, Mingze Li, Xiong-Hui Chen, Bowen Yu, Chang Gao, Kai Dang, Yuqiong Liu, Rui Men, An Yang, et al. Group sequence policy optimization. *arXiv preprint arXiv:2507.18071*, 2025.
- Huaiyu Zhu and Richard Rohwer. Information geometric measurements of generalisation. *Preprint*, 1995.

Appendix

A Related Work	11
B Preliminaries	12
B.1 KL Regularization in Policy Gradients	12
B.2 Group Relative Policy Optimization (GRPO)	13
C REINFORCE and Proximal Policy Optimization (PPO)	13
C.1 REINFORCE	13
C.2 Proximal Policy Optimization (PPO)	14
D Unnormalized Reverse KL Regularization	14
E REINFORCE-Style Regularized Policy Gradients	15
E.1 RPG-Style Clip: dual-clip truncation of importance ratios	15
F Equivalence of k_3 Estimator and Unnormalized KL Divergence	16
G Normalized KL Regularization	16
G.1 Forward KL Regularization	16
G.2 Reverse KL Regularization	17
H REINFORCE-Style Regularized Policy Gradients with Various KL Regularization Forms	17
H.1 Rationale for REINFORCE-Style Loss Formulation	17
H.2 REINFORCE-Style RPG with Forward KL Regularization	17
H.3 REINFORCE-Style RPG with Unnormalized Forward KL Regularization	18
H.4 REINFORCE-Style RPG with Reverse KL Regularization	18
H.5 REINFORCE-Style RPG with Unnormalized Reverse KL Regularization	18
I More on Algorithmic Details	18
I.1 Stabilization Techniques for Regularized Policy Gradients	18
I.2 Stabilization Techniques for REINFORCE-Style Regularized Policy Gradients	19
J Detailed Experimental Setup	21
K More Experiment Results	23
K.1 The performance with 2k context length	23
K.2 Ablation Study	24
L Proofs of Theorem B.1 (Generalized Policy Gradient Theorem)	27
M Proofs for Regularized Policy Gradients	27
M.1 Proof of Proposition G.1 (Policy Gradient and Differentiable Loss for Normalized Forward KL)	27

M.2	Proof of Proposition 2.2 (Policy Gradient and Differentiable Loss for Unnormalized Forward KL)	28
M.3	Proof of Proposition G.3 (Policy Gradient and Differentiable Loss for Normalized Reverse KL)	29
M.4	Proof of Proposition D.3 (Policy Gradient and Differentiable Loss for Unnormalized Reverse KL)	30
N	Proofs for REINFORCE-Style Regularized Policy Gradients	31
N.1	Proof of Proposition H.1 (REINFORCE-style Policy Gradient for Forward KL) . .	32
N.2	Proof of Proposition H.3 ((REINFORCE-style Policy Gradient for Unnormalized Forward KL)	32
N.3	Proof of Proposition H.4 (REINFORCE-Style Loss)	33
N.4	Proof of Proposition H.5 (REINFORCE-Style Loss for Unnormalized Reverse KL)	33
O	Broader Impact and Limitations	33
P	NeurIPS Paper Checklist	34

A Related Work

Fine-tuning large language models (LLMs) using human feedback has become a critical step in developing capable and aligned AI systems. Broadly, methods fall into two main categories: those relying on policy optimization using an explicit reward model learned from feedback, and those directly optimizing policies based on preference data.

RLHF via Policy Optimization. The classic RLHF involves training a reward model (RM) $r_\phi(x, y)$ to predict human preferences and then using reinforcement learning to optimize the language model policy π_θ to maximize the expected reward from the RM, often regularizing against deviating too far from an initial reference policy π_{ref} . This approach was pioneered by Christiano et al. [2017] and gained widespread prominence with its application to LLMs like InstructGPT [Ouyang et al., 2022] and ChatGPT [OpenAI, 2022], which utilized Proximal Policy Optimization (PPO) [Schulman et al., 2017]. PPO became a workhorse due to its relative stability, achieved by constraining policy updates via a clipped surrogate objective. The standard PPO setup for RLHF involves the policy π_θ , a value function V_ψ , the RM r_ϕ , and the reference policy π_{ref} .

RLHF via Direct Preference Optimization. An alternative and increasingly popular approach bypasses explicit reward modeling by directly optimizing the policy π_θ based on preference data, typically pairwise comparisons (y_w, y_l) indicating that response y_w is preferred over y_l for a given prompt x . Inspired by the Bradley-Terry model [Bradley and Terry, 1952], Direct Preference Optimization (DPO) [Rafailov et al., 2023] derived a simple loss function directly relating preference probabilities to policy likelihoods under π_θ and a reference policy π_{ref} . DPO maximizes the relative likelihood of preferred responses using a logistic loss: $\mathcal{L}_{\text{DPO}} \propto -\mathbb{E}[\log \sigma(\beta \Delta \log p)]$, where $\Delta \log p$ is the difference in log-probabilities of y_w and y_l between π_θ and π_{ref} . DPO’s simplicity and effectiveness led to its wide adoption in models like Llama-3 [Grattafiori et al., 2024], Qwen2 [Yang et al., 2024], and Phi-3 [Abdin et al., 2024]. Numerous variants have followed: SLiC-HF [Zhao et al., 2023] uses a pairwise hinge loss for calibration; IPO [Azar et al., 2024] uses an identity link function; SimPO [Meng et al., 2024] offers a simpler objective focusing on the margin; KTO [Ethayarajh et al., 2024] handles binary (good/bad) feedback; DQO [Ji et al., 2024] incorporates direct Q-value modeling; RAFT [Dong et al., 2023], RSO [Liu et al., 2024] and RFT [Yuan et al., 2023] use a rejection sampling perspective. Recognizing that preferences might evolve, iterative methods like Iterative DPO [Xiong et al., 2024], PCO [Xu et al., 2023] and SPIN [Chen et al., 2024] alternate between generation/preference learning and policy updates, often using the current policy’s outputs in a self-improvement loop. Game theory offers another lens, with Nash Learning from Human Feedback (NLHF) [Munos et al., 2024] framing RLHF as finding a Nash equilibrium between policies. Self-play ideas appear in SPPO [Wu et al., 2025] and GPO [Zhang et al., 2025], where the policy generates pairs for comparison. Methods like GPM [Zhang et al., 2025] aim to handle more general preference structures efficiently using latent embeddings beyond pairwise comparisons.

RL for Enhancing LLM Reasoning. Beyond general alignment with human preferences, RL techniques are increasingly explored to specifically enhance the multi-step reasoning capabilities of LLMs in domains like mathematics, coding, and complex instruction following. In these contexts, RL optimizes the policy to generate sequences (e.g., chain-of-thought, code blocks) that lead to successful outcomes, often using rewards derived from external feedback like unit test results, execution outcomes, or correctness checks by an automated judge or specialized reward model trained on reasoning quality. For instance, the DeepSeekMath model [Shao et al., 2024] employed the GRPO algorithm, a value-free PPO variant, demonstrating significant improvements in mathematical problem-solving benchmarks through RL fine-tuning. DeepSeek-R1 [Guo et al., 2025] represents efforts in applying advanced techniques potentially involving RL for complex tasks, although specific methods might vary. Furthermore, preference-based methods like SPPO and GPO have been applied to reasoning-specialized models such as Kimi-1.5 [Team et al., 2025], and the resulting improvements observed on benchmarks involving coding and math suggest that preference-based RLHF can also contribute to refining reasoning abilities, potentially by optimizing implicit properties related to logical consistency and correctness within the preference data. The need for a value function (critic model) used in PPO incurs significant computational costs, and standard PPO can face stability challenges with sparse rewards common in LLM tasks. Addressing these issues has driven recent work. Several methods aim to improve efficiency by removing the value network: ReMax [Li et al., 2024] adapts REINFORCE [Williams, 1992] using Monte Carlo returns and normalization; GRPO [Shao et al., 2024] uses a group-average reward baseline and adds a k_3 -based KL penalty to the objective; and VinePPO [Kazemnejad et al., 2024] uses MC sampling from intermediate

steps. Other approaches focus on stability and alternative baselines, such as RLOO [Ahmadian et al., 2024], which uses leave-one-out statistics within a group, and REINFORCE++ [Hu, 2025], which enhances REINFORCE with token-level KL penalties (using the k_2 estimator) and normalization. Dr. GRPO [Liu et al., 2025] identifies and corrects a bias found in GRPO’s advantage estimators, DAPO [Yu et al., 2025] introduces strategies like Clip-Higher, reward over-sampling, and a token-level loss to handle long sequences and entropy collapse, while VAPO [Yuan et al., 2025] builds upon it with length-adaptive advantage estimation. Recently, GSPO [Zheng et al., 2025] was proposed with sequence-level rewards and used in the Qwen3 model series [Team, 2025].

Our contribution is to make the off-policy weighting and estimator equivalences explicit across normalized/unnormalized variants, to identify a bias introduced when these weights are omitted (as in the GRPO KL term), and to provide corrected surrogates that are gradient-equivalent to the intended objectives. The design-space view makes transparent how several recent algorithms arise as special cases.

B Preliminaries

Policy gradient (PG) methods are a cornerstone of modern reinforcement learning (RL), optimizing parameterized policies π_θ by estimating the gradient of an expected objective function $J(\theta)$ with respect to the policy parameters θ . Typically, $J(\theta)$ represents the expected cumulative discounted reward over trajectories $\tau = (s_0, a_0, r_0, s_1, \dots, s_T, a_T, r_T)$ generated by the policy: $J(\theta) = \mathbb{E}_{\tau \sim \pi_\theta} [G(\tau)]$, where $G(\tau) = \sum_{t=0}^T \gamma^t r_t$ is the trajectory return (with discount factor γ), and the expectation is taken over the trajectories sampled according to the policy $\pi_\theta(a|s)$ and the environment dynamics $p(s'|s, a)$. The Generalized Policy Gradient Theorem (GPPT) provides a foundation for deriving these gradients (see Appendix L for the proof).

Proposition B.1 (Generalized Policy Gradient Theorem). Let $\pi_\theta(x)$ be a probability density or mass function parameterized by θ , representing the probability of sampling item x . Let $f(x, \theta)$ be a scalar-valued function associated with x , potentially depending on θ . Under suitable regularity conditions, the gradient of the expectation $\mathbb{E}_{x \sim \pi_\theta} [f(x, \theta)]$ with respect to θ is:

$$\nabla_\theta \mathbb{E}_{x \sim \pi_\theta} [f(x, \theta)] = \mathbb{E}_{x \sim \pi_\theta} [f(x, \theta) \nabla_\theta \log \pi_\theta(x) + \nabla_\theta f(x, \theta)]. \quad (\text{B.1})$$

The term $\mathbb{E}[f \nabla \log \pi]$ reflects how changes in θ affect the probability of sampling x , while $\mathbb{E}[\nabla f]$ reflects how changes in θ directly affect the function value f .

The classic REINFORCE algorithm [Williams, 1992] applies the GPPT to the standard RL objective $J(\theta) = \mathbb{E}_{\tau \sim \pi_\theta} [G(\tau)]$. In this case, $f(\tau, \theta) = G(\tau)$, the total trajectory return, which does not depend directly on θ (i.e., $\nabla_\theta G(\tau) = 0$). The theorem simplifies, and the gradient can be expressed using per-timestep contributions [Sutton et al., 1998]:

$$\nabla_\theta J(\theta) = \mathbb{E}_{\tau \sim \pi_\theta} \left[\sum_{t=0}^T G_t \nabla_\theta \log \pi_\theta(a_t | s_t) \right],$$

where $G_t = \sum_{k=t}^T \gamma^{k-t} r_k$ is the return-to-go from timestep t . Due to space limit, we defer the detailed introduction of REINFORCE to Appendix C.1.

B.1 KL Regularization in Policy Gradients

A common technique to stabilize policy optimization, especially in off-policy settings or when fine-tuning large models, is regularization. The Kullback-Leibler (KL) divergence is frequently used to penalize the deviation of the learned policy π_θ from a reference policy π_{ref} (which could be π_{old} , an initial supervised fine-tuned model, or another prior). $\text{KL}(P \| Q) \geq 0$ with equality iff $P = Q$ almost everywhere. It is asymmetric (i.e., $\text{KL}(P \| Q) \neq \text{KL}(Q \| P)$). Minimizing the forward KL $\text{KL}(\pi_{\text{ref}} \| \pi_\theta)$ encourages π_θ to cover the support of π_{ref} (zero-forcing), while minimizing the reverse KL $\text{KL}(\pi_\theta \| \pi_{\text{ref}})$ encourages π_θ to be concentrated where π_{ref} has high probability mass (mode-seeking).

Adding a KL penalty to the RL objective, such as $J(\theta) = \mathbb{E}_{\pi_\theta} [R] - \beta \text{KL}(\pi_\theta \| \pi_{\text{ref}})$, helps control the policy update size, prevents large deviations from π_{ref} , encourages exploration near known good policies, and can mitigate issues like catastrophic forgetting or overly confident outputs, particularly relevant in LLM fine-tuning [Ouyang et al., 2022]. For PPO (see Appendix C.2), this penalty can be incorporated implicitly via reward shaping: $r'_t = r_t - \beta \log(\pi_\theta(a_t | s_t) / \pi_{\text{ref}}(a_t | s_t))$. Alternatively, it

can be added explicitly to the objective function, as in GRPO. The specific form of the KL divergence (forward/reverse), whether distributions are normalized (KL vs. UKL), and the choice of estimator (e.g., standard log-ratio vs. k_3 estimator [Schulman, 2020]) can vary, leading to different properties (mode seeking v.s. zero-forcing) and gradient estimators, as explored later in this paper (Sections 2 and E).

B.2 Group Relative Policy Optimization (GRPO)

Group Relative Policy Optimization (GRPO) [Shao et al., 2024] adapts the PPO framework for training LLMs, notably by eliminating the need for a learned value function (critic). Instead of using GAE, GRPO estimates the advantage $\hat{A}_{i,t}$ at token t of output o_i based on the relative rewards within a group of G outputs $\{o_1, \dots, o_G\}$ sampled from the old policy $\pi_{\theta_{\text{old}}}$ for the same prompt q .

Crucially, GRPO modifies the PPO objective by explicitly adding a KL regularization term directly to the objective function. Its objective (simplified notation) is:

$$\mathcal{J}_{\text{GRPO}}(\theta) = \mathbb{E}_{q \sim P(Q), \{o_i\} \sim \pi_{\text{old}}} \left[\frac{1}{G} \sum_{i=1}^G \frac{1}{|o_i|} \sum_{t=1}^{|o_i|} \left(J_{i,t}^{\text{Clip}}(\theta) - \beta \cdot \text{KL}_{\text{est}}(\pi_{\theta}(\cdot|h_{i,t}) \parallel \pi_{\text{ref}}(\cdot|h_{i,t})) \right) \right],$$

where $h_{i,t} = (q, o_{i,<t})$ is the history, $J_{i,t}^{\text{Clip}}(\theta)$ represents the PPO-Clip term from Eq. (C.3) applied using the group-relative advantage estimate $\hat{A}_{i,t}$, and π_{ref} is a reference model (e.g., the initial SFT model). For the KL penalty, GRPO employs the k_3 estimator form [Schulman, 2020], evaluated per token $o_{i,t}$:

$$\text{KL}_{\text{est}}(\pi_{\theta} \parallel \pi_{\text{ref}}) \approx k_3 \left(\frac{\pi_{\text{ref}}(o_{i,t} | h_{i,t})}{\pi_{\theta}(o_{i,t} | h_{i,t})} \right) = \frac{\pi_{\text{ref}}(o_{i,t} | h_{i,t})}{\pi_{\theta}(o_{i,t} | h_{i,t})} - \log \frac{\pi_{\text{ref}}(o_{i,t} | h_{i,t})}{\pi_{\theta}(o_{i,t} | h_{i,t})} - 1.$$

This uses the functional form $k_3(y) = y - \log y - 1$ as discussed in Schulman [2020], applied with $y = \pi_{\text{ref}}(o_{i,t} | h_{i,t}) / \pi_{\theta}(o_{i,t} | h_{i,t})$. This form is related to the unnormalized reverse KL divergence, $\text{UKL}(\pi_{\theta} \parallel \pi_{\text{ref}})$ (see Section D and Appendix F for a detailed discussion).

$w_{i,t} = \frac{\pi_{\theta}(o_{i,t} | h_{i,t})}{\pi_{\text{old}}(o_{i,t} | h_{i,t})}$ multiplying the k_3 term. The direct subtraction without this weight means the gradient derived from GRPO’s objective does not, in general, correspond to the gradient of the intended off-policy objective $\mathcal{J}^{\text{Clip}} - \beta \text{UKL}(\pi_{\theta} \parallel \pi_{\text{ref}})$. For clarity, a corrected off-policy estimator for the GRPO KL component at history $h_{i,t}$ is

$$\widehat{\text{KL}}_{\text{GRPO-corrected}}(h_{i,t}; \theta) = \mathbb{E}_{o_{i,t} \sim \pi_{\text{old}}(\cdot|h_{i,t})} \left[w_{i,t} k_3 \left(\frac{\pi_{\text{ref}}(o_{i,t} | h_{i,t})}{\pi_{\theta}(o_{i,t} | h_{i,t})} \right) \right],$$

which is consistent with URKL/UKL depending on direction (see Section 2 and Appendix F). Our results in Section 2 provide derivations for KL-regularized objectives that explicitly account for off-policy sampling via importance weights. Related work is detailed in Appendix A.

C REINFORCE and Proximal Policy Optimization (PPO)

C.1 REINFORCE

REINFORCE performs Monte Carlo (MC) updates after sampling a complete trajectory, using the sampled return G_t as an unbiased estimate of the state-action value function $Q^{\pi_{\theta}}(s_t, a_t)$. However, these MC estimates often exhibit high variance, leading to slow and unstable learning.

To reduce variance, a state-dependent baseline $b(s_t)$ (commonly an estimate of the state value function, $V^{\pi_{\theta}}(s_t)$) is subtracted from the return-to-go:

$$\nabla_{\theta} J(\theta) = \mathbb{E}_{\tau \sim \pi_{\theta}} \left[\sum_{t=0}^T (G_t - b(s_t)) \nabla_{\theta} \log \pi_{\theta}(a_t | s_t) \right] = \mathbb{E}_{\tau \sim \pi_{\theta}} \left[\sum_{t=0}^T \hat{A}_t \nabla_{\theta} \log \pi_{\theta}(a_t | s_t) \right]. \quad (\text{C.1})$$

Here, $\hat{A}_t = G_t - b(s_t)$ is an estimate of the advantage function $A^{\pi_{\theta}}(s_t, a_t) = Q^{\pi_{\theta}}(s_t, a_t) - V^{\pi_{\theta}}(s_t)$. Subtracting a baseline that only depends on the state s_t does not bias the gradient estimate, since $\mathbb{E}_{a_t \sim \pi_{\theta}(\cdot|s_t)} [b(s_t) \nabla_{\theta} \log \pi_{\theta}(a_t | s_t)] = b(s_t) \nabla_{\theta} \sum_{a_t} \pi_{\theta}(a_t | s_t) = b(s_t) \nabla_{\theta} 1 = 0$. REINFORCE with baseline is typically implemented by minimizing the loss:

$$\mathcal{L}_{\text{REINFORCE}}(\theta) = -\mathbb{E}_{\tau \sim \pi_{\theta}} \left[\sum_{t=0}^T \text{SG}(\hat{A}_t) \log \pi_{\theta}(a_t | s_t) \right], \quad (\text{C.2})$$

using the stop-gradient operator $\text{SG}(\cdot)$ to prevent gradients from flowing into the advantage estimate $\hat{A}_t$. As REINFORCE uses samples collected under the current policy π_θ for gradient estimation, it is an on-policy algorithm.

C.2 Proximal Policy Optimization (PPO)

On-policy methods like REINFORCE can be sample-inefficient, requiring new trajectories for each gradient update. Proximal Policy Optimization (PPO) [Schulman et al., 2017] improves stability and sample efficiency by enabling multiple updates using the same batch of data collected under a slightly older policy π_{old} . This makes PPO effectively off-policy. PPO achieves this by optimizing a surrogate objective function that discourages large deviations between the current policy π_θ and the old policy π_{old} . The most widely used variant, PPO-Clip, employs a clipped objective:

$$J^{\text{PPO-Clip}}(\theta) = \mathbb{E}_t \left[\min \left(w_t(\theta) \hat{A}_t, \text{clip}(w_t(\theta), 1 - \epsilon, 1 + \epsilon) \hat{A}_t \right) \right], \quad (\text{C.3})$$

where the expectation $\mathbb{E}_t$ is taken over timesteps in the collected batch sampled from π_{old} . Here, $w_t(\theta) = \frac{\pi_\theta(a_t|s_t)}{\pi_{\text{old}}(a_t|s_t)}$ is the importance sampling ratio. $\hat{A}_t$ is an advantage estimate, typically computed using Generalized Advantage Estimation (GAE) [Schulman et al., 2016], which leverages observed rewards and a learned state-value function $V(s)$ to reduce variance.

Notably, in many practical implementations, especially in Reinforcement Learning from Human Feedback (RLHF) for large language models [Ouyang et al., 2022], a KL divergence penalty against a reference policy π_{ref} (e.g., the initial supervised model) is often incorporated implicitly by modifying the reward signal before calculating the advantage. For example, the reward used for GAE calculation might become $r'_t = r_t - \beta \log(\pi_\theta(a_t|s_t)/\pi_{\text{ref}}(a_t|s_t))$. When this r'_t is used within GAE to compute $\hat{A}_t$, the KL penalty term is effectively folded into the advantage estimate that multiplies the importance weight $w_t(\theta)$ in the objective function. This approach contrasts with adding the KL penalty as a separate term to the final objective, as seen in GRPO (Section B.2) or the formal derivations in Section 2.

The hyperparameter ϵ (e.g., 0.2) defines the clipping range $[1 - \epsilon, 1 + \epsilon]$ for the importance ratio $w_t(\theta)$. This clipping limits the influence of potentially noisy importance weights when the policy changes significantly, preventing destructive updates and further stabilizing the off-policy training. PPO optimizes the policy π_θ by maximizing $J^{\text{PPO-Clip}}(\theta)$.

D Unnormalized Reverse KL Regularization

Similar to the forward case, we can define an unnormalized reverse KL divergence, relaxing the normalization constraint on the reference distribution π_{old} . Let $\pi_{\text{old}}(x)$ be a potentially unnormalized reference measure with total mass $Z_{\text{old}} = \int \pi_{\text{old}}(x) dx$. Let $\tilde{\pi}_{\text{old}}(x) = \pi_{\text{old}}(x)/Z_{\text{old}}$ be the corresponding normalized probability distribution.

Definition D.1 (Unnormalized Reverse KL). The unnormalized reverse KL divergence between the density π_θ and the measure π_{old} is defined as:

$$\text{UKL}(\pi_\theta \| \pi_{\text{old}}) = \underbrace{\int_x \pi_\theta(x) \log \frac{\pi_\theta(x)}{\pi_{\text{old}}(x)} dx}_{\text{Generalized KL}} + \underbrace{\int_x (\pi_{\text{old}}(x) - \pi_\theta(x)) dx}_{\text{Mass Correction}}.$$

The mass correction term simplifies to $Z_{\text{old}} - \int \pi_\theta(x) dx$.

Remark D.2. (Equivalence to k_3 estimator) The k_3 estimator [Schulman, 2020], often used for its empirical properties (e.g., in GRPO [Shao et al., 2024]), is defined for a density ratio $y(x)$ as:

$$k_3(y) := y - 1 - \log y. \quad (\text{D.1})$$

As shown in Appendix F, this functional form directly relates to unnormalized KL divergences. For instance, $\text{KL}_{k_3}(\pi_\theta \| \pi_{\text{old}}) := \mathbb{E}_{x \sim \pi_\theta} [k_3(\pi_{\text{old}}(x)/\pi_\theta(x))]$ is equivalent to $\text{UKL}(\pi_\theta \| \pi_{\text{old}})$. This equivalence relationship justifies the exploration of UKL/URKL formulations within our framework.

Consider the objective using URKL:

$$J_{\text{URKL}}(\theta) = \mathbb{E}_{x \sim \pi_\theta} [R(x)] - \beta \text{UKL}(\pi_\theta \| \pi_{\text{old}}), \quad (\text{D.2})$$

where UKL is defined above. As with UFKL, we derive the gradient and loss using expectations over the normalized reference $\tilde{\pi}_{\text{old}}$ and the importance weight $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$ (with unnormalized π_{old}). The results are summarized in Proposition D.3.

Proposition D.3 (Policy Gradient and Differentiable Loss for Unnormalized Reverse KL). Consider the reverse unnormalized KL regularized objective function in Eq. (D.2). The gradient of $J_{\text{URKL}}(\theta)$ is:

$$\nabla_{\theta} J_{\text{URKL}}(\theta) = Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[w(x) \left(R(x) - \beta \log w(x) \right) \nabla_{\theta} \log \pi_{\theta}(x) \right].$$

A corresponding surrogate loss for gradient descent optimization, estimated using samples $\{x_i\} \sim \tilde{\pi}_{\text{old}}$, is:

$$\mathcal{L}_{\text{URKL}}(\theta) = Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[-w(x) R(x) + \beta (w(x) \log w(x) - w(x)) \right],$$

satisfying $\nabla_{\theta} \mathcal{L}_{\text{URKL}}(\theta) = -\nabla_{\theta} J_{\text{URKL}}(\theta)$. The constant Z_{old} scales the loss and gradient and may be omitted in practice.

Remark D.4 (URKL Loss and Mass Correction). The surrogate loss $\mathcal{L}_{\text{URKL}}(\theta)$ is designed such that its gradient is $-\nabla_{\theta} J_{\text{URKL}}(\theta)$. Specifically, the term $Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\beta (w(x) \log w(x) - w(x))]$ in the loss directly relates to the off-policy estimation of the unnormalized reverse KL divergence $\beta \text{UKL}(\pi_{\theta} \parallel \pi_{\text{old}})$, omitting a constant related to the total mass Z_{old} which does not affect the gradient. The policy gradient’s effective reward scaling factor, $(R(x) - \beta \log w(x))$, is simpler than its normalized RKL counterpart.

E REINFORCE-Style Regularized Policy Gradients

In Section 2, we derived policy gradient estimators and corresponding fully differentiable surrogate losses $\mathcal{L}(\theta)$ for KL-regularized objectives. Those losses were constructed such that $\nabla_{\theta} \mathcal{L}(\theta) = -\nabla_{\theta} J(\theta)$ directly, typically by setting $\mathcal{L}(\theta) = -J_{\text{IS}}(\theta)$ (where J_{IS} is the importance-sampled objective) up to constants. Notice that the gradients derived in Section 2 (Theorems 2.2 through D.3) share a structural similarity with the REINFORCE estimator:

$$\nabla_{\theta} J(\theta) = \mathbb{E}_{x \sim \pi_{\text{sampling}}} [\text{Weight}(x, \theta) \nabla_{\theta} \log \pi_{\theta}(x)]$$

where π_{sampling} is π_{old} or its normalized version $\tilde{\pi}_{\text{old}}$, and $\text{Weight}(x, \theta)$ encapsulates the reward and KL regularization terms, differing for each specific objective.

Proposition E.1 (Gradient-Equivalence of Surrogates). For each KL-regularized objective $J(\theta)$ derived in Section 2, the corresponding REINFORCE-style losses in Table 3 satisfy $\nabla_{\theta} \mathcal{L}(\theta) = -\nabla_{\theta} J(\theta)$ under the standard regularity assumptions used in the policy-gradient theorem. In particular, the stop-gradient operator ensures that dependence of the weight on θ (through importance ratios) does not leak unintended gradients. A proof sketch follows directly from the policy-gradient theorem and is completed in Appendix N.

This structural similarity motivates an alternative REINFORCE-style implementation using the stop-gradient operator SG. The general form of such losses and the detailed rationale for how they yield the target gradient via automatic differentiation are presented in Appendix H.1 (see Eq. (H.1)).

We explore these REINFORCE-style estimators as part of our framework, as they offer an alternative implementation path and demonstrate competitive empirical performance (Section 3). Proofs are in Appendix N. In the main text, we tabulate the unnormalized REINFORCE-style losses; normalized counterparts are deferred to Appendix H.

Table 3: REINFORCE-style surrogate losses $\mathcal{L}(\theta)$ for unnormalized KL-regularized objectives using the stop-gradient operator (SG). These losses yield the target gradient via automatic differentiation. Compare with the fully differentiable losses in Table 1. Normalized versions are given in Appendix H.

Regularization (Unnormalized)	REINFORCE-style loss (sampling $x \sim \tilde{\pi}_{\text{old}}$)
Forward (UFKL)	$-\mathbb{E} \left[\text{SG}(Z_{\text{old}}(w(x)R(x) - \beta(w(x) - 1))) \log \pi_{\theta}(x) \right]$
Reverse (URKL)	$-\mathbb{E} \left[\text{SG}(Z_{\text{old}}w(x)(R(x) - \beta \log w(x))) \log \pi_{\theta}(x) \right]$

E.1 RPG-Style Clip: dual-clip truncation of importance ratios

Large importance ratios $w(x) = \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)}$ induce high variance and destabilize off-policy updates. Our **RPG-Style Clip** follows the dual-clip method implemented in Algorithm 1 in the appendix: we clip w into $[1 - \epsilon_1, 1 + \epsilon_2]$ and additionally impose a lower bound for negative advantages.

Table 4: Summary of fully differentiable surrogate losses for **normalized** KL-regularized objectives (counterparts to Table 1). Here $x \sim \pi_{\text{old}}$, $w(x) = \pi_{\theta}(x)/\pi_{\text{old}}(x)$.

Regularization (Normalized)	Surrogate loss (sampling $x \sim \pi_{\text{old}}$)
Forward (KL)	$\mathbb{E}[-w(x)R(x) - \beta \log \pi_{\theta}(x)]$
Reverse (KL)	$\mathbb{E}[w(x)(-R(x) + \beta \log w(x))]$

Let $\hat{A}(x; \theta)$ denote the regularized advantage analogue determined by the chosen objective (e.g., $\hat{A}_{\text{URKL}} = (R-b) - \beta \log w$, $\hat{A}_{\text{RKL}} = (R-b) - \beta(\log w + 1)$). The loss used in our implementation is

$$\mathcal{L}^{\text{RPG-Clip}}(x, \theta) = \begin{cases} \max(-w(x)\hat{A}(x; \theta), -\text{clip}(w(x), 1 - \epsilon_1, 1 + \epsilon_2)\hat{A}(x; \theta)), & \hat{A}(x; \theta) \geq 0, \\ \min(\max(-w(x)\hat{A}(x; \theta), -\text{clip}(w(x), 1 - \epsilon_1, 1 + \epsilon_2)\hat{A}(x; \theta)), -c\hat{A}(x; \theta)), & \hat{A}(x; \theta) < 0, \end{cases}$$

with $\epsilon_1, \epsilon_2 > 0$ and $c > 1$. The choice of $\hat{A}$ for each divergence (URKL/UFKL/RKL/FKL) matches the gradients in Section 2 and is instantiated in Algorithm 1.

F Equivalence of k_3 Estimator and Unnormalized KL Divergence

As mentioned in Section D, the k_3 estimator for KL divergence [Schulman, 2020] is equivalent to the unnormalized KL (UKL) divergence. The k_3 function is defined as $k_3(y) = y - 1 - \log y$.

Forward KL- k_3 and UKL($\pi_{\text{old}} \parallel \pi_{\theta}$): The forward KL- k_3 divergence is

$$\text{KL}_{k_3}(\pi_{\text{old}} \parallel \pi_{\theta}) := \mathbb{E}_{x \sim \pi_{\text{old}}} [k_3(\pi_{\theta}(x)/\pi_{\text{old}}(x))].$$

$$\begin{aligned} \mathbb{E}_{x \sim \pi_{\text{old}}} \left[k_3 \left(\frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} \right) \right] &= \mathbb{E}_{x \sim \pi_{\text{old}}} \left[\frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} - 1 - \log \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} \right] \\ &= \int_x \pi_{\text{old}}(x) \left(\frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} - 1 \right) dx - \int_x \pi_{\text{old}}(x) \log \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} dx \\ &= \int_x (\pi_{\theta}(x) - \pi_{\text{old}}(x)) dx + \int_x \pi_{\text{old}}(x) \log \frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} dx \\ &= \text{UKL}(\pi_{\text{old}} \parallel \pi_{\theta}). \end{aligned}$$

Reverse KL- k_3 and UKL($\pi_{\theta} \parallel \pi_{\text{old}}$): The reverse KL- k_3 divergence is

$$\text{KL}_{k_3}(\pi_{\theta} \parallel \pi_{\text{old}}) := \mathbb{E}_{x \sim \pi_{\theta}} [k_3(\pi_{\text{old}}(x)/\pi_{\theta}(x))].$$

$$\begin{aligned} \mathbb{E}_{x \sim \pi_{\theta}} \left[k_3 \left(\frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} \right) \right] &= \mathbb{E}_{x \sim \pi_{\theta}} \left[\frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} - 1 - \log \frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} \right] \\ &= \int_x \pi_{\theta}(x) \left(\frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} - 1 \right) dx - \int_x \pi_{\theta}(x) \log \frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} dx \\ &= \int_x (\pi_{\text{old}}(x) - \pi_{\theta}(x)) dx + \int_x \pi_{\theta}(x) \log \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} dx \\ &= \text{UKL}(\pi_{\theta} \parallel \pi_{\text{old}}). \end{aligned}$$

G Normalized KL Regularization

For completeness, we collect here the normalized KL formulations that were previously in the main text. Their proofs remain in Appendix M.

G.1 Forward KL Regularization

Consider the objective function with forward KL regularization:

$$J_{\text{FKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\theta}} [R(x)] - \beta \text{KL}(\pi_{\text{old}} \parallel \pi_{\theta}). \quad (\text{G.1})$$

Proposition G.1 (Policy Gradient and Differentiable Loss for Forward KL). The gradient of $J_{\text{FKL}}(\theta)$ with respect to θ is:

$$\nabla_{\theta} J_{\text{FKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} [(w(x)R(x) + \beta) \nabla_{\theta} \log \pi_{\theta}(x)],$$

where $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$. A corresponding surrogate loss is:

$$\mathcal{L}_{\text{FKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} [-w(x)R(x) - \beta \log \pi_\theta(x)],$$

which satisfies $\nabla_\theta \mathcal{L}_{\text{FKL}}(\theta) = -\nabla_\theta J_{\text{FKL}}(\theta)$.

Remark G.2 (Connection to Maximum Likelihood Estimation). If $R(x) = 0$, maximizing $J_{\text{FKL}}(\theta)$ reduces to minimizing $\beta \text{KL}(\pi_{\text{old}} \parallel \pi_\theta)$, i.e., MLE on samples from π_{old} .

G.2 Reverse KL Regularization

Consider the reverse KL objective:

$$J_{\text{RKL}}(\theta) = \mathbb{E}_{x \sim \pi_\theta} [R(x)] - \beta \text{KL}(\pi_\theta \parallel \pi_{\text{old}}). \quad (\text{G.2})$$

Proposition G.3 (Policy Gradient and Differentiable Loss for Reverse KL). The gradient of $J_{\text{RKL}}(\theta)$ is:

$$\nabla_\theta J_{\text{RKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} \left[w(x) \left(R(x) - \beta (\log w(x) + 1) \right) \nabla_\theta \log \pi_\theta(x) \right].$$

A corresponding surrogate loss is:

$$\mathcal{L}_{\text{RKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) (-R(x) + \beta \log w(x))],$$

with $\nabla_\theta \mathcal{L}_{\text{RKL}}(\theta) = -\nabla_\theta J_{\text{RKL}}(\theta)$.

REINFORCE-style RPG with normalized KL regularizations. REINFORCE-style losses for FKL/RKL appear in Appendix H (Table analogues to Table 3).

H REINFORCE-Style Regularized Policy Gradients with Various KL Regularization Forms

H.1 Rationale for REINFORCE-Style Loss Formulation

As noted in Section E of the main text, the derived off-policy policy gradients (Theorems G.1 through D.3) share a structural similarity with the REINFORCE estimator:

$$\nabla_\theta J(\theta) = \mathbb{E}_{x \sim \pi_{\text{sampling}}} [\text{Weight}(x, \theta) \nabla_\theta \log \pi_\theta(x)].$$

This structure suggests an alternative way to implement the gradient update, analogous to the REINFORCE-style approach used in the on-policy setting. Specifically, one could define a surrogate loss of the form:

$$\mathcal{L}_{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \pi_{\text{sampling}}} [\text{SG}(\text{Weight}(x, \theta)) \log \pi_\theta(x)]. \quad (\text{H.1})$$

The rationale is that applying automatic differentiation to this loss should yield:

$$\nabla_\theta \mathcal{L}_{\text{REINFORCE-style}}(\theta) \stackrel{\text{Autodiff}}{=} -\mathbb{E}_{x \sim \pi_{\text{sampling}}} [\text{SG}(\text{Weight}(x, \theta)) \nabla_\theta \log \pi_\theta(x)].$$

When this gradient is used for optimization, the stop-gradient SG is conceptually removed, resulting in an update aligned with $-\nabla_\theta J(\theta)$. This relies on SG preventing gradients from flowing through the θ -dependence within $\text{Weight}(x, \theta)$ (specifically, the dependence via the importance weight $w(x)$). The following subsections detail these REINFORCE-style loss formulations for each KL regularization type.

H.2 REINFORCE-Style RPG with Forward KL Regularization

We can convert Forward KL regularization of RPG to REINFORCE-style using the stop-gradient operator:

Proposition H.1 (REINFORCE-Style Loss for Forward KL). For the forward KL regularized objective function in Eq. (G.1), the corresponding REINFORCE-style surrogate loss function for gradient descent optimization via automatic differentiation is:

$$\mathcal{L}_{\text{FKL}}^{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \pi_{\text{old}}} [\text{SG}(w(x)R(x) + \beta) \log \pi_\theta(x)],$$

where $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$. This loss aims to produce the gradient $-\nabla_\theta J_{\text{FKL}}(\theta)$ via automatic differentiation.

Remark H.2. This REINFORCE-style loss requires SG to prevent backpropagation through $w(x)$ in the weight term. Baselines can be added to $R(x)$ inside SG for variance reduction (see Appendix I). In practice we further apply *RPG-Style Clip* (Section E.1) by replacing w with $\bar{w}$ and, when present, $\log w$ with $\log \bar{w}$ inside $\text{SG}(\cdot)$.

H.3 REINFORCE-Style RPG with Unnormalized Forward KL Regularization

Similarly, we can also transform the Unnormalized Forward KL Regularization of RPG into REINFORCE-style as follows:

Proposition H.3 (REINFORCE-Style Loss for Unnormalized Forward KL). For the objective $J_{\text{UFKL}}(\theta) = \mathbb{E}_{\pi_\theta}[R(x)] - \beta \text{UKL}(\pi_{\text{old}} \parallel \pi_\theta)$, whose gradient (sampling from $\tilde{\pi}_{\text{old}}$) is $\nabla_\theta J_{\text{UFKL}}(\theta) = \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}}[Z_{\text{old}}(w(x)R(x) - \beta(w(x) - 1))\nabla_\theta \log \pi_\theta(x)]$ (Proposition 2.2), a corresponding REINFORCE-style surrogate loss is:

$$\mathcal{L}_{\text{UFKL}}^{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}}[\text{SG}(Z_{\text{old}}(w(x)R(x) - \beta(w(x) - 1)))\log \pi_\theta(x)],$$

where $\tilde{\pi}_{\text{old}} = \pi_{\text{old}}/Z_{\text{old}}$ and $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$ (using unnormalized π_{old}). This loss aims to produce the gradient $-\nabla_\theta J_{\text{UFKL}}(\theta)$ via automatic differentiation.

H.4 REINFORCE-Style RPG with Reverse KL Regularization

Proposition H.4 (REINFORCE-Style Loss for Reverse KL). For the objective $J_{\text{RKL}}(\theta) = \mathbb{E}_{\pi_\theta}[R(x)] - \beta \text{KL}(\pi_\theta \parallel \pi_{\text{old}})$, whose gradient is $\nabla_\theta J_{\text{RKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}}[w(x)(R(x) - \beta(\log w(x) + 1))\nabla_\theta \log \pi_\theta(x)]$ (Proposition G.3), a corresponding REINFORCE-style surrogate loss is:

$$\mathcal{L}_{\text{RKL}}^{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \pi_{\text{old}}}[\text{SG}(w(x)(R(x) - \beta \log w(x) - \beta))\log \pi_\theta(x)], \quad (\text{H.2})$$

where $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$. This loss aims to produce the gradient $-\nabla_\theta J_{\text{RKL}}(\theta)$ via automatic differentiation.

H.5 REINFORCE-Style RPG with Unnormalized Reverse KL Regularization

Proposition H.5 (REINFORCE-Style Loss for Unnormalized Reverse KL). For the objective $J_{\text{URKL}}(\theta) = \mathbb{E}_{\pi_\theta}[R(x)] - \beta \text{UKL}(\pi_\theta \parallel \pi_{\text{old}})$, whose gradient (sampling from $\tilde{\pi}_{\text{old}}$) is $\nabla_\theta J_{\text{URKL}}(\theta) = \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}}[Z_{\text{old}}w(x)(R(x) - \beta \log w(x))\nabla_\theta \log \pi_\theta(x)]$ (Proposition D.3), a corresponding REINFORCE-style surrogate loss is:

$$\mathcal{L}_{\text{URKL}}^{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}}[\text{SG}(Z_{\text{old}}w(x)(R(x) - \beta \log w(x)))\log \pi_\theta(x)],$$

where $\tilde{\pi}_{\text{old}} = \pi_{\text{old}}/Z_{\text{old}}$ and $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$ (using unnormalized π_{old}). This loss aims to produce the gradient $-\nabla_\theta J_{\text{URKL}}(\theta)$ via automatic differentiation.

I More on Algorithmic Details

I.1 Stabilization Techniques for Regularized Policy Gradients

Practical implementations of off-policy policy gradient methods often require stabilization techniques to manage variance or prevent destructively large policy updates. Common techniques include:

- **Dual-Clip Objective:** This method adapts the clipping mechanism from PPO, with a modification for negative advantages proposed by Ye et al. [2020], to stabilize updates [Schulman et al., 2017]. The Dual Clip objective aims to maximize $J^{\text{DualClip}} = \mathbb{E}_{x \sim \pi_{\text{old}}}[L^{\text{DualClip}}(x, \theta)]$, where $\hat{A}(x)$ is an estimate of the advantage analogue (e.g., $R(x) - b$ or the full term derived from the regularized gradient), $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$ is the importance ratio, and $L^{\text{DualClip}}(x, \theta)$ is defined as:

- If $\hat{A}(x) \geq 0$: $L^{\text{DualClip}}(x, \theta) = \min(w(x)\hat{A}(x), \text{clip}(w(x), 1 - \epsilon_1, 1 + \epsilon_2)\hat{A}(x))$.
- If $\hat{A}(x) < 0$: $L^{\text{DualClip}}(x, \theta) = \max(\min(w(x)\hat{A}(x), \text{clip}(w(x), 1 - \epsilon_1, 1 + \epsilon_2)\hat{A}(x)), c\hat{A}(x))$.

where $\epsilon_1, \epsilon_2 > 0$ are clipping parameters and $c > 1$ provides a lower bound for negative advantages. To use this with gradient descent (which minimizes a loss $\mathcal{L}$), we minimize the negative of the Dual Clip objective term. Using $-\min(a, b) = \max(-a, -b)$ and $-\max(a, b) = \min(-a, -b)$, the corresponding loss term for a single sample x is:

- If $\hat{A}(x) \geq 0$: $\mathcal{L}^{\text{DualClip}}(x, \theta) = \max(-w(x)\hat{A}(x), -\text{clip}(w(x), 1 - \epsilon_1, 1 + \epsilon_2)\hat{A}(x))$.
- If $\hat{A}(x) < 0$: Let $L_{\text{clip}} = \max(-w(x)\hat{A}(x), -\text{clip}(w(x), 1 - \epsilon_1, 1 + \epsilon_2)\hat{A}(x))$. Then, $\mathcal{L}^{\text{DualClip}}(x, \theta) = \min(L_{\text{clip}}, -c\hat{A}(x))$.

Here, $\hat{A}(x)$ should represent the advantage or an analogous term derived from the gradient of the original (non-negated) regularized objective (e.g., Proposition G.3). The overall loss is $\mathcal{L}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} [\mathcal{L}^{\text{DualClip}}(x, \theta)]$. This loss function is differentiable with respect to θ (which appears in $w(x)$ and potentially $\hat{A}(x)$ if it includes terms like $\log w(x)$).

This loss formulation ensures that updates are conservative. For positive advantages, it acts like standard PPO-Clip. For negative advantages, it prevents the objective from becoming arbitrarily large (loss becoming arbitrarily small) by introducing the lower bound $c\hat{A}(x)$ on the objective (upper bound $-\hat{A}(x)$ on the loss).

- **Baseline Subtraction:** Used to define the advantage $\hat{A}(x) = R(x) - b(x)$, reducing the variance of the gradient estimates. The baseline $b(x)$ should ideally not depend strongly on θ . A common choice is a value function estimate $V(x)$ or simply the batch average reward $b = \frac{1}{N} \sum R(x_i)$. The definition of $\hat{A}(x)$ might also incorporate regularization terms depending on the base objective chosen (see RKL example below).

For instance, applying Dual Clip to stabilize the reverse KL objective (Proposition G.3). The gradient involves the term $w(x) \underbrace{((R(x) - b) - \beta(\log w(x) + 1))}_{\text{Analogue to } \hat{A}_{\text{RKL}}(x, w; b)} \nabla \log \pi_{\theta}$. Using this $\hat{A}_{\text{RKL}}$ in the Dual Clip

loss structure $\mathcal{L}_{\text{RKL}}^{\text{DualClip}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} [\mathcal{L}_{\text{RKL}}^{\text{DualClip}}(x, \theta)]$ where:

- If $\hat{A}_{\text{RKL}}(x, w; b) \geq 0$:

$$\mathcal{L}_{\text{RKL}}^{\text{DualClip}}(x, \theta) = \max \left(-w(x)\hat{A}_{\text{RKL}}, -\text{clip}(w(x), 1 - \epsilon_1, 1 + \epsilon_2)\hat{A}_{\text{RKL}} \right).$$

- If $\hat{A}_{\text{RKL}}(x, w; b) < 0$: Let $L_{\text{clip}} = \max \left(-w(x)\hat{A}_{\text{RKL}}, -\text{clip}(w(x), 1 - \epsilon_1, 1 + \epsilon_2)\hat{A}_{\text{RKL}} \right)$.

$$\mathcal{L}_{\text{RKL}}^{\text{DualClip}}(x, \theta) = \min \left(L_{\text{clip}}, -c\hat{A}_{\text{RKL}} \right),$$

where $\hat{A}_{\text{RKL}}(x, w; b) = (R(x) - b) - \beta(\log w(x) + 1)$. Simpler approximations might use $\hat{A}(x) = R(x) - b$.

Using PPO-style clipping alters the optimization objective compared to the original KL-regularized objectives, trading strict adherence for enhanced stability. The choice of base objective structure, definition of $\hat{A}$, and stabilization techniques depends on the specific application.

I.2 Stabilization Techniques for REINFORCE-Style Regularized Policy Gradients

While the REINFORCE-style losses derived in this section (Table 3) provide theoretically grounded gradient estimators for the regularized objectives, practical implementations often benefit significantly from stabilization techniques common in policy gradient methods. These techniques aim to reduce variance and control the magnitude of policy updates, which is especially crucial in the off-policy setting where importance weights $w(x)$ and can exacerbate instability.

- **Baseline Subtraction and Regularized Advantage Definition:** This is a standard variance reduction technique. Critically, when combining with stabilization like PPO clipping in this REINFORCE-style context, the term playing the role of the advantage ($\hat{A}_t$) that gets clipped should ideally incorporate not just the baselined reward but also the regularization terms derived from the objective's gradient.

Recall the REINFORCE-style gradient structure $\nabla_{\theta} J(\theta) = \mathbb{E}_{x \sim \pi_{\text{sampling}}} [\text{Weight}(x, \theta) \nabla_{\theta} \log \pi_{\theta}(x)]$.

The PPO objective involves terms like $w_t \hat{A}_t$. To align these, we define the regularized advantage $\hat{A}_t$ such that $w_t \hat{A}_t$ approximates the key part of $\text{Weight}(x, \theta)$. For example:

- For RKL (Proposition H.4), $\text{Weight}_{\text{RKL}} = w(x)(R(x) - \beta(\log w(x) + 1))$. We define the regularized advantage as $\hat{A}_t^{\text{RKL}} = (R(x) - b(x)) - \beta(\log w(x) + 1)$.
- For URKL (Proposition H.5), $\text{Weight}_{\text{URKL}} = Z_{\text{old}} w(x)(R(x) - \beta \log w(x))$. Ignoring Z_{old} , we define $\hat{A}_t^{\text{URKL}} = (R(x) - b(x)) - \beta \log w(x)$.

Algorithm 1 RPG with Dual-Clip Stabilization

Require: Reference policy π_{old} , Reward function $R(x)$, Initial policy parameters θ_0
Require: Base objective structure J_{chosen} (implies regularization type), Regularization strength $\beta \geq 0$
Require: Learning rate $\alpha > 0$, Batch size $N > 0$, Number of epochs $K \geq 1$ per iteration
Require: Dual Clip parameters: $\epsilon_1 > 0, \epsilon_2 > 0, c > 1$
Require: Baseline method (e.g., batch/group average, value function V_ϕ)

- 1: Initialize policy parameters $\theta \leftarrow \theta_0$
- 2: Initialize value function parameters ϕ (if baseline uses V_ϕ)
- 3: **for** each training iteration **do**
- 4: Sample batch $\mathcal{D} = \{x_i\}_{i=1}^N \sim \pi_{\text{old}}$ ▷ Collect data using old policy
- 5: Compute R_i for $i = 1..N$
- 6: Compute baselines b_i for $i = 1..N$ (e.g., $b_i = \frac{1}{N} \sum_j R_j$ or $b_i = V_\phi(x_i)$)
- 7: **for** $k = 1$ to K **do** ▷ Multiple optimization epochs on the same batch
- 8: Initialize batch loss $\mathcal{L}_{\text{batch}} = 0$
- 9: **for** $i = 1$ to N **do**
- 10: $w_i = \frac{\pi_\theta(x_i)}{\pi_{\text{old}}(x_i)}, \log w_i = \log \pi_\theta(x_i) - \log \pi_{\text{old}}(x_i)$ ▷ Compute importance weight
- 11: Define Advantage analogue $\hat{A}_i$ based on $J_{\text{chosen}}, R_i, b_i, w_i, \beta$.
 ▷ Ex: For RKL, $\hat{A}_i = (R_i - b_i) - \beta(\log w_i + 1)$. Note: $\hat{A}_i$ depends on current θ via w_i
- 12: **if** Dual Clip enabled **then**
- 13: $\text{loss_term1}_i = -w_i \times \hat{A}_i$ ▷ Negative of unclipped term, gradient flows through w_i
- 14: $w_{i,\text{clipped}} = \text{clip}(w_i, 1 - \epsilon_1, 1 + \epsilon_2)$
- 15: $\text{loss_term2}_i = -w_{i,\text{clipped}} \times \hat{A}_i$ ▷ Negative of clipped term
- 16: $L_{\text{clip}}(i) = \max(\text{loss_term1}_i, \text{loss_term2}_i)$
- 17: **if** $\hat{A}_i \geq 0$ **then**
- 18: $\mathcal{L}_{\text{term}}(i) = L_{\text{clip}}(i)$
- 19: **else** ▷ $\hat{A}_i < 0$
- 20: $\text{loss_lower_bound}_i = -c \times \hat{A}_i$ ▷ Lower bound term
- 21: $\mathcal{L}_{\text{term}}(i) = \min(L_{\text{clip}}(i), \text{loss_lower_bound}_i)$
- 22: **end if**
- 23: **else**
- 24: ▷ Define base loss term (unclipped) based on chosen objective's negative gradient structure
- 25: ▷ Ex: For RKL loss (no clip): $\mathcal{L}_{\text{term}}(i) = w_i(-(R_i - b_i) + \beta \log w_i)$
- 26: $\mathcal{L}_{\text{term}}(i) = -w_i \times \hat{A}_i$
- 27: **end if**
- 28: $\mathcal{L}_{\text{batch}} = \mathcal{L}_{\text{batch}} + \mathcal{L}_{\text{term}}(i)$
- 29: **end for**
- 30: $\hat{\mathcal{L}}(\theta) = \frac{1}{N} \mathcal{L}_{\text{batch}}$ ▷ Compute final batch loss for minimization
- 31: $g \leftarrow \nabla_\theta \hat{\mathcal{L}}(\theta)$ ▷ Compute gradient (flows through w_i and $\hat{A}_i$)
- 32: $\theta \leftarrow \text{OptimizerUpdate}(\theta, g, \alpha)$ ▷ Update policy parameters
- 33: **if** using a learned baseline V_ϕ **then**
- 34: Update value function parameters ϕ (e.g., by minimizing $\mathbb{E}[(V_\phi(x_i) - R_i)^2]$ over the batch)
- 35: **end if**
- 36: **end for**
- 37: **end for**
- 38: **end for**
- 39: **return** Optimized policy parameters θ

- For FKL or UFKL, the structure might not cleanly separate into $w(x) \times (\dots)$. In such cases, a common simplification is to use $\hat{A}_t = R(x) - b(x)$ and accept that the clipping primarily stabilizes the reward term's contribution.

This calculated $\hat{A}_t$ (incorporating reward, baseline, and KL terms) is then treated as constant using the stop-gradient operator, $\text{SG}(\hat{A}_t)$, when plugged into the clipping loss function.

- **RPG-Style Objective Clipping (Dual-Clip Variant):** PPO [Schulman et al., 2017] introduces objective clipping to limit the impact of large importance ratios $w(x)$. The Dual-Clip variant [Ye et al., 2020] refines this, particularly for negative advantages, using a lower bound parameter $c > 1$. When applied in the REINFORCE-style setting, the PPO Dual-Clip objective aims to maximize (simplified notation, expectation over $t \sim \pi_{\text{old}}$):

$$J^{\text{DualClip}}(\theta) = \mathbb{E}_t[L_t^{\text{DualClip}}(\theta)]$$

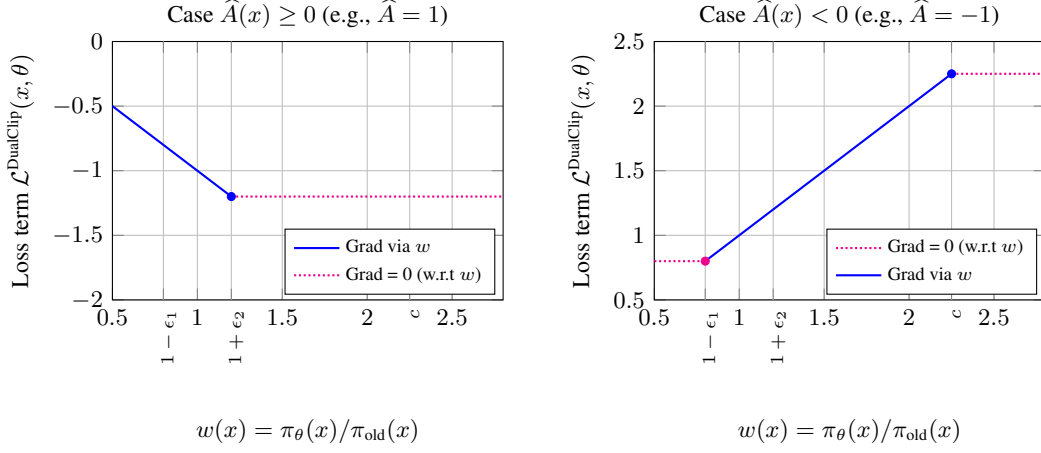


Figure 4: Visualization of the Dual-Clip loss term $\mathcal{L}^{\text{DualClip}}(x, \theta)$ vs. importance weight $w(x)$, as described in Section I.1 and Algorithm 1. This formulation is typically implemented as fully differentiable w.r.t θ (via $w(x)$ and potentially $\hat{A}(x)$ if $\hat{A}$ depends on θ , e.g., via $\log w(x)$), unlike REINFORCE-style implementations that use $\text{SG}(\hat{A})$ or $\text{SG}(\ell_i)$ within the loss. For visualization, $\hat{A}(x)$ is treated as constant ($\hat{A} = 1$ left, $\hat{A} = -1$ right) to isolate the effect of w . **Solid blue:** Loss depends linearly on w , gradient $\nabla_\theta \mathcal{L}$ flows via $w(x)$. **Dotted magenta:** Loss is constant w.r.t w , gradient $\nabla_\theta \mathcal{L}$ does not flow via $w(x)$ in this segment (though it might flow via $\hat{A}$ if $\hat{A}$ depends on θ). Left: Case $\hat{A} < 0$. Right: Case $\hat{A} \geq 0$.

where $\hat{A}_t$ is the regularized advantage defined above (incorporating R_t, b_t , and KL terms), $w_t(\theta) = \frac{\pi_\theta(a_t | s_t)}{\pi_{\text{old}}(a_t | s_t)}$, and $\mathcal{L}_t^{\text{DualClip}}(\theta)$ is defined based on the sign of $\text{SG}(\hat{A}_t)$:

- If $\text{SG}(\hat{A}_t) \geq 0$: $\mathcal{L}_t^{\text{DualClip}}(\theta) = \min(w_t(\theta) \text{SG}(\hat{A}_t), \text{clip}(w_t(\theta), 1 - \epsilon_1, 1 + \epsilon_2) \text{SG}(\hat{A}_t))$
- If $\text{SG}(\hat{A}_t) < 0$: $\mathcal{L}_t^{\text{DualClip}}(\theta) = \max(\min(w_t(\theta) \text{SG}(\hat{A}_t), \text{clip}(w_t(\theta), 1 - \epsilon_1, 1 + \epsilon_2) \text{SG}(\hat{A}_t)), c \text{SG}(\hat{A}_t))$

Here, ϵ_1, ϵ_2 are clipping hyperparameters, and c is the lower bound factor. Note that θ influences this objective only through $w_t(\theta)$, as $\hat{A}_t$ is detached via SG.

To implement this using gradient descent (minimizing a loss), we minimize the negative of the PPO Dual-Clip objective. The loss function becomes $\mathcal{L}^{\text{DualClip}}(\theta) = \mathbb{E}_t[\mathcal{L}_t^{\text{DualClip}}(\theta)]$, where $\mathcal{L}_t^{\text{DualClip}}(\theta) = -L_t^{\text{DualClip}}(\theta)$. Explicitly:

- If $\text{SG}(\hat{A}_t) \geq 0$: $\mathcal{L}_t^{\text{DualClip}}(\theta) = \max(-w_t(\theta) \text{SG}(\hat{A}_t), -\text{clip}(w_t(\theta), 1 - \epsilon_1, 1 + \epsilon_2) \text{SG}(\hat{A}_t))$.
- If $\text{SG}(\hat{A}_t) < 0$: Let $L_{\text{clip}} = \max(-w_t(\theta) \text{SG}(\hat{A}_t), -\text{clip}(w_t(\theta), 1 - \epsilon_1, 1 + \epsilon_2) \text{SG}(\hat{A}_t))$. Then, $\mathcal{L}_t^{\text{DualClip}}(\theta) = \min(L_{\text{clip}}, -c \text{SG}(\hat{A}_t))$.

This PPO Dual-Clip loss function $\mathcal{L}^{\text{DualClip}}(\theta)$ replaces the simpler REINFORCE-style losses derived earlier (like $\mathcal{L}_{\text{RKL}}^{\text{REINFORCE-style}}$ in Eq. (H.2)). The gradient $\nabla_\theta \mathcal{L}^{\text{DualClip}}(\theta)$ is computed via automatic differentiation, where the gradient flows through $w_t(\theta)$ but is stopped at $\hat{A}_t$. This approach uses the PPO objective structure with the appropriately defined regularized advantage for stabilization in an off-policy REINFORCE-style update. Algorithm 2 details this implementation.

J Detailed Experimental Setup

Base Models and Datasets. We conduct experiments using the Qwen3-4B large language models [Team, 2025]. For training, we utilize the DAPO-Math-17k dataset [Yu et al., 2025], filtered to include only English samples, resulting in a 13.9k sample training set. Model performance is evaluated on several mathematical reasoning benchmarks: AIME2024 [MAA, 2024a,b] and AIME2025 [MAA, 2025a,b].

Implementation and Framework. Experiments are implemented using the `verl` framework [Sheng et al., 2025] with the `vLLM` engine [Kwon et al., 2023] for efficient LLM serving and inference. For practical implementation of our RPG methods, we emphasize that the probabilities (or log-probabilities) from the last iteration’s model (π_{old}) for the sampled data can be pre-computed and stored. This allows the KL regularization terms to be calculated without needing to keep π_{old} in GPU memory during the training step of the current policy π_{θ} . Consequently, only *one* model (π_{θ}) needs to be actively managed in GPU memory for training, which is faster and more memory-efficient compared to approaches like GRPO that typically require access to at least two models (the current policy and a reference/sampling policy) during optimization.

Iterative reference updates. To further stabilize optimization, we adopt an *iterative reference-update* scheme: we periodically set $\pi_{\text{old}} \leftarrow \pi_{\theta}$ (every K optimizer steps, or when a moving average of token-level KL exceeds a target κ). This realizes a practical KL trust region while avoiding over-regularization toward the initial checkpoint. Further implementation details and hyperparameters (learning rate, β , clipping) are provided in Appendix J.

Stabilization and Advanced RL Techniques. Our RPG implementations (both fully differentiable and REINFORCE-style) incorporate stabilization techniques like baseline subtraction and PPO-style objective clipping (specifically, Dual-Clip [Ye et al., 2020, Schulman et al., 2017]), crucial for robust off-policy learning. Detailed algorithmic descriptions are provided in Appendix I (see Algorithm 1 for RPG with Dual-Clip and Algorithm 2 for the REINFORCE-style equivalent, along with Figures 2 and 4 for visualization). For PPO-style clipping, we set $(\epsilon_1, \epsilon_2) = (0.2, 0.28)$ for RPG, RPG-REINFORCE and DAPO. For GRPO, we use $(\epsilon_1, \epsilon_2) = (0.2, 0.2)$. Furthermore, to enhance training efficiency and data quality, we adopted techniques introduced by DAPO [Yu et al., 2025], including a dynamic sampling strategy with a group filtering mechanism (which oversamples challenging prompts and filters out those with near-perfect or near-zero accuracy based on initial rollouts) and an overlong punishment component in the reward shaping to discourage excessively verbose outputs. In addition, we enable *RPG-Style Clip* (Section E.1) for the REINFORCE-style estimators, which we found to be the most sample-efficient and stable variant for RL training at larger scales.

More Results and Discussion.

Similarly, Table 5 shows the experiment results with 2k context length. It can be observed that RPG and RPG-REINFORCE variants demonstrate robust performance, often competitive with or exceeding baselines. For example, RPG-REINFORCE-UFKL achieves the top “Best” scores for AIME24 (0.3625) and AIME25 (0.3083), and the top “Last” score of AIME25 (0.2927), while RPG-UFKL attain the top “Last” score of AIME24 (0.3427) and the second highest “Last” score of AIME25 (0.2833). Their training curves in Figure 5 generally indicate good stability and effective learning. The consistently high performance across various RPG formulations underscores the utility of the systematically derived KL-regularized objectives explored in this work.

Moreover, these algorithms generally exhibit stable training progressions regarding reward (critic score) and policy entropy, as shown in subfigures (c) and (d) in Figures 3 and 5, compared to some baselines like GRPO, which can show more volatility. This stability likely contributes to their robust benchmark performances (subfigures a-b). The response lengths (subfigure e) for RPG methods also appear well-controlled. These observations align with the strong final scores reported in Tables 2 and 5 for these variants.

Hyperparameters. Unless otherwise specified, all experiments use AdamW optimizer [Loshchilov and Hutter, 2019] with a learning rate of 1×10^{-6} , a weight decay of 0.1, and gradient clipping at 1.0. Training proceeds for 400 steps, including an initial 10 warm-up steps, after which a constant learning rate is maintained. The global training batch size is 512. For each sample in the batch, we roll out 16 responses using a temperature of 1.0. The per-GPU mini-batch size is 32, and experiments are conducted on 8 NVIDIA H100 GPUs. The maximum training and rollout length is set to 4,096 tokens for 2K context length and 8,192 tokens for 4K context length, with dynamic batching enabled. The KL regularization coefficient β is set to 1×10^{-4} .

Specific Clipping Parameters and Adopted Techniques. As mentioned in Section 3, we set $(\epsilon_1, \epsilon_2) = (0.2, 0.28)$ for RPG, RPG-REINFORCE and DAPO. For GRPO, we use $(\epsilon_1, \epsilon_2) = (0.1, 0.1)$.

Algorithm 2 REINFORCE-Style RPG with Dual-Clip Stabilization

Require: Reference policy π_{old} , Reward function $R(x)$, Initial policy parameters θ_0
Require: KL Component function $\text{Compute_KL_Component}(x, \theta, \pi_{\text{old}})$, KL Component Coefficient β
Require: Learning rate $\alpha > 0$, Batch size $N > 0$, Number of epochs $K \geq 1$ per iteration
Require: Dual Clip parameters: $\epsilon_1 > 0$ (low), $\epsilon_2 > 0$ (high), $c > 1$
Require: Baseline method (e.g., batch average, value function V_ϕ)

- 1: Initialize policy parameters $\theta \leftarrow \theta_0$
- 2: Initialize value function parameters ϕ (if baseline uses V_ϕ)
- 3: **for** each training iteration **do**
- 4: Sample batch $\mathcal{D} = \{x_i\}_{i=1}^N \sim \pi_{\text{old}}$
- 5: Compute rewards R_i for $i = 1..N$
- 6: Compute baselines b_i for $i = 1..N$ (e.g., $b_i = \frac{1}{N} \sum_j R_j$ or $b_i = V_\phi(x_i)$)
- 7: **for** $k = 1$ to K **do** ▷ Multiple optimization epochs on the same batch
- 8: Initialize batch loss $\mathcal{L}_{\text{batch}} = 0$
- 9: **for** $i = 1$ to N **do**
- 10: $w_i = \frac{\pi_\theta(x_i)}{\pi_{\text{old}}(x_i)}$ ▷ Importance weight
- 11: $\ell_i = -\log \pi_\theta(x_i)$ ▷ Negative log probability
- 12: $A_{R,i} = R_i - b_i$ ▷ Baseline-subtracted reward
- 13: $C_{\text{KL},i} = \beta \cdot \text{Compute_KL_Component}(x_i, \theta, \pi_{\text{old}}(x_i))$ ▷ KL component
- 14: $A'_i = A_{R,i} + \text{SG}(C_{\text{KL},i}) / \text{SG}(w_i)$ ▷ Effective advantage
- 15: $\psi_i = A'_i \times \ell_i$ ▷ Branching term
- 16: **if** $\psi_i \geq 0$ **then**
- 17: $w_{\text{high}} = 1 + \epsilon_2$
- 18: **if** $w_i < w_{\text{high}}$ **then**
- 19: $\mathcal{L}_i = \psi_i \times \text{SG}(w_i)$ ▷ Grad exists
- 20: **else** ▷ $w_i \geq w_{\text{high}}$
- 21: $A'_{\text{high}} = A_{R,i} + \text{SG}(C_{\text{KL},i}) / \text{SG}(w_{\text{high}})$
- 22: $\psi_{\text{high}} = A'_{\text{high}} \times \text{SG}(\ell_i)$
- 23: $\mathcal{L}_i = \psi_{\text{high}} \times \text{SG}(w_{\text{high}})$
- 24: **end if**
- 25: **else** ▷ $\psi_i \leq 0$
- 26: $w_{\text{low}} = 1 - \epsilon_1$
- 27: **if** $w_i \leq w_{\text{low}}$ **then**
- 28: $A'_{\text{low}} = A_{R,i} + \text{SG}(C_{\text{KL},i}) / \text{SG}(w_{\text{low}})$
- 29: $\psi_{\text{low}} = A'_{\text{low}} \times \text{SG}(\ell_i)$
- 30: $\mathcal{L}_i = \psi_{\text{low}} \times \text{SG}(w_{\text{low}})$
- 31: **else if** $w_i < c$ **then**
- 32: $\mathcal{L}_i = \psi_i \times \text{SG}(w_i)$ ▷ Grad exists
- 33: **else** ▷ $w_i \geq c$
- 34: $\mathcal{L}_i = A_{R,i} \times \text{SG}(\ell_i) \times c + \text{SG}(C_{\text{KL},i}) \times \text{SG}(\ell_i)$
- 35: **end if**
- 36: **end if**
- 37: $\mathcal{L}_{\text{batch}} = \mathcal{L}_{\text{batch}} + \mathcal{L}_i$
- 38: **end for**
- 39: $\mathcal{L}(\theta) = \frac{1}{N} \mathcal{L}_{\text{batch}}$ ▷ Compute average batch loss
- 40: $g \leftarrow \nabla_\theta \mathcal{L}(\theta)$ ▷ Compute gradient
- 41: $\theta \leftarrow \text{OptimizerUpdate}(\theta, g, \alpha)$ ▷ Update policy parameters
- 42: **if** using a learned baseline V_ϕ **then**
- 43: Update value function parameters ϕ
- 44: **end if**
- 45: **end for**
- 46: **end for**
- 47: **return** Optimized policy parameters θ

K More Experiment Results

K.1 The performance with 2k context length

We just display the last and best scores on AIME24 and AIME25 benchmarks in Table 5 for the experiments with 2K context length. The results also demonstrate the superiority of our algorithms over baselines, including GRPO and DAPO.

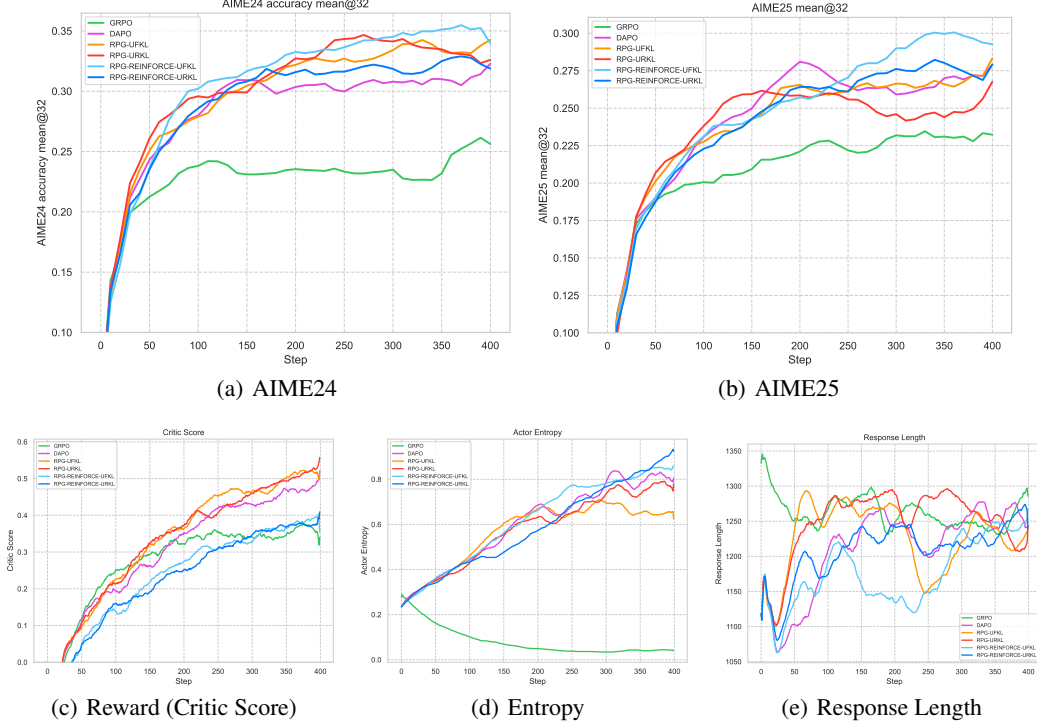


Figure 5: Training dynamics and benchmark performance for fully differentiable Regularized Policy Gradient (RPG) and REINFORCE-Style RPG (RPG-REINFORCE) compared to baselines (GRPO and DAPO) with 2k context length.

Table 5: Combined performance metrics with 2K context length on the AIME24, and AIME25 mathematical reasoning benchmarks, showing “Last” and “Best” scores. The “Last” score is from the 400th training step, assuming the training process remained stable to that point. The highest score in each column is **bolded**, and the second highest is underlined. RPG and RPG-REINFORCE methods are highlighted with light cyan and light green backgrounds, respectively.

Method	AIME24		AIME25	
	Last	Best	Last	Best
GRPO	0.2563	0.2708	0.2323	0.2479
DAPO	0.3229	0.3281	0.2792	0.2844
RPG-UFKL	0.3427	0.3479	<u>0.2833</u>	0.2833
RPG-URKL	0.3260	<u>0.3594</u>	<u>0.2677</u>	0.2677
RPG-REINFORCE-UFKL	<u>0.3396</u>	0.3625	0.2927	0.3083
RPG-REINFORCE-URKL	0.3188	0.3417	0.2792	<u>0.2938</u>

K.2 Ablation Study

To further investigate our algorithms, we implement an ablation study on the clip ratio and the effect of the KL regularization coefficient.

K.2.1 Ablation on clip ratio

We first implement experiments with different clip ratios on REINFORCE-style RPG algorithms. We choose (0.1, 0.1) and (0.2, 0.28) for (ϵ_1, ϵ_2) since they are 2 typical choices of clip ratios [Schulman et al., 2017, Yu et al., 2025], and the performance curves as well as key training dynamics are displayed in Figure 6. It can be observed that although the critic score and response length are similar for different settings, the actor entropy shows a huge difference in trend, demonstrating that an

adequately higher and clip-higher strategy proposed by DAPO may greatly contribute to the increase of performance by decreasing the actor entropy.

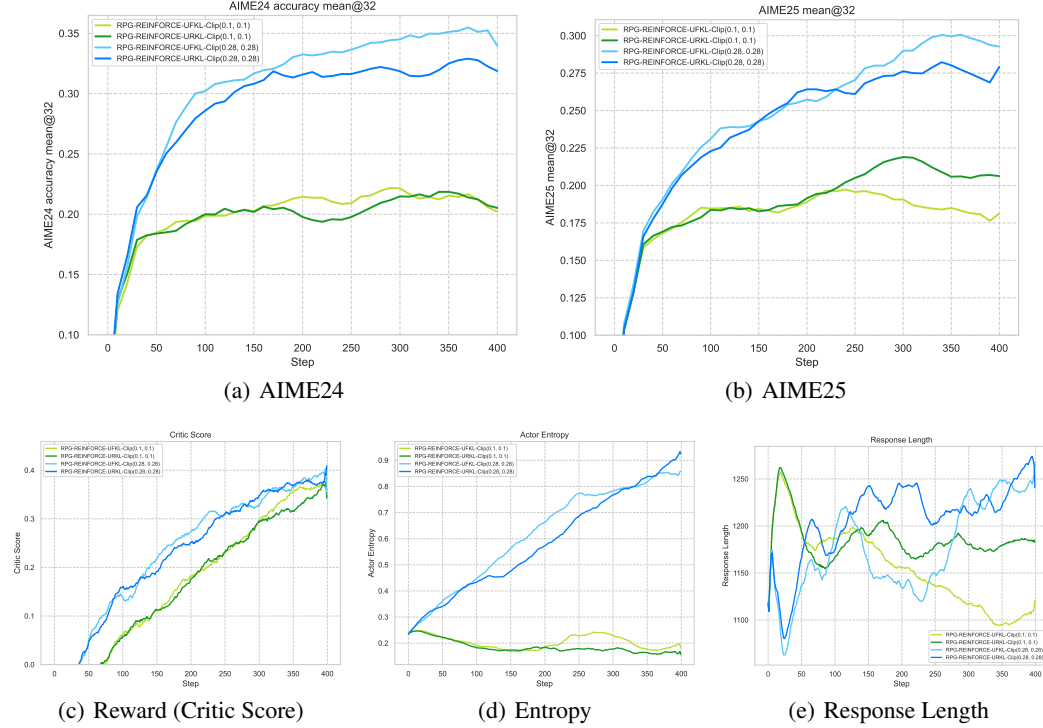


Figure 6: Performance of REINFORCE-Style Regularized Policy Gradient (RPG-REINFORCE) methods with different clip ratios with 2k context length. Plots display accuracy on mathematical reasoning benchmarks (AIME24, AIME25) and key training dynamics (reward, policy entropy, response length).

K.2.2 Ablation on KL regularization coefficient

We also implement ablation studies on the effect of the KL regularization coefficient. We implement experiments with REINFORCE-style RPG-UFKL (RPG-REINFORCE-UFKL) with $\beta = 1 \times 10^{-3}$ and 1×10^{-4} , and the results are shown in Figure 7. Figures 7(a) and 7(b) show that the coefficient 1×10^{-4} performs better than 1×10^{-3} , and the trend in response length conforms to the performance, indicating that longer response length may help with the improvement in performance.

We also dig into the effect of the iteratively updated reference model. We implement another experiment with no iteratively updated reference model, and display the performance and dynamics in Figure 7. It can be observed that the performance recovers with longer response length and much lower actor entropy, showing that longer response length can be an important factor and indicator of the performance on benchmarks.

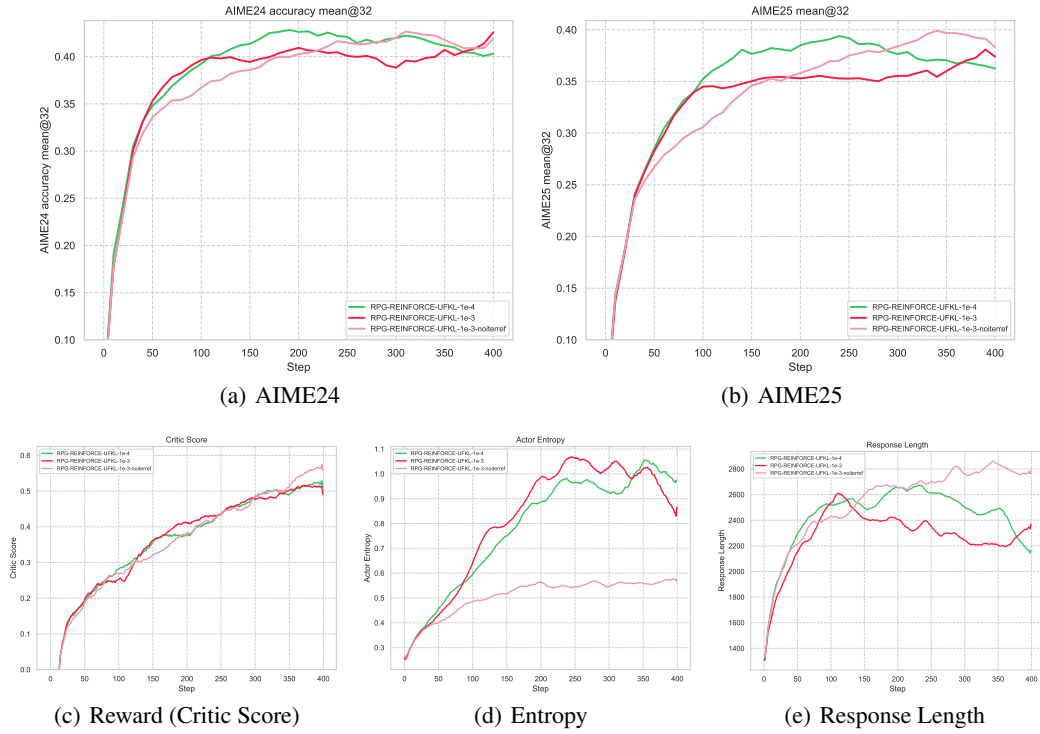


Figure 7: Performance of REINFORCE-Style Regularized Policy Gradient (RPG-REINFORCE) methods with different KL coefficients with 4K context length. Here "noiterref" indicates the model is trained with no iteratively updated reference model. Plots display accuracy on mathematical reasoning benchmarks (AIME24, AIME25) and key training dynamics (reward, policy entropy, response length).

L Proofs of Theorem B.1 (Generalized Policy Gradient Theorem)

Proof. The proof relies on the log-derivative trick, $\nabla_{\theta} \pi_{\theta}(x) = \pi_{\theta}(x) \nabla_{\theta} \log \pi_{\theta}(x)$, and the product rule under the integral sign:

$$\begin{aligned}
\nabla_{\theta} \mathbb{E}_{x \sim \pi_{\theta}} [f(x, \theta)] &= \nabla_{\theta} \int \pi_{\theta}(x) f(x, \theta) dx \\
&= \int \nabla_{\theta} (\pi_{\theta}(x) f(x, \theta)) dx \quad (\text{Swap } \nabla, \int) \\
&= \int ((\nabla_{\theta} \pi_{\theta}(x)) f(x, \theta) + \pi_{\theta}(x) (\nabla_{\theta} f(x, \theta))) dx \\
&= \int (\pi_{\theta}(x) (\nabla_{\theta} \log \pi_{\theta}(x)) f(x, \theta) + \pi_{\theta}(x) (\nabla_{\theta} f(x, \theta))) dx \quad (\text{Log-derivative}) \\
&= \int \pi_{\theta}(x) (f(x, \theta) \nabla_{\theta} \log \pi_{\theta}(x) + \nabla_{\theta} f(x, \theta)) dx \\
&= \mathbb{E}_{x \sim \pi_{\theta}} [f(x, \theta) \nabla_{\theta} \log \pi_{\theta}(x) + \nabla_{\theta} f(x, \theta)].
\end{aligned}$$

□

M Proofs for Regularized Policy Gradients

This section provides detailed proofs for the theorems presented in Section 2, demonstrating that the gradients of the proposed fully differentiable off-policy surrogate losses correspond to the negative gradients of the respective original objectives. The core tool used is the policy gradient theorem: $\nabla_{\theta} \mathbb{E}_{x \sim \pi_{\theta}} [f(x, \theta)] = \mathbb{E}_{x \sim \pi_{\theta}} [f(x, \theta) \nabla_{\theta} \log \pi_{\theta}(x) + \nabla_{\theta} f(x, \theta)]$. We use the notation $w(x) = \pi_{\theta}(x) / \pi_{\text{old}}(x)$ for the importance weight.

M.1 Proof of Proposition G.1 (Policy Gradient and Differentiable Loss for Normalized Forward KL)

Proof. We start by rewriting the objective function $J_{\text{FKL}}(\theta)$ using expectations with respect to the fixed reference policy π_{old} . The first term, the expected reward under π_{θ} , can be rewritten using importance sampling:

$$\mathbb{E}_{x \sim \pi_{\theta}} [R(x)] = \int \pi_{\theta}(x) R(x) dx = \int \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} \pi_{\text{old}}(x) R(x) dx = \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) R(x)].$$

The second term is the forward KL divergence:

$$\begin{aligned}
\text{KL}(\pi_{\text{old}} \parallel \pi_{\theta}) &= \mathbb{E}_{x \sim \pi_{\text{old}}} \left[\log \frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} \right] \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [\log \pi_{\text{old}}(x) - \log \pi_{\theta}(x)] \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [-\log \pi_{\theta}(x)] + \mathbb{E}_{x \sim \pi_{\text{old}}} [\log \pi_{\text{old}}(x)].
\end{aligned}$$

Substituting these into the objective function:

$$\begin{aligned}
J_{\text{FKL}}(\theta) &= \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) R(x)] - \beta (\mathbb{E}_{x \sim \pi_{\text{old}}} [-\log \pi_{\theta}(x)] + \mathbb{E}_{x \sim \pi_{\text{old}}} [\log \pi_{\text{old}}(x)]) \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) R(x) + \beta \log \pi_{\theta}(x)] - \beta \mathbb{E}_{x \sim \pi_{\text{old}}} [\log \pi_{\text{old}}(x)].
\end{aligned}$$

Since $\pi_{\text{old}}(x)$ does not depend on θ , the term $\beta \mathbb{E}_{x \sim \pi_{\text{old}}} [\log \pi_{\text{old}}(x)]$ is a constant with respect to θ . Now we compute the gradient $\nabla_{\theta} J_{\text{FKL}}(\theta)$. Assuming we can swap gradient and expectation (standard assumption in policy gradient methods):

$$\begin{aligned}
\nabla_{\theta} J_{\text{FKL}}(\theta) &= \nabla_{\theta} \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) R(x) + \beta \log \pi_{\theta}(x)] \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [\nabla_{\theta} (w(x) R(x) + \beta \log \pi_{\theta}(x))] \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [(\nabla_{\theta} w(x)) R(x) + \beta \nabla_{\theta} \log \pi_{\theta}(x)].
\end{aligned}$$

We use the identity for the gradient of the importance weight:

$$\nabla_{\theta} w(x) = \nabla_{\theta} \left(\frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} \right)$$

$$\begin{aligned}
&= \frac{1}{\pi_{\text{old}}(x)} \nabla_{\theta} \pi_{\theta}(x) \\
&= \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} \frac{\nabla_{\theta} \pi_{\theta}(x)}{\pi_{\theta}(x)} \\
&= w(x) \nabla_{\theta} \log \pi_{\theta}(x).
\end{aligned}$$

Substituting this back into the gradient expression:

$$\begin{aligned}
\nabla_{\theta} J_{\text{FKL}}(\theta) &= \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x)(\nabla_{\theta} \log \pi_{\theta}(x))R(x) + \beta \nabla_{\theta} \log \pi_{\theta}(x)] \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [(w(x)R(x) + \beta) \nabla_{\theta} \log \pi_{\theta}(x)].
\end{aligned}$$

This proves the first part of the theorem.

Now, consider the surrogate loss function:

$$\mathcal{L}_{\text{FKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} [-w(x)R(x) - \beta \log \pi_{\theta}(x)].$$

We compute its gradient:

$$\begin{aligned}
\nabla_{\theta} \mathcal{L}_{\text{FKL}}(\theta) &= \nabla_{\theta} \mathbb{E}_{x \sim \pi_{\text{old}}} [-w(x)R(x) - \beta \log \pi_{\theta}(x)] \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [\nabla_{\theta} (-w(x)R(x) - \beta \log \pi_{\theta}(x))] \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [-(\nabla_{\theta} w(x))R(x) - \beta \nabla_{\theta} \log \pi_{\theta}(x)] \\
&= \mathbb{E}_{x \sim \pi_{\text{old}}} [-w(x)(\nabla_{\theta} \log \pi_{\theta}(x))R(x) - \beta \nabla_{\theta} \log \pi_{\theta}(x)] \\
&= -\mathbb{E}_{x \sim \pi_{\text{old}}} [(w(x)R(x) + \beta) \nabla_{\theta} \log \pi_{\theta}(x)].
\end{aligned}$$

Comparing this with the gradient of the objective function, we see that $\nabla_{\theta} \mathcal{L}_{\text{FKL}}(\theta) = -\nabla_{\theta} J_{\text{FKL}}(\theta)$. This confirms that minimizing $\mathcal{L}_{\text{FKL}}(\theta)$ corresponds to maximizing $J_{\text{FKL}}(\theta)$ using gradient-based methods. $\square$

M.2 Proof of Proposition 2.2 (Policy Gradient and Differentiable Loss for Unnormalized Forward KL)

Proof. We start by expressing the components of $J_{\text{UFKL}}(\theta)$ using expectations over the normalized reference distribution $\tilde{\pi}_{\text{old}}(x) = \pi_{\text{old}}(x)/Z_{\text{old}}$. The importance weight is $w(x) = \pi_{\theta}(x)/\pi_{\text{old}}(x)$, which implies $\pi_{\theta}(x) = w(x)\pi_{\text{old}}(x) = w(x)Z_{\text{old}}\tilde{\pi}_{\text{old}}(x)$.

The expected reward term:

$$\begin{aligned}
\mathbb{E}_{x \sim \pi_{\theta}} [R(x)] &= \int \pi_{\theta}(x)R(x)dx = \int w(x)\pi_{\text{old}}(x)R(x)dx \\
&= \int w(x)Z_{\text{old}}\tilde{\pi}_{\text{old}}(x)R(x)dx = Z_{\text{old}}\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x)].
\end{aligned}$$

The unnormalized KL divergence term $\text{UKL}(\pi_{\text{old}} \parallel \pi_{\theta})$ has two parts. Part 1 (Generalized KL):

$$\begin{aligned}
\int \pi_{\text{old}}(x) \log \frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} dx &= \int Z_{\text{old}}\tilde{\pi}_{\text{old}}(x) \log \frac{\pi_{\text{old}}(x)}{\pi_{\theta}(x)} dx \\
&= Z_{\text{old}}\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\log \frac{1}{w(x)} \right] = Z_{\text{old}}\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-\log w(x)].
\end{aligned}$$

Part 2 (Mass Correction):

$$\begin{aligned}
\int (\pi_{\theta}(x) - \pi_{\text{old}}(x))dx &= \int (w(x)\pi_{\text{old}}(x) - \pi_{\text{old}}(x))dx \\
&= \int (w(x) - 1)\pi_{\text{old}}(x)dx = \int (w(x) - 1)Z_{\text{old}}\tilde{\pi}_{\text{old}}(x)dx \\
&= Z_{\text{old}}\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x) - 1] = Z_{\text{old}}\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)] - Z_{\text{old}}.
\end{aligned}$$

Combining these parts for the UKL term:

$$\text{UKL}(\pi_{\text{old}} \parallel \pi_{\theta}) = Z_{\text{old}}\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-\log w(x)] + Z_{\text{old}}\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)] - Z_{\text{old}}.$$

Now, substitute everything into the objective $J_{\text{UFKL}}(\theta)$:

$$\begin{aligned} J_{\text{UFKL}}(\theta) &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x)] - \beta (Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-\log w(x)] + Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)] - Z_{\text{old}}) \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x) + \beta \log w(x) - \beta w(x) + \beta]. \end{aligned}$$

To compute the gradient $\nabla_{\theta} J_{\text{UFKL}}(\theta)$, we differentiate the terms inside the expectation. The constant term βZ_{old} (arising from β inside the expectation) vanishes upon differentiation.

$$\begin{aligned} \nabla_{\theta} J_{\text{UFKL}}(\theta) &= \nabla_{\theta} (Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x) + \beta \log w(x) - \beta w(x)]) \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\nabla_{\theta} (w(x)R(x) + \beta \log w(x) - \beta w(x))]. \end{aligned}$$

We need the gradients of $w(x)$ and $\log w(x)$:

$$\begin{aligned} \nabla_{\theta} w(x) &= w(x) \nabla_{\theta} \log \pi_{\theta}(x) \quad (\text{as derived in Proposition G.1 proof}) \\ \nabla_{\theta} \log w(x) &= \nabla_{\theta} (\log \pi_{\theta}(x) - \log \pi_{\text{old}}(x)) = \nabla_{\theta} \log \pi_{\theta}(x). \end{aligned}$$

Substituting these into the gradient expression:

$$\begin{aligned} \nabla_{\theta} J_{\text{UFKL}}(\theta) &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [(\nabla_{\theta} w(x))R(x) + \beta \nabla_{\theta} \log \pi_{\theta}(x) - \beta (\nabla_{\theta} w(x))] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x) \nabla_{\theta} \log \pi_{\theta}(x) + \beta \nabla_{\theta} \log \pi_{\theta}(x) - \beta w(x) \nabla_{\theta} \log \pi_{\theta}(x)] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [(w(x)R(x) - \beta w(x) + \beta) \nabla_{\theta} \log \pi_{\theta}(x)] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\left(w(x)R(x) - \beta (w(x) - 1) \right) \nabla_{\theta} \log \pi_{\theta}(x) \right]. \end{aligned}$$

This proves the first part of the theorem.

Now, consider the surrogate loss function:

$$\mathcal{L}_{\text{UFKL}}(\theta) = Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-w(x)R(x) + \beta (w(x) - \log w(x) - 1)].$$

We compute its gradient:

$$\begin{aligned} \nabla_{\theta} \mathcal{L}_{\text{UFKL}}(\theta) &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\nabla_{\theta} (-w(x)R(x) + \beta (w(x) - \log w(x) - 1))] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-(\nabla_{\theta} w(x))R(x) + \beta (\nabla_{\theta} w(x) - \nabla_{\theta} \log w(x))] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-w(x)R(x) \nabla_{\theta} \log \pi_{\theta}(x) + \beta (w(x) \nabla_{\theta} \log \pi_{\theta}(x) - \nabla_{\theta} \log \pi_{\theta}(x))] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\left(-w(x)R(x) + \beta (w(x) - 1) \right) \nabla_{\theta} \log \pi_{\theta}(x) \right] \\ &= -Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\left(w(x)R(x) - \beta (w(x) - 1) \right) \nabla_{\theta} \log \pi_{\theta}(x) \right]. \end{aligned}$$

Comparing this with the gradient of the objective function, we find $\nabla_{\theta} \mathcal{L}_{\text{UFKL}}(\theta) = -\nabla_{\theta} J_{\text{UFKL}}(\theta)$. This confirms the surrogate loss function. Note that the constant -1 inside the logarithm term in the loss $\mathcal{L}_{\text{UFKL}}$ corresponds to the constant βZ_{old} in the objective J_{UFKL} and does not affect the gradient. $\square$

M.3 Proof of Proposition G.3 (Policy Gradient and Differentiable Loss for Normalized Reverse KL)

Proof. We rewrite the objective function $J_{\text{RKL}}(\theta)$ using expectations with respect to π_{old} . The expected reward term is $\mathbb{E}_{x \sim \pi_{\theta}} [R(x)] = \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x)R(x)]$, as shown previously. The reverse KL divergence term is:

$$\begin{aligned} \text{KL}(\pi_{\theta} \parallel \pi_{\text{old}}) &= \mathbb{E}_{x \sim \pi_{\theta}} \left[\log \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} \right] \\ &= \mathbb{E}_{x \sim \pi_{\theta}} [\log w(x)] \\ &= \int \pi_{\theta}(x) \log w(x) dx \\ &= \int \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} \pi_{\text{old}}(x) \log w(x) dx \\ &= \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) \log w(x)]. \end{aligned}$$

Substituting these into the objective function:

$$J_{\text{RKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x)R(x)] - \beta \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) \log w(x)] = \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x)R(x) - \beta w(x) \log w(x)].$$

Now we compute the gradient $\nabla_{\theta} J_{\text{RKL}}(\theta)$:

$$\begin{aligned}\nabla_{\theta} J_{\text{RKL}}(\theta) &= \nabla_{\theta} \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x)R(x) - \beta w(x) \log w(x)] \\ &= \mathbb{E}_{x \sim \pi_{\text{old}}} [\nabla_{\theta} (w(x)R(x)) - \beta \nabla_{\theta} (w(x) \log w(x))].\end{aligned}$$

We need the gradient of $w(x) \log w(x)$:

$$\begin{aligned}\nabla_{\theta} (w(x) \log w(x)) &= (\nabla_{\theta} w(x)) \log w(x) + w(x) \nabla_{\theta} (\log w(x)) \\ &= (w(x) \nabla_{\theta} \log \pi_{\theta}(x)) \log w(x) + w(x) (\nabla_{\theta} \log \pi_{\theta}(x)) \\ &= w(x) \nabla_{\theta} \log \pi_{\theta}(x) (\log w(x) + 1).\end{aligned}$$

Substituting this and $\nabla_{\theta} w(x) = w(x) \nabla_{\theta} \log \pi_{\theta}(x)$ into the gradient expression for $J_{\text{RKL}}(\theta)$:

$$\begin{aligned}\nabla_{\theta} J_{\text{RKL}}(\theta) &= \mathbb{E}_{x \sim \pi_{\text{old}}} [(\nabla_{\theta} w(x))R(x) - \beta w(x) \nabla_{\theta} \log \pi_{\theta}(x) (\log w(x) + 1)] \\ &= \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) (\nabla_{\theta} \log \pi_{\theta}(x)) R(x) - \beta w(x) (\log w(x) + 1) \nabla_{\theta} \log \pi_{\theta}(x)] \\ &= \mathbb{E}_{x \sim \pi_{\text{old}}} \left[w(x) \left(R(x) - \beta (\log w(x) + 1) \right) \nabla_{\theta} \log \pi_{\theta}(x) \right].\end{aligned}$$

This proves the first part of the theorem.

Now, consider the surrogate loss function:

$$\mathcal{L}_{\text{RKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} [w(x) (-R(x) + \beta \log w(x))].$$

We compute its gradient:

$$\begin{aligned}\nabla_{\theta} \mathcal{L}_{\text{RKL}}(\theta) &= \nabla_{\theta} \mathbb{E}_{x \sim \pi_{\text{old}}} [-w(x)R(x) + \beta w(x) \log w(x)] \\ &= \mathbb{E}_{x \sim \pi_{\text{old}}} [\nabla_{\theta} (-w(x)R(x)) + \beta \nabla_{\theta} (w(x) \log w(x))] \\ &= \mathbb{E}_{x \sim \pi_{\text{old}}} [-(\nabla_{\theta} w(x))R(x) + \beta w(x) \nabla_{\theta} \log \pi_{\theta}(x) (\log w(x) + 1)] \\ &= \mathbb{E}_{x \sim \pi_{\text{old}}} [-w(x) (\nabla_{\theta} \log \pi_{\theta}(x)) R(x) + \beta w(x) (\log w(x) + 1) \nabla_{\theta} \log \pi_{\theta}(x)] \\ &= \mathbb{E}_{x \sim \pi_{\text{old}}} \left[w(x) \left(-R(x) + \beta (\log w(x) + 1) \right) \nabla_{\theta} \log \pi_{\theta}(x) \right] \\ &= -\mathbb{E}_{x \sim \pi_{\text{old}}} \left[w(x) \left(R(x) - \beta (\log w(x) + 1) \right) \nabla_{\theta} \log \pi_{\theta}(x) \right].\end{aligned}$$

Comparing this with the gradient of the objective function, we confirm that $\nabla_{\theta} \mathcal{L}_{\text{RKL}}(\theta) = -\nabla_{\theta} J_{\text{RKL}}(\theta)$. $\square$

M.4 Proof of Proposition D.3 (Policy Gradient and Differentiable Loss for Unnormalized Reverse KL)

Proof. We again express the objective components using expectations over the normalized reference distribution $\tilde{\pi}_{\text{old}}(x) = \pi_{\text{old}}(x)/Z_{\text{old}}$, with $w(x) = \pi_{\theta}(x)/\pi_{\text{old}}(x)$.

The expected reward term: $\mathbb{E}_{x \sim \pi_{\theta}} [R(x)] = Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x)]$.

The unnormalized reverse KL divergence $\text{UKL}(\pi_{\theta} \parallel \pi_{\text{old}})$ has two parts. Part 1 (Generalized KL):

$$\begin{aligned}\int \pi_{\theta}(x) \log \frac{\pi_{\theta}(x)}{\pi_{\text{old}}(x)} dx &= \int \pi_{\theta}(x) \log w(x) dx \\ &= \int w(x) \pi_{\text{old}}(x) \log w(x) dx \\ &= \int w(x) Z_{\text{old}} \tilde{\pi}_{\text{old}}(x) \log w(x) dx \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x) \log w(x)].\end{aligned}$$

Part 2 (Mass Correction):

$$\begin{aligned}\int (\pi_{\text{old}}(x) - \pi_{\theta}(x)) dx &= \int \pi_{\text{old}}(x) dx - \int \pi_{\theta}(x) dx \\ &= Z_{\text{old}} - \int w(x) \pi_{\text{old}}(x) dx \\ &= Z_{\text{old}} - \int w(x) Z_{\text{old}} \tilde{\pi}_{\text{old}}(x) dx\end{aligned}$$

$$= Z_{\text{old}} - Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)].$$

Combining these for the UKL term:

$$\text{UKL}(\pi_\theta \| \pi_{\text{old}}) = Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x) \log w(x)] + Z_{\text{old}} - Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)].$$

Now, substitute into the objective $J_{\text{URKL}}(\theta)$:

$$\begin{aligned} J_{\text{URKL}}(\theta) &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x)] - \beta (Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x) \log w(x)] + Z_{\text{old}} - Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)]) \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x) - \beta w(x) \log w(x) - \beta + \beta w(x)]. \end{aligned}$$

We compute the gradient $\nabla_\theta J_{\text{URKL}}(\theta)$. The constant term $-\beta Z_{\text{old}}$ vanishes upon differentiation.

$$\begin{aligned} \nabla_\theta J_{\text{URKL}}(\theta) &= \nabla_\theta (Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x) - \beta w(x) \log w(x) + \beta w(x)]) \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\nabla_\theta (w(x)R(x)) - \beta \nabla_\theta (w(x) \log w(x)) + \beta \nabla_\theta w(x)]. \end{aligned}$$

Using the previously derived gradients $\nabla_\theta w(x) = w(x) \nabla_\theta \log \pi_\theta(x)$ and $\nabla_\theta (w(x) \log w(x)) = w(x) \nabla_\theta \log \pi_\theta(x) (\log w(x) + 1)$:

$$\begin{aligned} \nabla_\theta J_{\text{URKL}}(\theta) &= \nabla_\theta (Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x) - \beta w(x) \log w(x) + \beta w(x)]) \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\nabla_\theta (w(x)R(x)) - \beta \nabla_\theta (w(x) \log w(x)) + \beta \nabla_\theta w(x)] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [(\nabla_\theta w(x))R(x) - \beta w(x) \nabla_\theta \log \pi_\theta(x) (\log w(x) + 1) + \beta (\nabla_\theta w(x))] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [w(x)R(x) \nabla_\theta \log \pi_\theta(x) - \beta w(x) (\log w(x) + 1) \nabla_\theta \log \pi_\theta(x) \\ &\quad + \beta w(x) \nabla_\theta \log \pi_\theta(x)] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[w(x) \nabla_\theta \log \pi_\theta(x) \left(R(x) - \beta (\log w(x) + 1) + \beta \right) \right] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[w(x) \nabla_\theta \log \pi_\theta(x) \left(R(x) - \beta \log w(x) \right) \right] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[w(x) \left(R(x) - \beta \log w(x) \right) \nabla_\theta \log \pi_\theta(x) \right]. \end{aligned}$$

This proves the first part of the theorem.

Now, consider the surrogate loss function:

$$\mathcal{L}_{\text{URKL}}(\theta) = Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-w(x)R(x) + \beta (w(x) \log w(x) - w(x))].$$

We compute its gradient:

$$\begin{aligned} \nabla_\theta \mathcal{L}_{\text{URKL}}(\theta) &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\nabla_\theta (-w(x)R(x)) + \beta \nabla_\theta (w(x) \log w(x) - w(x))] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-(\nabla_\theta w(x))R(x) + \beta (\nabla_\theta (w(x) \log w(x)) - \nabla_\theta w(x))] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-w(x)R(x) \nabla_\theta \log \pi_\theta(x) \\ &\quad + \beta (w(x) (\log w(x) + 1) \nabla_\theta \log \pi_\theta(x) - w(x) \nabla_\theta \log \pi_\theta(x))] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [-w(x)R(x) \nabla_\theta \log \pi_\theta(x) + \beta w(x) \log w(x) \nabla_\theta \log \pi_\theta(x)] \\ &= Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[w(x) \left(-R(x) + \beta \log w(x) \right) \nabla_\theta \log \pi_\theta(x) \right] \\ &= -Z_{\text{old}} \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[w(x) \left(R(x) - \beta \log w(x) \right) \nabla_\theta \log \pi_\theta(x) \right]. \end{aligned}$$

Comparing this with the gradient of the objective function, we confirm that $\nabla_\theta \mathcal{L}_{\text{URKL}}(\theta) = -\nabla_\theta J_{\text{URKL}}(\theta)$. The constant term $+1$ (corresponding to $-\beta Z_{\text{old}}$ in the objective) that appeared in the derivation in Section D does not affect the gradient and is often omitted from the final loss expression used in practice. $\square$

N Proofs for REINFORCE-Style Regularized Policy Gradients

This section provides justifications for the REINFORCE-style surrogate loss functions presented in Section E (Theorems H.1 to H.5). These proofs demonstrate how automatic differentiation applied to the proposed losses, utilizing the stop-gradient operator SG, yields the correct gradient direction (negative of the objective gradient derived in Section 2).

The core idea relies on the operational definition of the stop-gradient operator $\text{SG}(\cdot)$ within automatic differentiation frameworks: $\nabla_\theta \text{SG}(f(\theta)) = 0$, while the forward computation uses the value of $f(\theta)$. We use the notation $w(x) = \pi_\theta(x)/\pi_{\text{old}}(x)$.

N.1 Proof of Proposition H.1 (REINFORCE-style Policy Gradient for Forward KL)

Proof. The objective is $J_{\text{FKL}}(\theta) = \mathbb{E}_{\pi_\theta}[R(x)] - \beta \text{KL}(\pi_{\text{old}} \parallel \pi_\theta)$. From Proposition G.1, its gradient is:

$$\nabla_\theta J_{\text{FKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} \left[\underbrace{(w(x)R(x) + \beta)}_{\text{Weight}_{\text{FKL}}(x, \theta)} \nabla_\theta \log \pi_\theta(x) \right].$$

The proposed REINFORCE-style surrogate loss is:

$$\mathcal{L}_{\text{FKL}}^{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \pi_{\text{old}}} [\text{SG}(w(x)R(x) + \beta) \log \pi_\theta(x)].$$

We compute the gradient of this loss as it would be computed by an automatic differentiation system. Assuming the gradient can be swapped with the expectation:

$$\begin{aligned} \nabla_\theta \mathcal{L}_{\text{FKL}}^{\text{REINFORCE-style}}(\theta) &= -\mathbb{E}_{x \sim \pi_{\text{old}}} [\nabla_\theta (\text{SG}(w(x)R(x) + \beta) \log \pi_\theta(x))] \\ &= -\mathbb{E}_{x \sim \pi_{\text{old}}} \left[\underbrace{(\nabla_\theta \text{SG}(w(x)R(x) + \beta)) \log \pi_\theta(x)}_{=0 \text{ by definition of SG}} \right. \\ &\quad \left. + \text{SG}(w(x)R(x) + \beta) (\nabla_\theta \log \pi_\theta(x)) \right] \\ &= -\mathbb{E}_{x \sim \pi_{\text{old}}} [\text{SG}(w(x)R(x) + \beta) \nabla_\theta \log \pi_\theta(x)]. \end{aligned}$$

This gradient expression, when used in an optimization algorithm (where SG is conceptually removed), corresponds to applying updates proportional to:

$$-(\mathbb{E}_{x \sim \pi_{\text{old}}} [(w(x)R(x) + \beta) \nabla_\theta \log \pi_\theta(x)]) = \nabla_\theta J_{\text{FKL}}(\theta).$$

Thus, minimizing $\mathcal{L}_{\text{FKL}}^{\text{REINFORCE-style}}(\theta)$ using gradient descent with automatic differentiation effectively performs gradient ascent on the original objective $J_{\text{FKL}}(\theta)$. $\square$

N.2 Proof of Proposition H.3 ((REINFORCE-style Policy Gradient for Unnormalized Forward KL)

Proof. The objective is $J_{\text{UFKL}}(\theta) = \mathbb{E}_{\pi_\theta}[R(x)] - \beta \text{UKL}(\pi_{\text{old}} \parallel \pi_\theta)$. From Proposition 2.2, its gradient is:

$$\nabla_\theta J_{\text{UFKL}}(\theta) = \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\underbrace{Z_{\text{old}}(w(x)R(x) - \beta(w(x) - 1))}_{\text{Weight}_{\text{UFKL}}(x, \theta)} \nabla_\theta \log \pi_\theta(x) \right].$$

The proposed REINFORCE-style surrogate loss is:

$$\mathcal{L}_{\text{UFKL}}^{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\text{SG}(Z_{\text{old}}(w(x)R(x) - \beta(w(x) - 1))) \log \pi_\theta(x)].$$

Computing the gradient via automatic differentiation:

$$\begin{aligned} \nabla_\theta \mathcal{L}_{\text{UFKL}}^{\text{REINFORCE-style}}(\theta) &= -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\nabla_\theta (\text{SG}(Z_{\text{old}}(\dots)) \log \pi_\theta(x))] \\ &= -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\underbrace{(\nabla_\theta \text{SG}(Z_{\text{old}}(\dots))) \log \pi_\theta(x)}_{=0} + \text{SG}(Z_{\text{old}}(\dots)) (\nabla_\theta \log \pi_\theta(x)) \right] \\ &= -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\text{SG}(Z_{\text{old}}(w(x)R(x) - \beta(w(x) - 1))) \nabla_\theta \log \pi_\theta(x)]. \end{aligned}$$

This gradient corresponds to the update direction $-\nabla_\theta J_{\text{UFKL}}(\theta)$ when the SG is dropped. Minimizing this loss achieves gradient ascent on $J_{\text{UFKL}}(\theta)$. If Z_{old} is omitted, the same argument applies to the proportionally scaled objective and loss. $\square$

N.3 Proof of Proposition H.4 (REINFORCE-Style Loss)

Proof. The objective is $J_{\text{RKL}}(\theta) = \mathbb{E}_{\pi_\theta}[R(x)] - \beta \text{KL}(\pi_\theta \parallel \pi_{\text{old}})$. From Proposition G.3, its gradient is:

$$\nabla_\theta J_{\text{RKL}}(\theta) = \mathbb{E}_{x \sim \pi_{\text{old}}} \left[\underbrace{w(x) \left(R(x) - \beta (\log w(x) + 1) \right)}_{\text{Weight}_{\text{RKL}}(x, \theta)} \nabla_\theta \log \pi_\theta(x) \right].$$

The proposed REINFORCE-style surrogate loss is:

$$\mathcal{L}_{\text{RKL}}^{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \pi_{\text{old}}} [\text{SG}(w(x) (R(x) - \beta \log w(x) - \beta)) \log \pi_\theta(x)].$$

Computing the gradient via automatic differentiation:

$$\begin{aligned} \nabla_\theta \mathcal{L}_{\text{RKL}}^{\text{REINFORCE-style}}(\theta) &= -\mathbb{E}_{x \sim \pi_{\text{old}}} [\nabla_\theta (\text{SG}(w(x)(\dots)) \log \pi_\theta(x))] \\ &= -\mathbb{E}_{x \sim \pi_{\text{old}}} \left[\underbrace{(\nabla_\theta \text{SG}(w(x)(\dots)))}_{=0} \log \pi_\theta(x) + \text{SG}(w(x)(\dots)) (\nabla_\theta \log \pi_\theta(x)) \right] \\ &= -\mathbb{E}_{x \sim \pi_{\text{old}}} [\text{SG}(w(x) (R(x) - \beta \log w(x) - \beta)) \nabla_\theta \log \pi_\theta(x)]. \end{aligned}$$

This gradient corresponds to the update direction $-\nabla_\theta J_{\text{RKL}}(\theta)$ when the SG is dropped. Minimizing this loss achieves gradient ascent on $J_{\text{RKL}}(\theta)$. $\square$

N.4 Proof of Proposition H.5 (REINFORCE-Style Loss for Unnormalized Reverse KL)

Proof. The objective is $J_{\text{URKL}}(\theta) = \mathbb{E}_{\pi_\theta}[R(x)] - \beta \text{UKL}(\pi_\theta \parallel \pi_{\text{old}})$. From Proposition D.3, its gradient is:

$$\nabla_\theta J_{\text{URKL}}(\theta) = \mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\underbrace{Z_{\text{old}} w(x) \left(R(x) - \beta \log w(x) \right)}_{\text{Weight}_{\text{URKL}}(x, \theta)} \nabla_\theta \log \pi_\theta(x) \right].$$

The proposed REINFORCE-style surrogate loss is:

$$\mathcal{L}_{\text{URKL}}^{\text{REINFORCE-style}}(\theta) = -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\text{SG}(Z_{\text{old}} w(x) (R(x) - \beta \log w(x))) \log \pi_\theta(x)].$$

Computing the gradient via automatic differentiation:

$$\begin{aligned} \nabla_\theta \mathcal{L}_{\text{URKL}}^{\text{REINFORCE-style}}(\theta) &= -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\nabla_\theta (\text{SG}(Z_{\text{old}} w(x)(\dots)) \log \pi_\theta(x))] \\ &= -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} \left[\underbrace{(\nabla_\theta \text{SG}(Z_{\text{old}} w(x)(\dots)))}_{=0} \log \pi_\theta(x) \right. \\ &\quad \left. + \text{SG}(Z_{\text{old}} w(x)(\dots)) (\nabla_\theta \log \pi_\theta(x)) \right] \\ &= -\mathbb{E}_{x \sim \tilde{\pi}_{\text{old}}} [\text{SG}(Z_{\text{old}} w(x) (R(x) - \beta \log w(x))) \nabla_\theta \log \pi_\theta(x)]. \end{aligned}$$

This gradient corresponds to the update direction $-\nabla_\theta J_{\text{URKL}}(\theta)$ when the SG is dropped. Minimizing this loss achieves gradient ascent on $J_{\text{URKL}}(\theta)$. If Z_{old} is omitted, the same argument applies to the proportionally scaled objective and loss. $\square$

O Broader Impact and Limitations

The methods developed in this paper contribute to the broader effort of enhancing the reasoning capabilities of large language models. Improved reasoning in LLMs has the potential to significantly benefit various fields, including scientific discovery, education, and complex problem-solving in engineering and medicine. By providing more stable and efficient training algorithms, our work can facilitate the development of more reliable and capable AI systems. However, as with any advancement in AI capabilities, it is crucial to consider the ethical implications and ensure responsible development and deployment of these technologies to mitigate potential misuse.

P 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: We describe all the contributions and scope in the abstract and introduction parts.

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: We discussed aspects like using a fixed KL regularization coefficient (β).

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: We list all the assumptions and proofs in Appendix [L](#), [M](#), and [N](#).

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: We listed all the experiment details in Section [3](#) for reproduction of our work.

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: The code and data required to reproduce the main experimental results are provided at <https://anonymous.4open.science/r/verl-neo-pub-3D2D>. The supplemental material will contain instructions for their use.

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: We just list all the training and test details in Section 3.

Guidelines:

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

7. Experiment statistical significance

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

Answer: [No]

Justification: The error bars are not reported because it would be too computationally expensive for repeated experiments on LLMs.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)

- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

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

Answer: [Yes]

Justification: We just list all the computer resources in Section 3.

Guidelines:

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

9. Code of ethics

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

Answer: [Yes]

Justification: Our research only explores more efficient methods for fine-tuning large language models with reinforcement learning approaches. Therefore, the research conducted in the paper conform, in every respect, 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]

Section O discusses potential positive societal benefits of improved LLM reasoning and also acknowledges the need to consider ethical implications and mitigate potential misuse.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.

- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

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

Answer: [NA]

Justification: Our work proposes a systematic framework for fine-tuning large language models using reinforcement learning methods. To our knowledge, this work has no direct path to any negative applications.

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 add the citation to all the codes (verl: Apache-2.0 License), models (Qwen2.5-7B-Instruct and Qwen2.5-Math-7B: Apache-2.0 License) and datasets (DAPO-Math-17k: Apache-2.0 License; AMC23, AIME24 and AIME25 are downloaded from huggingface: <https://huggingface.co/math-ai>, which sources from the website of Mathematical Association of America's American Mathematics Competitions) that we used in this work.

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: The code implementing our proposed RPG framework and experimental setup is released at <https://anonymous.4open.science/r/verl-neo-pub-3D2D>. This code will be documented to facilitate understanding and use by other researchers. No new datasets or pre-trained models are introduced beyond the code for the methods.

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: The paper only use open-source codes, checkpoints and datasets which do 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: The paper only use open-source codes, checkpoints and datasets which do 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: This work aims at exploring more efficient methods for fine-tuning large language models with reinforcement learning approaches. Therefore, LLMs, including Qwen2.5-7B-Instruct and Qwen2.5-Math-7B are used and well described in the main part of this paper.

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.