

Whatever Remains Must Be True: Filtering Drives Reasoning in LLMs, Shaping Diversity

Germán Kruszewski¹² Pierre Erbacher¹ Jos Rozen¹ Marc Dymetman³

¹NAVER Labs Europe ²Universitat Pompeu Fabra ³Independent Researcher

{german.kruszewski,pierre.erbacher,jos.rozen}@naverlabs.com marc.dymetman@gmail.com

Abstract

Reinforcement Learning (RL) has become the *de facto* standard for tuning LLMs to solve tasks involving reasoning. However, growing evidence shows that models trained in such way often suffer from a significant loss in diversity. We argue that this arises because RL implicitly optimizes the “mode-seeking” or “zero-forcing” *Reverse KL* to a target distribution causing the model to concentrate mass on certain high-probability regions of the target while neglecting others. In this work, we instead begin from an explicit target distribution, obtained by filtering out incorrect answers while preserving the relative probabilities of correct ones. Starting from a pre-trained LLM, we approximate this target distribution using the α -divergence family, which unifies prior approaches and enables direct control of the precision–diversity trade-off by interpolating between mode-seeking and mass-covering divergences. On a LEAN theorem-proving benchmark, our method achieves state-of-the-art performance along the coverage–precision Pareto frontier, outperforming all prior methods on the coverage axis.

1. Introduction

“How often have I said to you that when you have eliminated the impossible, whatever remains, however improbable, must be the truth?”

— Arthur Conan Doyle, *The Sign of Four*

Large Language Models (LLMs) have made striking progress on reasoning tasks. A leading approach is Reinforcement Learning from Verifiable Rewards (RLVR) [16,30], where policy-gradient methods such as PPO [49] or GRPO [50] optimize against a reward that combines a binary verifier of correctness with a KL penalty to keep the tuned model close to its base distribution.

While RLVR has been credited with enabling exploration of new solutions [16], recent studies challenge this view. In particular, they find that base models already contain these solutions given a sufficient sampling budget, and that tuned models often exploit additional samples less effectively due to reduced output diversity [14,19,61,68]. Earlier with RLHF [9,71], similar diversity reductions were observed sometimes described as “mode collapse” [26,42].

We argue that this loss of diversity stems from the *implicit objective* of RL-based training: optimizing the *Reverse KL* divergence to a target distribution that favors correct answers [29]. Reverse KL is “mode-seeking” or “zero-forcing,” emphasizing precision on a subset of solutions while ignoring others [5,22,33]. This explains why RLVR models become accurate but less diverse.

To address this, we explicitly define the desired target distribution: one that *always* outputs correct solutions while remaining as close as possible to the base model, thus preserving any solution included therein [24,25]. Direct sampling is infeasible, but we can approximate it with an autoregressive policy using Distributional Policy Gradient Algorithms [DPG 18,24,43]. Concretely, we apply f -DPG [18], which minimizes an f -divergence [45] to this target. Different divergences trade off precision (probability of sampling a correct solution) and coverage (probability of sampling at least one correct given a sufficiently large sampling budget): Reverse KL emphasizes the former, Forward KL the latter [5]. To interpolate between them, we introduce α -DPG, based on α -divergences [3,12,48]. Notably, this method unifies RLVR (Reverse KL) and both KL-DPG [18,24,43] and Rejection Sampling

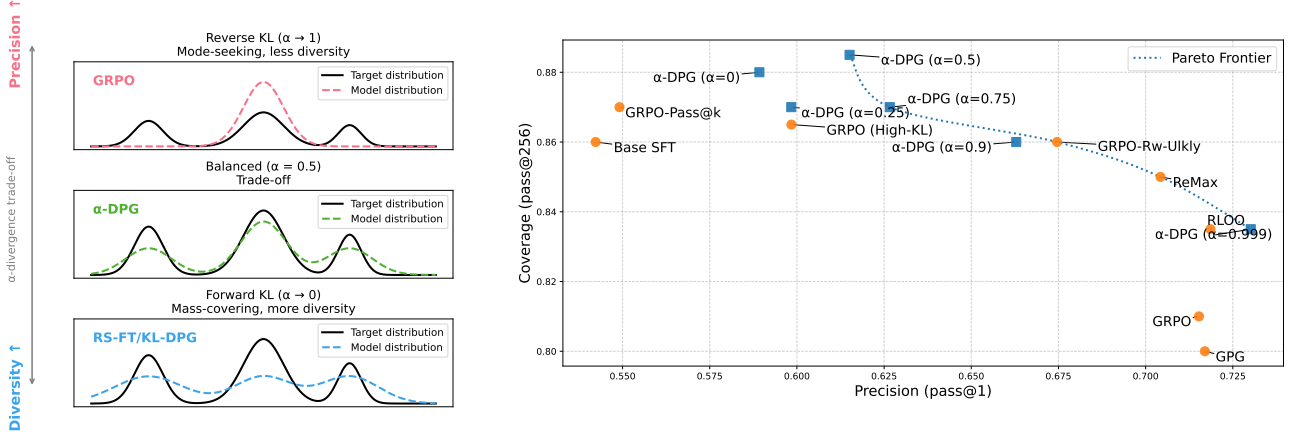


Figure 1: Left: Illustrative representation of our method. GRPO/PPO and other policy-gradient methods used in RLVR focus the model on a small region of the target distribution. Other methods, such as KL-DPG, recover more of the diversity at the cost of putting probability mass to low-quality regions. α -DPG allows to strike a balance between the two. Right: Estimates of models precision (pass@1) and coverage (pass@256). α -DPG models sit along a Pareto frontier.

Fine-Tuning (Forward KL) [66,69] under a single umbrella. We coin this approach Distributional Matching with Verifiable Rewards (DMVR), positioning it under the general framework of Distributional Matching [24,28,29].

We evaluate α -DPG in LEAN, a proof assistant that automatically verifies formal mathematical proofs. In this setting, success requires not only correct candidates but also diversity across proof attempts, since harder theorems may only be solved by rare derivations. We find that α -DPG achieves state-of-the-art performance, producing models that lie on the Pareto frontier between precision (pass@1) and coverage (pass@256) and that surpass prior methods in coverage.

Contributions.

- We introduce the DMVR framework, which trains models by approximating an explicitly defined verifier-based target distribution.
- We clarify how the implicit dynamics of RL-based methods lead to reduced diversity.
- We highlight the role of the divergence family in trading off precision and diversity, and propose α -DPG to smoothly interpolate between Forward and Reverse KL.
- We show on the LEAN benchmark results that are Pareto-optimal along the coverage-precision frontier, with the α parameter allowing to optimally trade-off between precision and coverage. Moreover, α -DPG achieves the best coverage results among all considered methods when using intermediate values of α .

2. Background

Reinforcement Learning with Verifiable Rewards (RLVR) Let $\pi_\theta(\cdot|x) : \mathcal{Y} \rightarrow \mathbb{R}$ be an LLM with parameters θ defining a probability distribution over sequences $y \in \mathcal{Y}$ conditioned on a prompt x . A verifier $v(y, x) \in \{0, 1\}$ is a binary function that discriminates between correct and incorrect responses to x . Given a dataset $\mathcal{D}$ of prompts (and, optionally, ground-truth answers for the verifier), RLVR uses a standard reinforcement learning objective derived from previous work on Reinforcement Learning from Human Feedback (RLHF) [9,71]:

$$\mathcal{J}_{\text{RLVR}} = \arg \max_{\theta} \mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_\theta(\cdot|x)} R_\theta(y, x), \quad (1)$$

where the “pseudo-reward” [29] R_θ is defined as follows:

$$R_\theta(y, x) = v(y, x) - \beta \log \frac{\pi_\theta(y|x)}{\pi_{\text{base}}(y|x)}. \quad (2)$$

Here, π_{base} is the base LLM and β is a tunable parameter that controls the trade-off between maximizing the reward and minimizing divergence from the original model. Some authors fall back to setting $\beta = 0$, just optimizing the expected verifier reward [38,39,65].

Policy Gradient (PG) Algorithms We can maximize Eq. 2 following multiple algorithms. One of the simplest ones is KL-Control [23,57], with the following gradient

$$\nabla_{\theta} \mathcal{J}_{\text{KL-Control}}(\theta) = \mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_{\theta}} \hat{A}(y, x) \nabla_{\theta} \log \pi_{\theta}(y|x), \quad (3)$$

where $\hat{A}(y, x) = R_{\theta}(y, x) - B(x)$ is the advantage function and $B(x)$ is a baseline that doesn't depend on y to keep the objective unbiased used for variance reduction [53, Chapter 2, Section 7]. Some options for an unbiased baseline could include a constant, a critic [54], or a leave-one-out average in a batch of rewards [RLOO; 1,27]¹. When $\beta = 0$, Eq. 3 reduces to the original policy gradient algorithm REINFORCE [60]. Another popular choice is PPO [49], which optimizes the following clipped surrogate objective:

$$\mathcal{J}^{\text{PPO}}(\theta) = \mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_{\theta_{\text{old}}}} \left[\sum_{t=1}^{|y|} \min \left(\rho(y_t|x, y_{<t}) \hat{A}(y, x), \text{clip}(\rho(y_t|x, y_{<t}), 1 - \epsilon, 1 + \epsilon) \hat{A}(y, x) \right) \right], \quad (4)$$

where $\rho(y_t|x, y_{<t}) = \pi_{\theta}(y_t|x, y_{<t}) / \pi_{\theta_{\text{old}}}(y_t|x, y_{<t})$ is the ratio between the current and the previous policy token probabilities and ϵ is a small hyperparameter. PPO discourages large policy updates that would move ρ outside the interval $[1 - \epsilon, 1 + \epsilon]$, which improves training stability by constraining the new policy to be close to the old one. GRPO [16] and other recent methods use the same clipping strategy, but differ in the form of the baseline. While PPO is commonly understood to use a critic model as a baseline, GRPO uses the group-level reward average.

Distribution Matching (DM) DM [18,24,25,28,29] is a family of techniques for aligning LLMs to meet certain constraints. For instance, let's suppose that we want to constrain the LM so that all sequences y meet $r(x, y) = 1$ for some binary filter $r(x, y) \in \{0, 1\}$ ². The method expresses this as a target distribution we want to approximate, which in this case is given by

$$p_x(y) \propto \pi_{\text{base}}(y|x) r(y, x). \quad (5)$$

This distribution is the only distribution p' that fulfills the following two desirable conditions: (i) it satisfies the constraint $r(y, x) = 1, \forall y \in \text{Supp}(p'(\cdot|x))$, and (ii) it is the closest to $\pi_{\text{base}}(\cdot|x)$ in terms of $D_{\text{KL}}(p'(\cdot|x) || \pi_{\text{base}}(\cdot|x))$. The second condition guarantees the preservation of all the diversity contained in the original model. In information geometric terms, p_x is the I-projection of $\pi_{\text{base}}(\cdot|x)$ into the manifold of all distributions that satisfy the constraint given by r [13].

Distributional Policy Gradient (DPG) Algorithms Given a target distribution, we can optimize an autoregressive policy π_{θ} to approximate it. Khalifa et al. [24] used for this the DPG algorithm [28,29,43], later denoted KL-DPG, which minimizes the Forward KL $D_{\text{KL}}(p_x || \pi_{\theta})$ to the target distribution p_x :

$$\nabla_{\theta} \mathcal{L}^{\text{DPG}}(\theta) = \nabla_{\theta} D_{\text{KL}}(p_x || \pi_{\theta}) = -\mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_{\theta}(\cdot|x)} \left(\frac{p_x(y)}{\pi_{\theta}(y|x)} - 1 \right) \nabla_{\theta} \log \pi_{\theta}(y|x). \quad (6)$$

The negative one term acts as a baseline. Computing $p_x(y|x)$ requires estimating a normalization constant (partition function) Z_x , which can be estimated by importance sampling [41] (see App. E for additional details). Note that in the case of a binary constraint, Z_x is the acceptance rate of the constraint when sampling from the base model [25]: $Z_x = \sum_{y \in \mathcal{Y}} a(y|x) r(y, x) = \mathbb{E}_{y \sim a(\cdot|x)} r(y, x) = \mathbb{P}_{y \sim a(\cdot|x)} [r(y, x) = 1]$. Later, Go et al. [18] introduced f -DPG, effectively generalizing this technique to any f -divergence by minimizing $D_f(\pi_{\theta}, p_x)$:

$$\nabla_{\theta} \mathcal{L}^{f\text{-DPG}}(\theta) = \nabla_{\theta} \mathbb{E}_{x \sim \mathcal{D}} [D_f(\pi_{\theta}(\cdot|x) || p_x)] = \mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_{\theta}(\cdot|x)} [-\hat{A}^f(y, x) \nabla_{\theta} \log \pi_{\theta}(y|x)], \quad (7)$$

where $\hat{A}^f(y, x) = R_{\theta}^f(y, x) - B(x)$, $R_{\theta}^f(y, x) \doteq -f' \left(\frac{\pi_{\theta}(y|x)}{p_x(y)} \right)$ is a "pseudo-reward", $B(x)$ is a context-dependent baseline, and f is a convex function that parametrizes the f -divergence where $f : (0, \infty) \rightarrow \mathbb{R}$ s.t. $f(1) = 0$. If $p_x(y) = 0$, then, $R_{\theta}^f(y, x) = -f'(\infty)$, where $f'(\infty) \doteq \lim_{t \rightarrow 0} t f(\frac{1}{t})$.

¹Note that the average of all samples, as is used in GRPO, is biased instead.

²If the reader notices a strong similarity between a constraint $r(x, y)$ and a verifier $v(x, y)$, this is no coincidence: It is exactly this connection that we will exploit in this paper.

3. Distributional Matching with Verifiable Rewards (DMVR)

Here, we adopt the DM framework, and propose that the ideal target distribution for tuning a language model $\pi_{\text{base}}(\cdot|x)$ to solve problems $x \sim \mathcal{D}$ by means of a verifier $v(y, x) \in \{0, 1\}$ is simply the result of applying the verifier as a binary constraint, i.e. $r(y, x) = v(y, x)$, as follows:

$$p_x(y) \propto \pi_{\text{base}}(y|x)v(y, x) \quad \forall x \in \mathcal{D}. \quad (8)$$

This distribution *filters out all incorrect responses*, leaving out only correct ones with the same relative probabilities as the reference LLM. This is the single distribution that (i) always answers correctly to x , and (ii) it is closest to the base model π_{base} as measured by $D_{\text{KL}}(\cdot||\pi_{\text{base}})$ [24]. In the following discussion, we argue that (1) the distribution approximated by RLVR is closely related, becoming equivalent to p_x in the limit $\beta \rightarrow 0$, (2) for any fixed β , RLVR optimizes the Reverse KL to a target distribution, which has a mode-seeking behavior, incentivizing the policy to put high probability mass in small regions of high reward at the cost of diversity, and (3) we propose an alternative based on f -DPG parametrized with α -divergences to trade-off precision and diversity. Of these, points (1) and (3) are original to our work, whereas point (2) is reproduced from Korbak et al. [29]. Moreover, the target distribution and the f -DPG technique to approximate it were defined in prior art [18,24], as detailed in Section 2.

3.1. From RLVR to DMVR

Consider the following lemma, reproducing the argument from Korbak et al. [29]:

Lemma 1. *Define*

$$p_{x,\beta}(y) = \frac{1}{Z_x(\beta)} \pi_{\text{base}}(y|x) \exp(v(y, x)/\beta), \quad Z_x(\beta) = \sum_y \pi_{\text{base}}(y|x) \exp(v(y, x)/\beta).$$

Then,

$$\nabla_{\theta} \mathbb{E}_x[\text{KL}(\pi_{\theta}||p_{x,\beta})] = -\mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_{\theta}} \left[\frac{1}{\beta} v(y, x) - \log \frac{\pi_{\theta}(y|x)}{\pi_{\text{base}}(y|x)} \right] \nabla_{\theta} \log \pi_{\theta}(y|x) \quad (9)$$

$$= -\frac{1}{\beta} \mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_{\theta}} \left[\underbrace{v(y, x) - \beta \log \frac{\pi_{\theta}(y|x)}{\pi_{\text{base}}(y|x)}}_{\text{RLVR pseudo-reward}} \right] \nabla_{\theta} \log \pi_{\theta}(y|x). \quad (10)$$

Proof is in App. B. From Eq. 10, we can see that the gradient of the KL-Control’s objective (right-hand-side, optimizing the expected RLVR pseudo-reward) is proportional to the gradient of the Reverse KL of π_{θ} to $p_{x,\beta}$ (left-hand-side), up to a negative constant that flips the direction of optimization. Therefore, maximizing the regularized reward implies minimizing the divergence, and vice versa. Furthermore, the distribution $p_{x,\beta}$ is a smooth approximation to the ideal distribution defined in Eq. 8, p_x , converging to it as $\beta \rightarrow 0^+$:

Lemma 2.

$$\lim_{\beta \rightarrow 0} p_{x,\beta} = p_x. \quad (11)$$

(Proof in App. C).

However, the gradient of the Reverse KL is dominated by $-\frac{1}{\beta} \nabla_{\theta} \mathbb{E}_{y \sim \pi_{\theta}} [v(y, x)]$ (Eq. 9), revealing why minimizing Reverse KL becomes an increasingly aggressive mode-seeking proxy for maximizing expected reward, at the cost of preserving diversity even if the diversity of responses is well-captured by the target p_x . In the limit, with $\beta = 0$, the Reverse KL becomes undefined, the KL-Control algorithm reduces to plain REINFORCE, and no safeguard remains to preserve diversity.

3.2. And back (as a special case of α -DPG)

Having defined the target distribution p_x , we are left with the task of picking a divergence to train a policy π_{θ} to approximate it. As we have seen, the Reverse KL $D_{\text{KL}}(\pi_{\theta}||p)$ implicitly employed by RLVR, is a mode-seeking or zero-forcing divergence [5,22,34]. This means that it will penalize placing probability mass in regions where p

assigns little or none, but it is relatively indifferent to ignoring modes of p altogether. As a result, the learned policy tends to concentrate on a small subset of high-probability modes while disregarding other plausible regions of the target distribution, which can reduce diversity and lead to brittle or degenerate behavior.

In contrast, the Forward KL, $D_{\text{KL}}(p \parallel \pi_\theta)$, is mass-covering. Here, the divergence becomes large whenever π_θ assigns insufficient probability to regions where p has support, encouraging the policy to cover all modes of the target distribution. While this can improve diversity and robustness, it often comes at the cost of assigning non-negligible probability mass to low-reward or unlikely regions, leading to less precise approximations of the target.

This tension between mode-seeking and mass-covering behavior motivates considering a broader family of divergences. The α -divergences family [3, 12, 48] provides exactly such a continuum, interpolating smoothly between the Forward KL (as $\alpha \rightarrow 0$) and the Reverse KL (as $\alpha \rightarrow 1$). In between (for $\alpha = 0.5$) lies the squared Hellinger distance [20]. By tuning α , one can balance the degree of mode-seeking versus mass-covering behavior, potentially capturing the benefits of both extremes while mitigating their drawbacks. We refer to Table 1 for a full characterization. We denote the parametrization of f -DPG with α -divergences as α -DPG, which has the following pseudo-reward:

$$R_\theta(y, x) = -f' \left(\frac{\pi_\theta(y|x)}{p_x(y)} \right) = \frac{1}{1-\alpha} \left(\left(\frac{p_x(y)}{\pi_\theta(y|x)} \right)^{1-\alpha} - 1 \right). \quad (12)$$

Because for low values of α , peaks in the p/π ratios can induce large variance in R_θ , we clip the parenthetical factor to a maximum M . We also rescale by discounting the constant $1/(1-\alpha)$:

$$\hat{R}_\theta(y, x) \doteq \min \left(\left(\frac{p_x(y)}{\pi_\theta(y|x)} \right)^{1-\alpha} - 1, M \right). \quad (13)$$

Finally, we use the gradient formula in Eq. 7, setting the baseline $B(x)$ to the leave-one-out per-context average of the pseudo rewards [1, 27].

It is interesting to note its behavior when setting $\alpha = 1 - \epsilon$ (i.e., close to the Reverse KL). Then,

$$\nabla_\theta D_{f_\alpha}(\pi_\theta(\cdot|x) \parallel p_x) = \mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_\theta} f'_\alpha \left(\frac{\pi_\theta(y|x)}{p_x(y)} \right) \nabla_\theta \log \pi_\theta(y|x) \quad (14)$$

$$\propto -\mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_\theta} \left(\frac{p_x(y)}{\pi_\theta(y|x)} \right)^\epsilon \nabla_\theta \log \pi_\theta(y|x) \quad (15)$$

$$\approx \mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_\theta} -v(y, x) \nabla_\theta \log \pi_\theta(y|x). \quad (16)$$

by discounting for simplicity the scaling constant and the baselines in Eq. 15, and noting that for $0 < \epsilon \ll 1$ and bounded $|x|$ then $x^\epsilon \approx \mathbb{I}[x \neq 0]$ and $p_x(y) \neq 0 \Leftrightarrow v(y, x) = 1$ in the last step. Thus, for α that is lower but very close to 1, we are again recovering the REINFORCE learning rule. In contrast, when $\alpha = 0$, then D_{f_α} becomes exactly the Forward KL, and thus we recover the original KL-DPG algorithm [24, 43], which conserves more diversity from the original distribution sacrificing some precision:

$$\nabla_\theta D_{f_\alpha}(\pi_\theta(\cdot|x) \parallel p_x) = -\mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_\theta} \left(\frac{p_x(y)}{\pi_\theta(y|x)} - 1 \right) \nabla_\theta \log \pi_\theta(y|x) \quad (17)$$

Notably, it is easy to see that if a fixed amount of samples are obtained just from π_{base} instead of π_θ , KL-DPG reduces to RS-FT [66, 69], as it optimizes the cross-entropy to samples from the base model filtered by the verifier (see App. D for more details).

4. Experiments

Informal vs Formal Mathematics Informal mathematics has become a common paradigm for training large language models (LLMs) on reasoning tasks, with widely used benchmarks such as MATH [21] and AIME [52]. These methods have achieved impressive performance [17], but they also face inherent challenges. Informal proofs and solutions often lack guarantees of rigor, making large-scale verification difficult and requiring heuristics like majority voting rather than provable correctness. While effective in many scenarios, these approaches may be less suited for tasks that aim to explore novel results or a wide variety of solutions.

Parameter	Name/Correspondence	Generator $f_\alpha(t)$ (generic)	$f'_\alpha(t)$ (generic)	$f'_\alpha(\pi/p)$	$f'_\alpha(\infty)$
$\alpha \neq 0, 1$	α -divergence (f -div.)	$\frac{t^\alpha - \alpha t - (1 - \alpha)}{\alpha(\alpha - 1)}$	$\frac{t^{\alpha-1} - 1}{\alpha - 1}$	$\frac{1}{\alpha - 1} \left(\left(\frac{p}{\pi} \right)^{1-\alpha} - 1 \right)$	$\alpha < 1 : \frac{1}{1-\alpha}$ $\alpha > 1 : +\infty$
$\alpha \rightarrow 1$	Reverse KL $\text{KL}(\pi \ p)$	$\lim_{\alpha \rightarrow 1} f_\alpha(t) = t \log t - t + 1$	$\lim_{\alpha \rightarrow 1} f'_\alpha(t) = \log t$	$\log(\pi/p)$	$+\infty$
$\alpha \rightarrow 0$	Forward KL $\text{KL}(p \ \pi)$	$\lim_{\alpha \rightarrow 0} f_\alpha(t) = -\log t + t - 1$	$\lim_{\alpha \rightarrow 0} f'_\alpha(t) = 1 - \frac{1}{t}$	$1 - p/\pi$	1
$\alpha = \frac{1}{2}$	Hellinger (exact)	$-4\sqrt{t} + 2t + 2$	$-\frac{2}{\sqrt{t}} + 2$	$-2\sqrt{p/\pi} + 2$	2

 Table 1: Parametrization of the α -divergence as an f -divergence $D_{f_\alpha}(\pi, p)$.

Formal methods provide a promising alternative [56]. Proof assistants such as LEAN [15], COQ [4], and ISABELLE [40] represent statements and proofs in a formally verifiable language. This not only allows for efficient and rigorous proof generation but also reduces dependence on labeled data. At inference, the paradigm shifts from achieving consensus over a large pool of candidates (eg., via majority voting) to ensuring sufficient diversity among them to guarantee larger coverage. Theorem provers must therefore go beyond producing a single correct proof but should generate diverse candidate proofs to more fully explore the solution space. Prior work in LEAN has applied reinforcement learning (e.g., GRPO [47, 58, 63]), but such methods can result in decreased diversity [19]. In this work, we present an approach to formal theorem proving in LEAN that seeks to improve both precision for better efficiency and diversity to guarantee high coverage, highlighting the potential of formal reasoning for scalable discovery.

Models For these experiments we consider DeepSeek-Prover-V1.5-SFT [63] a 7B parameters models based on DeepSeek model and further pre-trained on high-quality mathematics and code data, with a focus on formal languages such as LEAN, ISABELLE, and METAMATH, and finetuned on LEAN4 code completion datasets [63].

Dataset We follow the experimental configuration introduced in prior work [19]. The training set is composed of 10K solvable LEAN problems extracted and filtered from the LEAN Workbook dataset [62, 64], from which 200 problems are kept unseen as a test set.

Reward function and LEAN4 Verifier The reward function extracts the last LEAN4 code block in the generated sequence and verifies it automatically by the LEAN proof assistant. The sequence is given a reward of 1 if it is verified as correct and 0 otherwise. In our experiments we used the same LEAN4 and Mathlib4 version (lean4:v4.9.0) as used in DeepSeek-ProverV1.5 [63].

Baselines We compare against several baselines, focusing on critic-free methods, which have become standard thanks to not requiring an additional copy of the model. DeepSeek-Prover-SFT is the supervised fine-tuned model serving as the base for all methods. GRPO [16], for which we only consider its unbiased instantiation, Dr. GRPO [38], and other variants with diversity-preserving regularization: **High-KL** with a strong KL penalty ($\beta = 0.1$), and **Rw-Ulkly** [19] with a rank bias promoting diversity $\beta = 0.25$. Finally, **Pass@k training** [8, 55] directly optimizes pass@k via a leave-one-out advantage formulation that reduces variance. We further include in our comparison **GPG** [10], **ReMax** [36] and **RLOO** [1].

Training All trainings are done on a single node of 4xA100 with 28 CPUs dedicated for running proof assessment in parallel. Due to the high variance in LEAN compilation and verification time, the reward function is executed asynchronously. All training scripts are based on VERL [51]. By default, since we use a verifiable reward, we follow Liu et al. [38], Mistral-AI et al. [39] and disable advantage normalization while setting the KL divergence penalty to $\beta = 0$. For all α -DPG training, the clipping value is set to $M = 10$ and the partition function is constrained on the lower side as $Z_x \geq \epsilon$ with $\epsilon = 1e^{-4}$. For all problems in the training set, we pre-compute it as the average verifier score of 128 samples from the base model (see App. E for more details), setting it to ϵ if the problem was not solved. We fix the maximum response length to 1024 tokens. We use this last value as the normalization constant in the policy loss computation. For all α -DPG runs we set batch size to $b_z = 128$ and rollout size to $K = 4$ for total batch size of 512 sequences (which is the same for all baselines) and train for 200 iterations (≈ 3 epochs).

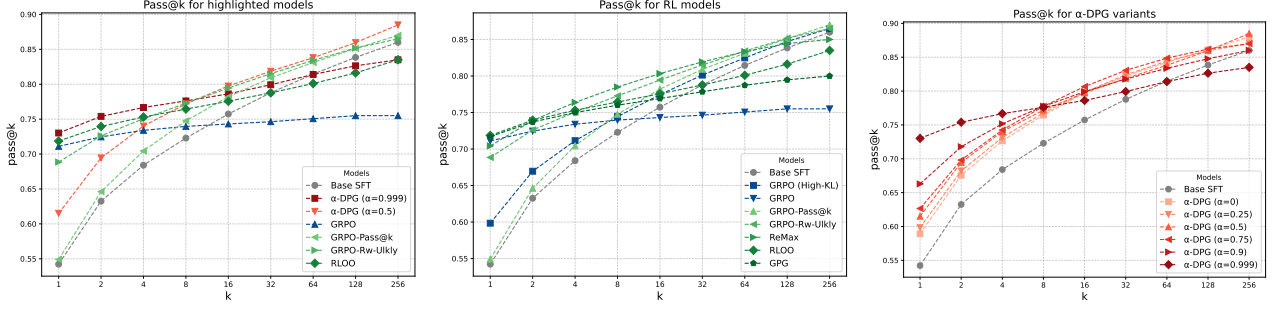


Figure 2: Pass@k curves on the test set for the Base-SFT model tuned with different methods.

Evaluation Metric To evaluate model performance on reasoning and generation tasks, we report **pass@k** [6], a standard metric widely used in code generation and problem-solving benchmarks. The metric measures the probability that at least one correct solution is found within the top- k sampled outputs of the model. For a problem with n generated outputs with $n \geq k$, of which c are correct, an unbiased estimate of the probability that at least one of k sampled outputs is correct is [6]

$$\text{pass@k} = 1 - \frac{\binom{n-c}{k}}{\binom{n}{k}}.$$

Averaging over problems gives the overall pass@k.

4.1. Results

Coverage vs. precision analysis We analyze models in terms of their capacity to cover a broad set of problems versus their precision at producing correct solutions. Coverage is captured by the pass@256 metric, while precision is reflected in pass@1. Results are summarized in Fig. 1. Models trained with GRPO and GPG achieve high pass@1, indicating high precision on solvable problems, but exhibit the lowest pass@256, suggesting limited overall coverage. RLOO and α -DPG ($\alpha = 0.999$), achieve comparable or better results in terms of pass@1, but better pass@256, thus being able to solve a wider spectrum of problems with comparable efficiency. In contrast, the base SFT model shows poor precision but can reach a wider range of problems. Methods such as Pass@k training, or GRPO with KL regularization improve in relation to the base model. α -DPG ($\alpha = 0.5$) achieves the highest coverage. Models trained with $\alpha \geq 0.5$ trace out a Pareto frontier, meaning that these models achieve the optimal trade-off between precision and coverage. The only other models that sit along this frontier are GRPO-Rw-Ulkly, ReMax and RLOO.

Pass@k curves Figure 2 illustrates the pass@k performance on the test set measured over $n = 256$ samples. First, looking at the left panel, we note that we have reproduced the results reported by He et al. [19], where the GRPO model starts with a much higher pass@1 score, but then the base model overpasses it as it reaches pass@16. Furthermore, we can see the performance of α -DPG with different values of α . As we can see, higher values of α give better pass@1 performance at the cost of lower pass@k. Interestingly, α -DPG with $\alpha = 0.999$ dominates completely the GRPO model, while $\alpha = 0.5$ dominates the base model by a significant margin, achieving the best performance on pass@256. Other baselines such as GRPO-Rw-Ulkly and also dominate the base model but by a much smaller margin in pass@256. RLOO, on the other hand, behaves very similar to $\alpha = 0.999$. The middle panel presents a more detailed comparison of the RL-based baselines, showing that KL regularization help prevent mode collapse, and that ReMax is a very competitive baseline. Also Pass@k training and Rw-Ulkly baselines offer competitive results. The rightmost panel is comparing different variants of α -DPG in relation to the base model. While the base model surpasses and matches $\alpha = 0.999$ and $\alpha = 0.9$, respectively, models with lower values of alpha dominate the base model across all values of k .

Problem difficulty analysis In this section, we examine how training affects problem solvability. We categorize problem difficulty based on model performance, measured as the proportion of correctly solved sequences. A problem is considered easy for a given model if at least 80% of the sampled sequences are correct, medium if

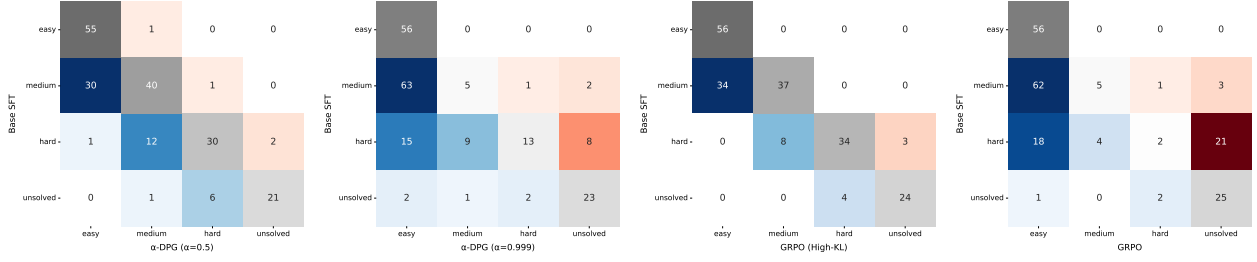
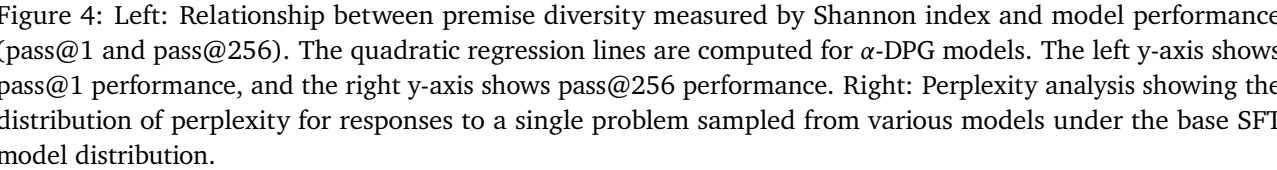


Figure 3: Problem Difficulty Transition Matrix from the Base-SFT to GRPO. The matrix shows the number of problems that transition from an initial difficulty classification under the base model (Base-SFT) (y-axis) to a final classification after post-training (x-axis). The results highlight a polarizing effect: α -DPG ($\alpha = 0.99$) and GRPO exhibit similar behavior, improving performance on a majority of medium-difficulty problems by making them easy, but also degrading performance on hard problems, causing nearly a half of them to become unsolved. α -DPG ($\alpha = 0.5$) and GRPO (High-KL) are more conservative, improving sample efficiency on fewer problems but harder problems remain solvable.

20–80% are correct, and hard if fewer than 20% are correct. This notion of difficulty has a direct relation with the efficiency of the model at solving a given problem as the number of samples until generating one correct solution follows a geometric distribution whose parameter is the model’s sampling accuracy. In Figure 3 we plot how the problems difficulties evolve after having trained the Base SFT model using various methods. Problems on the diagonal (grey) remain unaffected by training. Elements in the lower-left triangle (blue) represent problems for which solving efficiency improved, while elements in the upper-right triangle (red) indicate problems where sampling efficiency decreased. As we can see, many problems that were medium or hard for the base model became easy after training both for GRPO and $\alpha = 0.999$. However, this came at the cost of other problems in these same categories becoming unsolvable given the same sampling budget. α -DGP ($\alpha = 0.5$) and GRPO (High-KL) on the other hand, improve sample efficiency on fewer problems at the cost of just two or three problems becoming unsolvable.

Diversity Analysis In this section, we investigate the relationship between proof diversity and model performance. We decompose this analysis into two components: the *tactics* and *premises* used in candidate LEAN proofs. A *tactic* is a command that transforms a proof goal into simpler subgoals (e.g., `intro`, `apply`, `rw`), while a *premise* is a lemma or previously proven theorem that can be used within a proof (e.g., `mul_comm`, `mul_assoc`). As measures of diversity, we employ the Shannon index and the Gini-Simpson index. For each problem statement, we evaluate 256 generated proof sequences. At every proof state, we compute both the Simpson index and Shannon entropy over the choices of tactics and premises. Concretely, for a given problem, we count the occurrences of each premise and tactic, and compute the Simpson index as $D = 1 - \sum_{i=1}^S p_i^2$ and the Shannon index as $H = - \sum_{i=1}^S p_i \ln p_i$ where p_i is the relative abundance of premise or tactic i , and S is the total number of premises or tactics. These metrics are then aggregated across all problems to capture the overall diversity of candidate sequences. Higher diversity in tactics and premises in candidate proofs generally correlates with improved `pass@256` performance, whereas it is anticorrelated with `pass@1` as shown in the left panel of Figure 4 (and additionally, in Appendix Figure 14).

Perplexity Analysis Recent work on RLVR [67] shows that RL-trained models do not truly discover new solutions, instead the solutions they generate are already likely under the base model. To further investigate whether α -DPG stays close to the base model, we conduct a perplexity analysis. We sample a single problem from the test set and have each model generate 16 solutions. For each solution, we compute perplexity both under the model that generated it (self-perplexity) and under the base model (Base SFT). As shown in the right panel of Figure 4, across all models, the generated sequences are already highly probable according to the base model, with very similar perplexities. Note also that for this particular problem, GRPO collapsed and produces 16 identical sequences, which were also highly probable under the base model.



Improving test-time scaling (pass@k) Popular RL-based post-training methods, such as PPO [49] or GRPO [50], optimize inherently mode-seeking objectives, which can lead to mode collapse. The resulting models often achieve high accuracy, but exhibit very low entropy, generating less diverse sequences [68]. This issue is particularly pronounced when scaling test-time compute for solution search, with the SFT models often achieving better performances due to higher diversity [7, 19, 55, 67, 70]. Multiple approaches try to overcome this issue, mainly by adapting the advantage function. [19] proposed to add a rank bias penalty, that increases the advantage of unlikely sequences. [8] modify the reward from pass@1 equivalent to pass@k, allowing better sampling efficiency. They obtain a better pass@k at inference than training with entropy regularization. Tang et al. [55] propose a leave one out strategy to compute the advantage function to reduce variance and therefore improving pass@k at test-time.

6. Final Remarks and Conclusions

However, the core principle of filtering the base model is sound and can be independently motivated; it enforces correctness while preserving the multiplicity of valid responses. The true source of diversity loss, therefore, lies not in the target distribution itself, but in the divergence used to approximate it with gradient descent within a restricted parametric family.

By explicitly defining the target distribution, DMVR enables optimization with divergences that balance the two competing goals: correctness and diversity. In particular, we explored α -divergences, which smoothly interpolate between Forward and Reverse KL. This approach generates a Pareto frontier of models: setting α near the Reverse KL recovers models that match or exceed the performance of RL-based alternatives, while intermediate values of α produce models that preserve substantial diversity while still improving sampling precision over the base model.

The choice of divergence impacts the resulting models in at least two different ways: On one hand, different loss landscapes associated with different divergences can produce different training dynamics. On the other hand, within a restricted parametric family, each divergence can induce different optima. To disentangle the contribution that each of these factors has in the resulting models we could adopt a *curriculum* in which we gradually increase α during training, thus encouraging coverage at the beginning of training, and precision towards the end. If the resulting models preserve more diversity, this could be an indication that training dynamics matter. If, on the contrary, they reach as much coverage as using a high value of α from the start, then this would be an indication that it is the optima for the given parametric family that dominate. We will tackle these questions in future work.

Known Limitations Without clipping, α -DPG is unstable for small values of α (e.g., ≤ 0.5). Khalifa et al. [24] use an offline version of DPG where the sampling policy is only updated from the training policy if the divergence to the target probability is estimated to have improved. Here, we avoided keeping in memory a second copy of the model, and preferred to manage variance by clipping the pseudo rewards. Another point of caution is that our method relies on importance sampling weights, such as p/π ratios. For relatively long sequences, these ratios could suffer increased variance when working with the standard bfloat16. However, recent work suggests that this could have a simple solution [46]. We will dig deeper into these issues in future work.

Ethics statement

LLMs tuned using verifier feedback have attracted considerable attention from the research community and the general public because of their strong problem-solving abilities. The techniques introduced in this paper aim to preserve more of a model’s initial diversity than existing RL-based approaches at the cost of less strict adherence to the verifier’s feedback.

In the specific setting we study—training models to prove theorems—this trade-off carries little risk of harm. In broader applications, however, such as attempts to “align” language models with specific policies or behaviors, weaker adherence to constraints could be more concerning. We note that recent proposals within the distributional matching framework may help address this issue [25], even though they rely on the availability of a verifier at inference time.

More importantly, the distributional perspective makes explicit the central choice of a target distribution to optimize. Although the training techniques used for approximation also influence model behavior, we believe that making the choice of target distribution open and transparent can promote accountability and clarity. We therefore encourage this practice when tuning models for specific goals.

Reproducibility statement

We will make publicly available all the code to reproduce our experiments after publication. The LEAN Workbook dataset we used in our experiments is already open [62,64], as so it is the base model DeepSeek-Prover-V1.5-SFT we use for training [63]. In addition to the experiment details we describe in Section 4, we report additional details such as hyperparameter choices in App. F.

Acknowledgements

We thank Thibaut Thonet for the very helpful comments on this paper and anonymous reviewers for suggestions that helped improve this work.

References

- [1] Arash Ahmadian, Chris Cremer, Matthias Gallé, Marzieh Fadaee, Julia Kreutzer, Olivier Pietquin, Ahmet Üstün, and Sara Hooker. Back to Basics: Revisiting REINFORCE Style Optimization for Learning from Human Feedback in LLMs, February 2024. URL <http://arxiv.org/abs/2402.14740>. arXiv:2402.14740 [cs]. 3, 5, 6
- [2] AlphaProof and AlphaGeometry teams. AI achieves silver-medal standard solving International Mathematical Olympiad problems. Blog post, July 2024. URL: <https://deepmind.google/discover/blog/ai-solves-imo-problems-at-silver-medal-level/>. 9
- [3] Shun-ichi Amari. α -Divergence and α -Projection in Statistical Manifold. In Shun-ichi Amari (ed.), *Differential-Geometrical Methods in Statistics*, pp. 66–103. Springer, New York, NY, 1985. ISBN 978-1-4612-5056-2. doi: 10.1007/978-1-4612-5056-2_3. 1, 5
- [4] Bruno Barras, Samuel Boutin, Cristina Cornes, Judicaël Courant, Jean-Christophe Filliâtre, Eduardo Giménez, Hugo Herbelin, Gérard Huet, César Muñoz, Chetan Murthy, Catherine Parent, Christine Paulin-Mohring, Amokrane Saïbi, and Benjamin Werner. The Coq Proof Assistant Reference Manual : Version 6.1. Research Report RT-0203, INRIA, May 1997. URL <https://inria.hal.science/inria-00069968>. Projet COQ. 6, 9
- [5] Christopher M. Bishop. *Pattern recognition and machine learning*. Information science and statistics. Springer, New York, 2006. ISBN 978-0-387-31073-2. 1, 4
- [6] Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Pondé de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Joshua Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. Evaluating large language models trained on code. *CoRR*, abs/2107.03374, 2021. URL <http://dblp.uni-trier.de/db/journals/corr/corr2107.html>. 7
- [7] Zhipeng Chen, Xiaobo Qin, Youbin Wu, Yue Ling, Qinghao Ye, Wayne Xin Zhao, and Guang Shi. Pass@k training for adaptively balancing exploration and exploitation of large reasoning models. *arXiv preprint arXiv:2508.10751*, 2025. 9
- [8] Zhipeng Chen, Xiaobo Qin, Youbin Wu, Yue Ling, Qinghao Ye, Wayne Xin Zhao, and Guang Shi. Pass@k Training for Adaptively Balancing Exploration and Exploitation of Large Reasoning Models, August 2025. URL <http://arxiv.org/abs/2508.10751>. arXiv:2508.10751 [cs] version: 1. 6, 9
- [9] Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. Deep Reinforcement Learning from Human Preferences. In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. URL https://proceedings.neurips.cc/paper_files/paper/2017/hash/d5e2c0adad503c91f91df240d0cd4e49-Abstract.html. 1, 2
- [10] Xiangxiang Chu, Hailang Huang, Xiao Zhang, Fei Wei, and Yong Wang. GPG: A Simple and Strong Reinforcement Learning Baseline for Model Reasoning, May 2025. URL <http://arxiv.org/abs/2504.02546>. arXiv:2504.02546 [cs]. 6
- [11] Andrzej Cichocki and Shun-ichi Amari. Families of alpha-, beta- and gamma-divergences: Flexible and robust measures of similarities. *Entropy*, 12(6):1532–1568, 2010. doi: 10.3390/e12061532. URL <https://www.mdpi.com/1099-4300/12/6/1532>. 25
- [12] Noel Cressie and Timothy R. C. Read. Multinomial Goodness-of-Fit Tests. *Journal of the Royal Statistical Society. Series B (Methodological)*, 46(3):440–464, 1984. ISSN 0035-9246. URL <https://www.jstor.org/stable/2345686>. Publisher: [Royal Statistical Society, Oxford University Press]. 1, 5
- [13] I. Csiszár and P. C. Shields. Information Theory and Statistics: A Tutorial. *Foundations and Trends® in Communications and Information Theory*, 1(4):417–528, December 2004. ISSN 1567-2190, 1567-2328. doi: 10.1561/01000000004. URL <https://www.nowpublishers.com/article/Details/CIT-004>. Publisher: Now Publishers, Inc. 3

- [14] Xingyu Dang, Christina Baek, J. Zico Kolter, and Aditi Raghunathan. Assessing Diversity Collapse in Reasoning. February 2025. URL <https://openreview.net/forum?id=AMiKsHLjQh>. 1
- [15] Leonardo Mendonça de Moura, Soonho Kong, Jeremy Avigad, Floris van Doorn, and Jakob von Raumer. The lean theorem prover (system description). In *CADE*, 2015. URL <https://api.semanticscholar.org/CorpusID:232990>. 6, 9
- [16] DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jiashi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaoju Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiushi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng Zhou, Shuting Pan, S. S. Li, Shuang Zhou, Shaoqing Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun, T. Wang, Wangding Zeng, Wanjia Zhao, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin, Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxiang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang, Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yuan Ou, Yuduan Wang, Yue Gong, Yuheng Zou, Yujia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu, Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu, Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu Zhang, and Zhen Zhang. DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning, January 2025. URL <http://arxiv.org/abs/2501.12948>. arXiv:2501.12948 [cs]. 1, 3, 6, 9
- [17] Gemini team. Advanced version of gemini with deep think officially achieves gold medal standard at the international mathematical olympiad. Blog post, July 2025. URL: <https://deepmind.google/discover/blog/advanced-version-of-gemini-with-deep-think-officially-achieves-gold-medal-standard-at-the-international-mathematical-olympiad>. 5
- [18] Dongyoung Go, Tomasz Korbak, Germán Kruszewski, Jos Rozen, Nahyeon Ryu, and Marc Dymetman. Aligning Language Models with Preferences through f-divergence Minimization, June 2023. URL <http://arxiv.org/abs/2302.08215>. arXiv:2302.08215 [cs]. 1, 3, 4
- [19] Andre He, Daniel Fried, and Sean Welleck. Rewarding the Unlikely: Lifting GRPO Beyond Distribution Sharpening, June 2025. URL <http://arxiv.org/abs/2506.02355>. arXiv:2506.02355 [cs]. 1, 6, 7, 9
- [20] E. Hellinger. Neue Begründung der Theorie quadratischer Formen von unendlichvielen Veränderlichen. *Journal für die reine und angewandte Mathematik*, 1909(136):210–271, July 1909. ISSN 1435-5345, 0075-4102. doi: 10.1515/crll.1909.136.210. URL <https://www.degruyter.com/document/doi/10.1515/crll.1909.136.210/html>. 5
- [21] Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. Measuring Mathematical Problem Solving With the MATH Dataset, November 2021. URL <http://arxiv.org/abs/2103.03874>. arXiv:2103.03874 [cs]. 5, 19
- [22] Ferenc Huszár. How (not) to Train your Generative Model: Scheduled Sampling, Likelihood, Adversary?, November 2015. URL <http://arxiv.org/abs/1511.05101>. arXiv:1511.05101 [stat]. 1, 4
- [23] B. Kappen, V. Gomez, and M. Opper. Optimal control as a graphical model inference problem. *Machine Learning*, 87(2):159–182, May 2012. ISSN 0885-6125, 1573-0565. doi: 10.1007/s10994-012-5278-7. URL <http://arxiv.org/abs/0901.0633>. arXiv:0901.0633 [math]. 3

- [24] Muhammad Khalifa, Hady Elsahar, and Marc Dymetman. A Distributional Approach to Controlled Text Generation, May 2021. URL <http://arxiv.org/abs/2012.11635>. arXiv:2012.11635 [cs]. 1, 2, 3, 4, 5, 10
- [25] Minbeom Kim, Thibaut Thonet, Jos Rozen, Hwaran Lee, Kyomin Jung, and Marc Dymetman. Guaranteed Generation from Large Language Models, July 2025. URL <http://arxiv.org/abs/2410.06716>. arXiv:2410.06716 [cs]. 1, 3, 10
- [26] Robert Kirk, Ishita Mediratta, Christoforos Nalmpantis, Jelena Luketina, Eric Hambro, and Edward Grefenstette. Understanding the Effects of RLHF on LLM Generalisation and Diversity. 2024. 1
- [27] Wouter Kool. Buy 4 REINFORCE Samples, Get a Baseline for Free!, July 2019. URL <https://wouterkool.github.io/publication/buy-4-samples-free-baseline/>. 3, 5
- [28] Tomasz Korbak, Hady Elsahar, German Kruszewski, and Marc Dymetman. Controlling Conditional Language Models without Catastrophic Forgetting, June 2022. URL <http://arxiv.org/abs/2112.00791>. arXiv:2112.00791 [cs]. 2, 3
- [29] Tomasz Korbak, Hady Elsahar, Germán Kruszewski, and Marc Dymetman. On Reinforcement Learning and Distribution Matching for Fine-Tuning Language Models with no Catastrophic Forgetting, November 2022. URL <http://arxiv.org/abs/2206.00761>. arXiv:2206.00761 [cs]. 1, 2, 3, 4, 9
- [30] Nathan Lambert, Jacob Morrison, Valentina Pyatkin, Shengyi Huang, Hamish Ivison, Faeze Brahman, Lester James V. Miranda, Alisa Liu, Nouha Dziri, Shane Lyu, Yuling Gu, Saumya Malik, Victoria Graf, Jena D. Hwang, Jiangjiang Yang, Ronan Le Bras, Oyvind Tafjord, Chris Wilhelm, Luca Soldaini, Noah A. Smith, Yizhong Wang, Pradeep Dasigi, and Hannaneh Hajishirzi. Tulu 3: Pushing Frontiers in Open Language Model Post-Training, April 2025. URL <http://arxiv.org/abs/2411.15124>. arXiv:2411.15124 [cs]. 1
- [31] Guillaume Lample, Timothee Lacroix, Marie anne Lachaux, Aurelien Rodriguez, Amaury Hayat, Thibaut Lavril, Gabriel Ebner, and Xavier Martinet. Hypertree proof search for neural theorem proving. In Alice H. Oh, Alekh Agarwal, Danielle Belgrave, and Kyunghyun Cho (eds.), *Advances in Neural Information Processing Systems*, 2022. URL <https://openreview.net/forum?id=J4pX8Q8cxHH>. 9
- [32] Aitor Lewkowycz, Anders Andreassen, David Dohan, Ethan Dyer, Henryk Michalewski, Vinay Ramasesh, Ambrose Slone, Cem Anil, Imanol Schlag, Theo Gutman-Solo, Yuhuai Wu, Behnam Neyshabur, Guy Gur-Ari, and Vedant Misra. Solving Quantitative Reasoning Problems With Language Models. NIPS’22: Proceedings of the 36th International Conference on Neural Information Processing Systems, 2022. 20
- [33] Cheuk Ting Li and Farzan Farnia. Mode-Seeking Divergences: Theory and Applications to GANs. In *Proceedings of The 26th International Conference on Artificial Intelligence and Statistics*, pp. 8321–8350. PMLR, April 2023. URL <https://proceedings.mlr.press/v206/ting-li23a.html>. ISSN: 2640-3498. 1
- [34] Long Li, Jiaran Hao, Jason Klein Liu, Zhijian Zhou, Xiaoyu Tan, Wei Chu, Zhe Wang, Shirui Pan, Chao Qu, and Yuan Qi. The Choice of Divergence: A Neglected Key to Mitigating Diversity Collapse in Reinforcement Learning with Verifiable Reward, September 2025. URL <http://arxiv.org/abs/2509.07430>. arXiv:2509.07430 [cs] version: 1. 4
- [35] Yingzhen Li and Yarin Gal. Dropout inference in bayesian neural networks with alpha-divergences, 2017. URL <https://arxiv.org/abs/1703.02914>. 26
- [36] 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, May 2024. URL <http://arxiv.org/abs/2310.10505>. arXiv:2310.10505 [cs]. 6
- [37] Friedrich Liese and Igor Vajda. On divergences and informations in statistics and information theory. *IEEE Transactions on Information Theory*, 52(10):4394–4412, 2006. doi: 10.1109/TIT.2006.881731. 26
- [38] Zichen Liu, Changyu Chen, Wenjun Li, Penghui Qi, Tianyu Pang, Chao Du, Wee Sun Lee, and Min Lin. Understanding r1-zero-like training: A critical perspective. In *Conference on Language Modeling (COLM)*, 2025. 2, 6
- [39] Mistral-AI, Abhinav Rastogi, Albert Q. Jiang, Andy Lo, Gabrielle Berrada, Guillaume Lample, Jason Rute, Joep Barmantlo, Karmesh Yadav, Kartik Khandelwal, Khyathi Raghavi Chandu, Léonard Blier, Lucile Saulnier, Matthieu Dinot, Maxime Darrin, Neha Gupta, Roman Soletskyi, Sagar Vaze, Teven Le Scao, Yihan Wang, Adam Yang, Alexander H. Liu, Alexandre Sablayrolles, Amélie Héliou, Amélie Martin, Andy Ehrenberg, Anmol Agarwal, Antoine Roux, Arthur Darcet, Arthur Mensch, Baptiste Bout, Baptiste Rozière, Baudouin De Monicault, Chris Bamford, Christian Wallenwein, Christophe Renaudin, Clémence Lanfranchi, Darius Dabert,

- Devon Mizelle, Diego de las Casas, Elliot Chane-Sane, Emilien Fugier, Emma Bou Hanna, Gauthier Delerce, Gauthier Guinet, Georgii Novikov, Guillaume Martin, Himanshu Jaju, Jan Ludziejewski, Jean-Hadrien Chabran, Jean-Malo Delignon, Joachim Studnia, Jonas Amar, Josselin Somerville Roberts, Julien Denize, Karan Saxena, Kush Jain, Lingxiao Zhao, Louis Martin, Luyu Gao, L  lio Renard Lavaud, Marie Pellat, Mathilde Guillaumin, Mathis Felardos, Maximilian Augustin, Micka  l Seznec, Nikhil Raghuraman, Olivier Duchenne, Patricia Wang, Patrick von Platen, Patryk Saffer, Paul Jacob, Paul Wambergue, Paula Kurylowicz, Pavankumar Reddy Muddireddy, Philom  ne Chagniot, Pierre Stock, Pravesh Agrawal, Romain Sauvestre, R  mi Delacourt, Sanchit Gandhi, Sandeep Subramanian, Shashwat Dalal, Siddharth Gandhi, Soham Ghosh, Srijan Mishra, Sumukh Aithal, Szymon Antoniak, Thibault Schueller, Thibaut Lavril, Thomas Robert, Thomas Wang, Timoth  e Lacroix, Valeriia Nemychnikova, Victor Paltz, Virgile Richard, Wen-Ding Li, William Marshall, Xuanyu Zhang, and Yunhao Tang. Magistral, June 2025. URL <http://arxiv.org/abs/2506.10910>. arXiv:2506.10910 [cs]. 2, 6
- [40] Tobias Nipkow, Markus Wenzel, and Lawrence C. Paulson. *Isabelle/HOL: a proof assistant for higher-order logic*. Springer-Verlag, Berlin, Heidelberg, 2002. ISBN 3540433767. 6, 9
- [41] Art B Owen. Importance Sampling. In *Monte Carlo theory, methods and examples*. Chapter 9, 2013. URL <https://statweb.stanford.edu/~Eowen/mc/Ch-var-is.pdf>. 3, 18
- [42] Laura O’Mahony, Leo Grinsztajn, Hailey Schoelkopf, and Stella Biderman. ATTRIBUTING MODE COLLAPSE IN THE FINE-TUNING OF LARGE LANGUAGE MODELS. 2024. 1
- [43] Tetiana Parshakova, Jean-Marc Andreoli, and Marc Dymetman. Distributional Reinforcement Learning for Energy-Based Sequential Models, December 2019. URL <http://arxiv.org/abs/1912.08517>. arXiv:1912.08517 [cs]. 1, 3, 5
- [44] Stanislas Polu, Jesse Michael Han, Kunhao Zheng, Mantas Baksys, Igor Babuschkin, and Ilya Sutskever. Formal mathematics statement curriculum learning. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=-P7G-8dmSh4>. 9
- [45] Yuri Polyanskiy and Yihong Wu. *Information Theory: From Coding to Learning*. Cambridge University Press, 2025. ISBN 978-1108832908. 1, 27
- [46] Penghui Qi, Zichen Liu, Xiangxin Zhou, Tianyu Pang, Chao Du, Wee Sun Lee, and Min Lin. Defeating the Training-Inference Mismatch via FP16, October 2025. URL <http://arxiv.org/abs/2510.26788>. arXiv:2510.26788 [cs]. 10, 21
- [47] Z. Z. Ren, Zhihong Shao, Junxiao Song, Huajian Xin, Haocheng Wang, Wanbiao Zhao, Liye Zhang, Zhe Fu, Qihao Zhu, Dejian Yang, Z. F. Wu, Zhibin Gou, Shirong Ma, Hongxuan Tang, Yuxuan Liu, Wenjun Gao, Daya Guo, and Chong Ruan. Deepseek-prover-v2: Advancing formal mathematical reasoning via reinforcement learning for subgoal decomposition, 2025. URL <https://arxiv.org/abs/2504.21801>. 6, 9
- [48] Alfr  d R  nyi. On Measures of Entropy and Information. In *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability, Volume 1: Contributions to the Theory of Statistics*, volume 4.1, pp. 547–562. University of California Press, January 1961. URL <https://projecteuclid.org/ebooks/berkeley-symposium-on-mathematical-statistics-and-probability/Proceedings-of-the-Fourth-Berkeley-Symposium-on-Mathematical-Statistics-and/chapter/On-Measures-of-Entropy-and-Information/bsmsp/1200512181>. 1, 5
- [49] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal Policy Optimization Algorithms, August 2017. URL <http://arxiv.org/abs/1707.06347>. arXiv:1707.06347 [cs]. 1, 3, 9
- [50] Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. DeepSeekMath: Pushing the Limits of Mathematical Reasoning in Open Language Models, April 2024. URL <http://arxiv.org/abs/2402.03300>. arXiv:2402.03300 [cs]. 1, 9
- [51] Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. Hybridflow: A flexible and efficient rlhf framework. *arXiv preprint arXiv: 2409.19256*, 2024. 6
- [52] Haoxiang Sun, Yingqian Min, Zhipeng Chen, Wayne Xin Zhao, Lei Fang, Zheng Liu, Zhongyuan Wang, and Ji-Rong Wen. Challenging the Boundaries of Reasoning: An Olympiad-Level Math Benchmark for Large Language Models, May 2025. URL <http://arxiv.org/abs/2503.21380>. arXiv:2503.21380 [cs] version: 2. 5
- [53] Richard S. Sutton and Andrew Barto. *Reinforcement learning: an introduction*. Adaptive computation and

- machine learning. The MIT Press, Cambridge, Massachusetts London, England, second edition edition, 2020. ISBN 978-0-262-03924-6. 3
- [54] Richard S Sutton, David McAllester, Satinder Singh, and Yishay Mansour. Policy Gradient Methods for Reinforcement Learning with Function Approximation. In *Advances in Neural Information Processing Systems*, volume 12. MIT Press, 1999. URL https://proceedings.neurips.cc/paper_files/paper/1999/hash/464d828b85b0bed98e80ade0a5c43b0f-Abstract.html. 3
- [55] Yunhao Tang, Kunhao Zheng, Gabriel Synnaeve, and Remi Munos. Optimizing language models for inference time objectives using reinforcement learning. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=ZVWJ05YTz4>. 6, 9
- [56] Terence Tao. Machine-assisted proof. *Notices of the American Mathematical Society*, 72:1, 01 2025. doi: 10.1090/noti3041. 6
- [57] Emanuel Todorov. Linearly-solvable Markov decision problems. In *Advances in Neural Information Processing Systems*, volume 19. MIT Press, 2006. URL https://papers.nips.cc/paper_files/paper/2006/hash/d806ca13ca3449af72a1ea5aedbed26a-Abstract.html. 3
- [58] Haiming Wang, Mert Unsal, Xiaohan Lin, Mantas Baksys, Junqi Liu, Marco Dos Santos, Flood Sung, Marina Vinyes, Zhenzhe Ying, Zekai Zhu, Jianqiao Lu, Hugues de Saxcé, Bolton Bailey, Chendong Song, Chenjun Xiao, Dehao Zhang, Ebony Zhang, Frederick Pu, Han Zhu, Jiawei Liu, Jonas Bayer, Julien Michel, Longhui Yu, Léo Dreyfus-Schmidt, Lewis Tunstall, Luigi Pagani, Moreira Machado, Pauline Bourigault, Ran Wang, Stanislas Polu, Thibaut Barroyer, Wen-Ding Li, Yazhe Niu, Yann Fleureau, Yangyang Hu, Zhouliang Yu, Zihan Wang, Zhilin Yang, Zhengying Liu, and Jia Li. Kimina-prover preview: Towards large formal reasoning models with reinforcement learning, 2025. URL <https://arxiv.org/abs/2504.11354>. 6, 9
- [59] Haiming Wang, Mert Unsal, Xiaohan Lin, Mantas Baksys, Junqi Liu, Marco Dos Santos, Flood Sung, Marina Vinyes, Zhenzhe Ying, Zekai Zhu, Jianqiao Lu, Hugues de Saxcé, Bolton Bailey, Chendong Song, Chenjun Xiao, Dehao Zhang, Ebony Zhang, Frederick Pu, Han Zhu, Jiawei Liu, Jonas Bayer, Julien Michel, Longhui Yu, Léo Dreyfus-Schmidt, Lewis Tunstall, Luigi Pagani, Moreira Machado, Pauline Bourigault, Ran Wang, Stanislas Polu, Thibaut Barroyer, Wen-Ding Li, Yazhe Niu, Yann Fleureau, Yangyang Hu, Zhouliang Yu, Zihan Wang, Zhilin Yang, Zhengying Liu, and Jia Li. Kimina-Prover Preview: Towards Large Formal Reasoning Models with Reinforcement Learning, April 2025. URL <http://arxiv.org/abs/2504.11354>. arXiv:2504.11354 [cs]. 21
- [60] Ronald J. Williams. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine Learning*, 8(3):229–256, May 1992. ISSN 1573-0565. doi: 10.1007/BF00992696. URL <https://doi.org/10.1007/BF00992696>. 3
- [61] Fang Wu, Weihao Xuan, Ximing Lu, Mingjie Liu, Yi Dong, Zaid Harchaoui, and Yejin Choi. The Invisible Leash: Why RLVR May or May Not Escape Its Origin, September 2025. URL <http://arxiv.org/abs/2507.14843>. arXiv:2507.14843 [cs]. 1, 9
- [62] Zijian Wu, Suozhi Huang, Zhejian Zhou, Huaiyuan Ying, Jiayu Wang, Dahua Lin, and Kai Chen. Internlm2.5-stepprover: Advancing automated theorem proving via expert iteration on large-scale lean problems, 2024. URL <https://arxiv.org/abs/2410.15700>. 6, 9, 10
- [63] Huajian Xin, Z.Z. Ren, Junxiao Song, Zhihong Shao, Wanxia Zhao, Haocheng Wang, Bo Liu, Liyue Zhang, Xuan Lu, Qiushi Du, Wenjun Gao, Haowei Zhang, Qihao Zhu, Dejian Yang, Zhibin Gou, Z.F. Wu, Fuli Luo, and Chong Ruan. Deepseek-prover-v1.5: Harnessing proof assistant feedback for reinforcement learning and monte-carlo tree search. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=I4YAIwrsXa>. 6, 9, 10
- [64] Huaiyuan Ying, Zijian Wu, Yihan Geng, Jiayu Wang, Dahua Lin, and Kai Chen. Lean workbook: A large-scale lean problem set formalized from natural language math problems. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2024. URL <https://openreview.net/forum?id=Vcw3vzjHDb>. 6, 10
- [65] Qiying Yu, Zheng Zhang, Ruofei Zhu, Yufeng Yuan, Xiaochen Zuo, Yu Yue, Tiantian Fan, Gaohong Liu, Lingjun Liu, Xin Liu, Haibin Lin, Zhiqi Lin, Bole Ma, Guangming Sheng, Yuxuan Tong, Chi Zhang, Mofan Zhang, Wang Zhang, Hang Zhu, Jinhua Zhu, Jiaze Chen, Jiangjie Chen, Chengyi Wang, Hongli Yu, Weinan Dai, Yuxuan Song, Xiangpeng Wei, Hao Zhou, Jingjing Liu, Wei-Ying Ma, Ya-Qin Zhang, Lin Yan, Mu Qiao, Yonghui Wu, and Mingxuan Wang. DAPO: An Open-Source LLM Reinforcement Learning System at Scale, March 2025. URL <http://arxiv.org/abs/2503.14476>. arXiv:2503.14476 [cs] version: 1. 2

- [66] 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, September 2023. URL <http://arxiv.org/abs/2308.01825>. arXiv:2308.01825 [cs]. 2, 5, 18
- [67] Yang Yue, Zhiqi Chen, Rui Lu, Andrew Zhao, Zhaokai Wang, Yang Yue, Shiji Song, and Gao Huang. Does reinforcement learning really incentivize reasoning capacity in llms beyond the base model? *arXiv preprint arXiv:2504.13837*, 2025. 8, 9
- [68] Yang Yue, Zhiqi Chen, Rui Lu, Andrew Zhao, Zhaokai Wang, Yang Yue, Shiji Song, and Gao Huang. Does Reinforcement Learning Really Incentivize Reasoning Capacity in LLMs Beyond the Base Model?, May 2025. URL <http://arxiv.org/abs/2504.13837>. arXiv:2504.13837 [cs]. 1, 9, 20
- [69] Eric Zelikman, Yuhuai Wu, Jesse Mu, and Noah D. Goodman. STaR: Bootstrapping Reasoning With Reasoning, May 2022. URL <http://arxiv.org/abs/2203.14465>. arXiv:2203.14465 [cs]. 2, 5, 18
- [70] Xinyu Zhu, Mengzhou Xia, Zhepei Wei, Wei-Lin Chen, Danqi Chen, and Yu Meng. The surprising effectiveness of negative reinforcement in llm reasoning. *arXiv preprint arXiv:2506.01347*, 2025. 9
- [71] Daniel M. Ziegler, Nisan Stiennon, Jeffrey Wu, Tom B. Brown, Alec Radford, Dario Amodei, Paul Christiano, and Geoffrey Irving. Fine-Tuning Language Models from Human Preferences, January 2020. URL <http://arxiv.org/abs/1909.08593>. arXiv:1909.08593 [cs]. 1, 2

A. LLM Usage

We have required the assistance from LLMs (ChatGPT, Gemini) in the production of this paper on many stages during its development, also during the rebuttal stage. The most prevalent usage has been on writing code and debugging, but we have also discussed research directions and mathematical statements, used them to simplify and format mathematical proofs, proof-reading of the paper and improving the flow of the text, and they even assisted us in producing the cartoon illustrating our method. The authors take full responsibility for the contents in this paper.

B. Proof of the correspondence between RLVR and Reverse KL

We work in the contextual setting where contexts x are drawn from $\mu(x)$, the policy is $\pi_\theta(y | x)$, the reward is $v(y, x)$, and the prior (reference policy) is $\pi_{\text{ref}}(y | x)$. Define for each context x and inverse temperature $\beta > 0$

$$p_{x,\beta}(y) = \frac{1}{Z_x(\beta)} \pi_{\text{ref}}(y | x) \exp(v(y, x)/\beta), \quad Z_x(\beta) = \sum_y \pi_{\text{ref}}(y | x) \exp(v(y, x)/\beta).$$

Fix a context x . The reverse KL divergence (from $\pi_\theta(\cdot | x)$ to $p_{x,\beta}$) is

$$\text{KL}(\pi_\theta(\cdot | x) \| p_{x,\beta}) = \mathbb{E}_{y \sim \pi_\theta(\cdot | x)} \left[\log \frac{\pi_\theta(y | x)}{p_{x,\beta}(y)} \right].$$

Substituting $p_{x,\beta}(y) = \frac{1}{Z_x(\beta)} \pi_{\text{ref}}(y | x) e^{v(y,x)/\beta}$ gives

$$\begin{aligned} \text{KL}(\pi_\theta \| p_{x,\beta}) &= \mathbb{E}_{y \sim \pi_\theta} \left[\log \pi_\theta(y | x) - \log \left(\frac{1}{Z_x(\beta)} \pi_{\text{ref}}(y | x) e^{v(y,x)/\beta} \right) \right] \\ &= \mathbb{E}_{y \sim \pi_\theta} \left[\log \frac{\pi_\theta(y | x)}{\pi_{\text{ref}}(y | x)} - \frac{1}{\beta} v(y, x) \right] + \log Z_x(\beta). \end{aligned}$$

Equivalently (rearranging signs),

$$\text{KL}(\pi_\theta \| p_{x,\beta}) = -\mathbb{E}_{y \sim \pi_\theta} \left[\frac{1}{\beta} v(y, x) - \log \frac{\pi_\theta(y | x)}{\pi_{\text{ref}}(y | x)} \right] + \log Z_x(\beta).$$

Now take the gradient with respect to θ and average over contexts $x \sim \mu$. Since $Z_x(\beta)$ depends only on π_{ref} and v (not on θ), $\nabla_\theta \log Z_x(\beta) = 0$. Thus

$$\nabla_\theta \mathbb{E}_{x \sim \mu} [\text{KL}(\pi_\theta \| p_{x,\beta})] = -\nabla_\theta \mathbb{E}_{x \sim \mu} \mathbb{E}_{y \sim \pi_\theta(\cdot | x)} \left[\frac{1}{\beta} v(y, x) - \log \frac{\pi_\theta(y | x)}{\pi_{\text{ref}}(y | x)} \right].$$

To proceed we use the score-function (REINFORCE) identity: for any function $h(y, x)$ independent of θ ,

$$\nabla_\theta \mathbb{E}_{y \sim \pi_\theta(\cdot | x)} [h(y, x)] = \mathbb{E}_{y \sim \pi_\theta(\cdot | x)} [h(y, x) \nabla_\theta \log \pi_\theta(y | x)].$$

Applying this identity (and noting the dependence of the $\log \pi_\theta$ term must be handled consistently) yields the explicit gradient form

$$\nabla_\theta \mathbb{E}_{x \sim \mu} [\text{KL}(\pi_\theta \| p_{x,\beta})] = -\mathbb{E}_{x \sim \mu} \mathbb{E}_{y \sim \pi_\theta(\cdot | x)} \left[\left(\frac{1}{\beta} v(y, x) - \log \frac{\pi_\theta(y | x)}{\pi_{\text{ref}}(y | x)} \right) \nabla_\theta \log \pi_\theta(y | x) \right]. \quad (18)$$

This is the standard policy-gradient / score-function expression for the gradient of the reverse-KL objective in the contextual case.

Finally, an equivalent compact form is obtained by moving the ∇_θ inside the expectation and noticing the factor $1/\beta$:

$$\nabla_\theta \mathbb{E}_x [\text{KL}(\pi_\theta \| p_{x,\beta})] = -\frac{1}{\beta} \nabla_\theta \mathbb{E}_{x \sim \mu, y \sim \pi_\theta} \left[v(y, x) - \beta \log \frac{\pi_\theta(y | x)}{\pi_{\text{ref}}(y | x)} \right].$$

□

Maximizing the expected reward under the conditional policy:

$$\mathbb{E}_{x, y \sim \pi_\theta} [v(y, x) - \beta \log \frac{\pi_\theta(y | x)}{\pi_{\text{ref}}(y | x)}]$$

is equivalent (up to constants) to minimizing:

$$\mathbb{E}_{x \sim \mu(x)} [D_{\text{KL}}(\pi_\theta(y | x) \| p_{x,\beta}(y | x))].$$

C. Proof that p_β becomes p as $\beta \rightarrow 0$

Proposition 1.

$$\lim_{\beta \rightarrow 0} p_{x,\beta} = p_x, \quad (19)$$

in the formal sense that $\|p_{x,\beta} - p_x\|_{\text{TV}} \rightarrow 0$ as $\beta \downarrow 0$. In particular $p_{x,\beta}(y) \rightarrow p_x(y)$ for every fixed y .

Proof. The proof is a direct adaptation, for a context x , of the following Lemma.

Lemma 3. Let π_{base} be a probability distribution on a countable set Y , let $r : Y \rightarrow \{0, 1\}$ and assume $Z := \sum_y \pi_{\text{base}}(y)r(y) \in (0, 1]$. Define

$$p(y) = \frac{\pi_{\text{base}}(y)r(y)}{Z}, \quad p_\beta(y) = \frac{\pi_{\text{base}}(y)e^{r(y)/\beta}}{Z_\beta}, \quad Z_\beta := \sum_y \pi_{\text{base}}(y)e^{r(y)/\beta}.$$

Then $\|p_\beta - p\|_{\text{TV}} \rightarrow 0$ as $\beta \downarrow 0$. In particular $p_\beta(y) \rightarrow p(y)$ for every fixed $y \in Y$.

Proof. Set $S_0 := \sum_{r(y)=0} \pi_{\text{base}}(y) = 1 - Z$ and $\varepsilon := e^{-1/\beta} S_0 = e^{-1/\beta} (1 - Z)$. A short computation gives $Z_\beta = e^{1/\beta} Z + S_0$ and hence, after multiplying numerator and denominator by $e^{-1/\beta}$,

$$p_\beta(y) = \begin{cases} \frac{\pi_{\text{base}}(y)}{Z + \varepsilon}, & r(y) = 1, \\ \frac{e^{-1/\beta} \pi_{\text{base}}(y)}{Z + \varepsilon}, & r(y) = 0. \end{cases}$$

Therefore

$$\sum_{r(y)=1} |p_\beta(y) - p(y)| = \frac{\varepsilon}{Z + \varepsilon}, \quad \sum_{r(y)=0} |p_\beta(y) - p(y)| = \frac{\varepsilon}{Z + \varepsilon},$$

so $\sum_y |p_\beta(y) - p(y)| = \frac{2\varepsilon}{Z + \varepsilon}$ and

$$\|p_\beta - p\|_{\text{TV}} = \frac{1}{2} \sum_y |p_\beta(y) - p(y)| = \frac{\varepsilon}{Z + \varepsilon} = \frac{e^{-1/\beta} (1 - Z)}{Z + e^{-1/\beta} (1 - Z)}.$$

Since $e^{-1/\beta} \rightarrow 0$ as $\beta \downarrow 0$, the right-hand side tends to 0, proving total-variation convergence. The pointwise convergence $p_\beta(y) \rightarrow p(y)$ follows immediately because $|p_\beta(y) - p(y)| \leq \sum_{y'} |p_\beta(y') - p(y')| = 2\|p_\beta - p\|_{\text{TV}}$. □

D. RS-FT optimizes the Forward KL

Rejection sampling fine tuning [66, 69] generates a sample set from the base model π_{base} , $\mathcal{S} \sim \pi_{\text{base}}(\cdot|x)$, filters them using the verifier to obtain $\mathcal{S}' = \{y_i : v(y_i, x) = 1, y_i \in \mathcal{S}\}$, and then trains the policy using standard cross-entropy on this filtered set:

$$\nabla_{\theta} \mathcal{L}^{\text{RS-FT}} = -\mathbb{E}_{y \sim \mathcal{S}'} \nabla_{\theta} \log \pi_{\theta}(y|x) \quad (20)$$

Now, starting from KL-DPG, we note that:

$$\nabla_{\theta} D_{\text{KL}}(p_x || \pi_{\theta}) = -\mathbb{E}_{y \sim \pi_{\theta}(\cdot|x)} \left(\frac{p_x(y)}{\pi_{\theta}(y|x)} - 1 \right) \nabla_{\theta} \log \pi_{\theta}(y|x). \quad (21)$$

$$= -\mathbb{E}_{y \sim \pi_{\theta}(\cdot|x)} \frac{p_x(y)}{\pi_{\theta}(y|x)} \nabla_{\theta} \log \pi_{\theta}(y|x). \quad (22)$$

$$= -\mathbb{E}_{y \sim \pi_{\theta}(\cdot|x)} \frac{\pi_{\text{base}}(y|x)}{\pi_{\text{base}}(y|x)} \frac{p_x(y)}{\pi_{\theta}(y|x)} \nabla_{\theta} \log \pi_{\theta}(y|x) \quad (23)$$

$$= -\mathbb{E}_{y \sim \pi_{\text{base}}(\cdot|x)} \frac{p_x(y)}{\pi_{\text{base}}(y|x)} \nabla_{\theta} \log \pi_{\theta}(y|x) \quad (24)$$

$$= -\frac{1}{Z_x} \mathbb{E}_{y \sim \pi_{\text{base}}(\cdot|x)} \frac{\pi_{\text{base}}(y|x) v(y, x)}{\pi_{\text{base}}(y|x)} \nabla_{\theta} \log \pi_{\theta}(y|x) \quad (25)$$

$$\propto \mathbb{E}_{y \sim \pi_{\text{base}}(\cdot|x)} v(y, x) \nabla_{\theta} \log \pi_{\theta}(y|x), \quad (26)$$

which exactly corresponds to the objective in Eq. 20. There are two crucial differences between the two: One is their sample efficiency: whereas RS-FT is limited to using samples from the base model, many of which may be rejected from the verifier, KL-DPG makes use of the updated policy which has a higher acceptance rate. The second one is that whereas RS-FT uses a finite pool of samples, and thus can over-fit to them, KL-DPG uses an unbounded number of samples.

E. Estimation of the partition function

Define

$$P_x(y) = \pi_{\text{base}}(y|x) v(y, x). \quad (27)$$

over a discrete space $\mathcal{Y}$, with partition function Z_x

$$Z_x = \sum_{y \in \mathcal{Y}} P_x(y). \quad (28)$$

Using importance sampling [41] we can estimate Z_x by using N samples $[y_i]_{i \in \{0 \dots N\}}$ generated from a proposal distribution q with $\text{Supp}(P) \subseteq \text{Supp}(q)$, as follows.

$$Z_x = \sum_{y \in \mathcal{Y}} P_x(y) = \sum_{y \in \mathcal{Y}} \frac{q(y|x)}{q(y|x)} P_x(y) = \mathbb{E}_{y \sim q(\cdot|x)} \frac{P_x(y)}{q(y|x)} \approx \frac{1}{N} \sum_i \frac{P_x(y_i)}{q(y_i|x)}. \quad (29)$$

Note that in the specific case that use $q = \pi_{\text{base}}$, then

$$Z_x = \mathbb{E}_{y \sim q(\cdot|x)} \frac{\pi_{\text{base}}(y|x) v(y, x)}{\pi_{\text{base}}(y|x)} = \mathbb{E}_{y \sim q(\cdot|x)} v(y, x) \approx \frac{1}{N} \sum_i v(y_i, x). \quad (30)$$

F. Additional Experimental Details and Hyperparameters

Method	Learning Rate	Batch Size	Rollout Size (N)	KL Penalty
Pass@k	1×10^{-6}	16	32	0.001
Rwrd. Unkly (rank pty=0.25)	2×10^{-6}	16	32	0.001
Dr. GRPO	2×10^{-6}	128	4	0.0
Dr. GRPO with High KL	2×10^{-6}	128	4	0.1
α -DPG	2×10^{-6}	128	4	-

Table 2: Summary of key hyper parameters for different training runs.

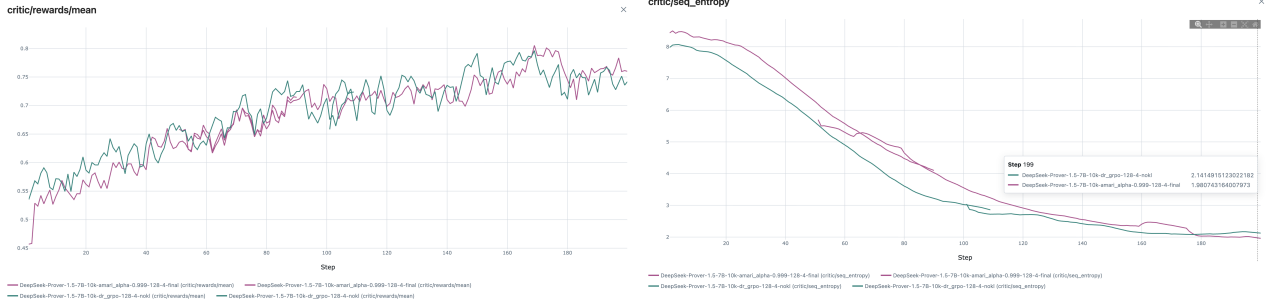
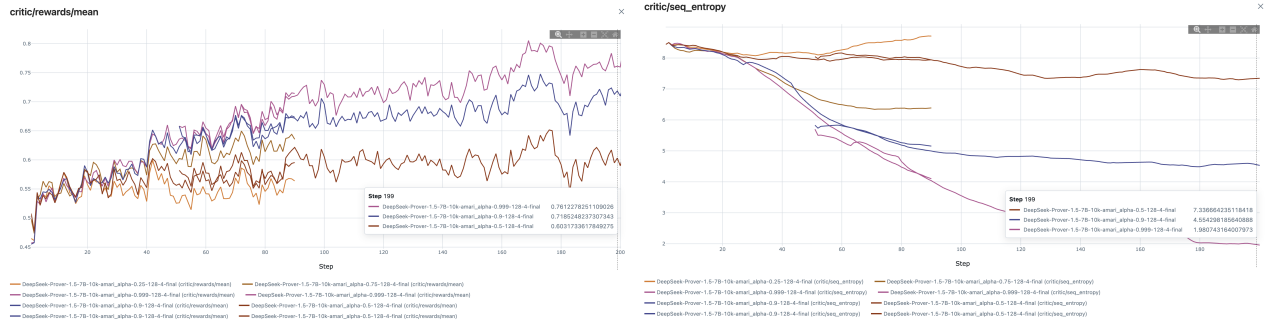
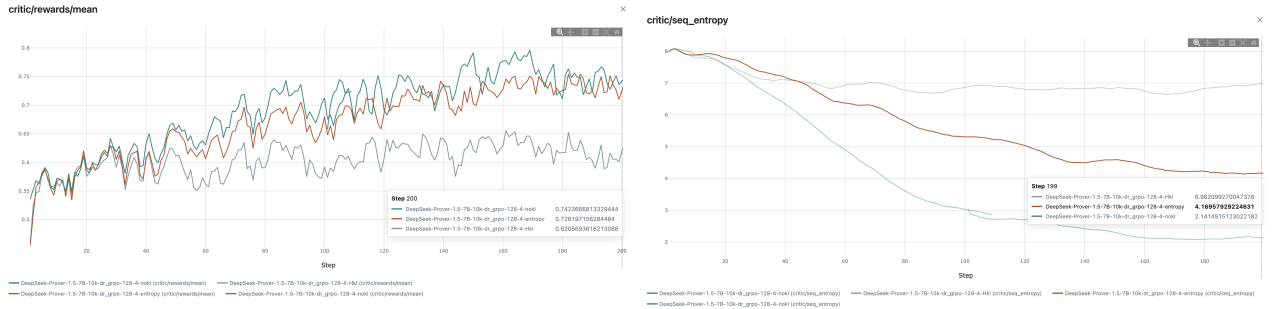

 Figure 5: Training curves of both α -DPG and dr-GRPO. Sequence entropy on the right and reward on the left

 Figure 6: Training curves of α -DPG for various alpha values. Sequence entropy on the right and reward on the left. (Truncated curves are runs that have been stopped and resumed)


Figure 7: Training curves various dr-grpo baselines (dr-grpo, +high KL, +entropy). Sequence entropy on the right and reward on the left. (Truncated curves are runs that have been stopped and resumed)

G. Additional Experimental Results

G.1. MATH + Minerva

To further validate the effectiveness of our approach on a different task and models, we repeat our experiments on an informal mathematical reasoning task. For this, we use the MATH dataset [21], selecting only the hardest problems in the set, which were annotated with the label “Level 5”. After this, we are left with 2304 problems.

We train all methods for 1000 iterations, with a maximum sequence length of 1024 tokens. For α -DPG, we compute the partition function online using each batch of samples so that there is no overhead in pre-computing it. Following Yue et al. [68], we evaluate the models on the Minerva dataset [32], inducing some domain shift.

We observe on pass@256 a diversity reduction effect but smaller for GRPO, which now has about the same coverage as the base model. Notably, Rewarding the Unlikely also shows a small effect, being on par with RLOO. $\alpha = 0.999$, consistently with the previous experiment, performs similarly to these models. On the other hand, $\alpha = 0.9$ achieves the highest pass@256, outperforming all other models, including the Pass@k-Training baseline. Lower values of alpha do not perform better in pass@256 and we conjecture that this comes from the fact that we did not pre-compute the partition function but rather computed it online on the basis of just 4 samples. Note that noise on the partition function should not affect higher values of alpha (see the analysis in App. H).

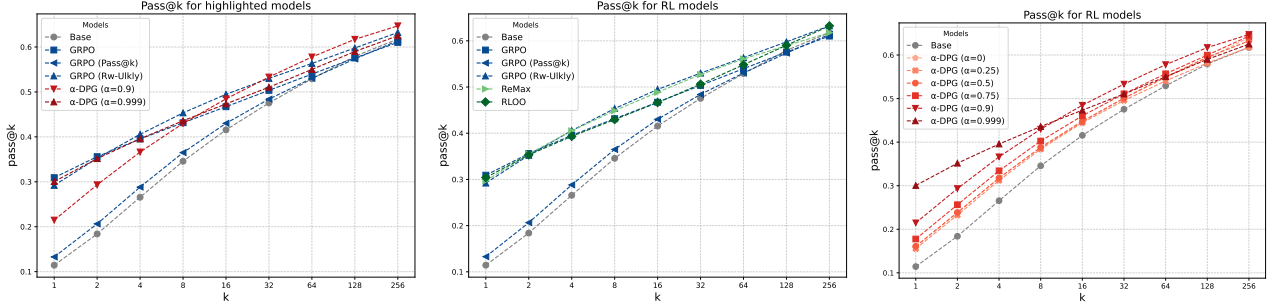


Figure 8: Pass@k curves on the Minerva dataset set for the Qwen-2.5-Math-1.5B model tuned with different methods.

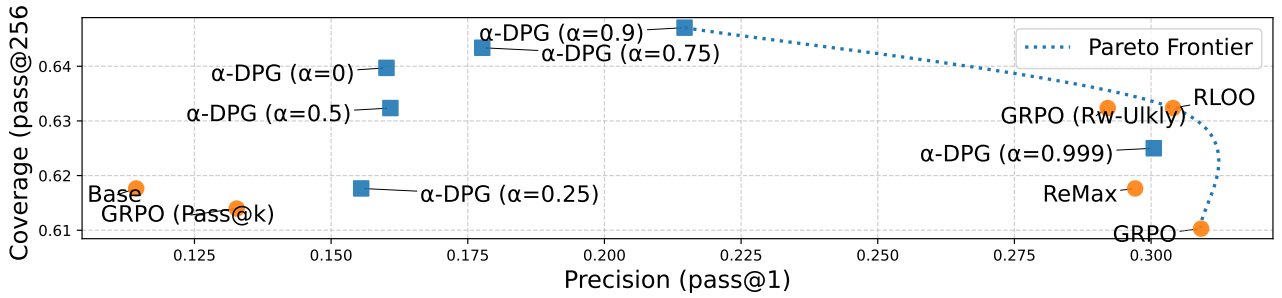


Figure 9: Pass@256 vs Pass@1 for Qwen-2.5-MATH-1.5B trained with different techniques and evaluated on the Minerva dataset

	Base	GRPO	GRPO (Pass@k)	GRPO (Rw-Ulkly)	ReMax	RLOO	α -DPG ($\alpha=0$)	α -DPG ($\alpha=0.25$)	α -DPG ($\alpha=0.5$)	α -DPG ($\alpha=0.75$)	α -DPG ($\alpha=0.9$)	α -DPG ($\alpha=0.999$)
Base	11.4/61.8	+0.7	+0.3	-1.5	+0.0	-1.5	-2.2	+0.0	-1.5	-2.6	-3.0	-0.8
GRPO	+19.5	30.9/61.0	-0.4	-2.2	-0.7	-2.2	-2.9	-0.7	-2.2	-3.3	-3.7	-1.5
GRPO (Pass@k)	+1.8	-17.6	13.3/61.4	-1.8	-0.3	-1.8	-2.6	-0.3	-1.8	-2.9	-3.3	-1.1
GRPO (Rw-Ulkly)	+17.8	-1.7	+15.9	29.2/63.3	+1.5	-0.0	-0.7	+1.5	+0.0	-1.1	-1.5	+0.7
ReMax	+18.3	-1.2	+16.4	+0.5	29.7/61.8	-1.5	-2.2	-0.0	-1.5	-2.6	-3.0	-0.8
RLOO	+19.0	-0.5	+17.1	+1.2	+0.7	30.4/63.3	-0.7	+1.5	+0.0	-1.1	-1.5	+0.7
α -DPG ($\alpha=0$)	+4.6	-14.9	+2.7	-13.2	-13.7	-14.4	16.0/64.0	+2.2	+0.7	-0.4	-0.7	+1.4
α -DPG ($\alpha=0.25$)	+4.1	-15.4	+2.3	-13.7	-14.2	-14.9	-0.5	15.5/61.8	-1.5	-2.6	-3.0	-0.8
α -DPG ($\alpha=0.5$)	+4.7	-14.8	+2.8	-13.1	-13.6	-14.3	+0.1	+0.5	16.1/63.2	-1.1	-1.5	+0.7
α -DPG ($\alpha=0.75$)	+6.3	-13.2	+4.5	-11.5	-11.9	-12.6	+1.8	+2.2	+1.7	17.8/64.4	-0.4	+1.8
α -DPG ($\alpha=0.9$)	+10.0	-9.5	+8.2	-7.7	-8.2	-8.9	+5.5	+5.9	+5.4	+3.7	21.5/64.7	+2.2
α -DPG ($\alpha=0.999$)	+18.6	-0.9	+16.8	+0.8	+0.3	-0.3	+14.0	+14.5	+14.0	+12.3	+8.6	30.1/62.5

Table 3: Pairwise performance comparison on Minerva dataset for models trained on MATH. **Diagonal**: Absolute Pass@1 / Pass@256 scores (%). **Upper Triangle**: Pass@256 differences (Row – Column). **Lower Triangle**: Pass@1 differences (Row – Column). **Bold** indicates statistical significance ($p < 0.05$) via paired bootstrap.

G.2. Kimina-Prover

We next train the Kimina-Prover-Distill-1.7B model using the Kimina-Prover-Promptset dataset [59], from which we reserve 200 problem prompts for testing. We train the models until convergence (see Figure 11) and allow sequence lengths of 8192 tokens maximum. Large sequence lengths destabilized training using bfloat16, and we resorted to float16 training instead [46]. Because of the computational load of these experiments, we have only preliminary results for GRPO and α -DPG ($\alpha=0.5$). The results are in line with earlier observations: The pass@k performance of GRPO flattens down for $k > 8$ and the base model quickly approaches the same levels of performance for higher values of k . α -DPG, on the other hand, outperforms GRPO for $k > 16$ and clearly dominates the base model for all considered values of k . We will complete these results in the final version to include all other models and baselines.

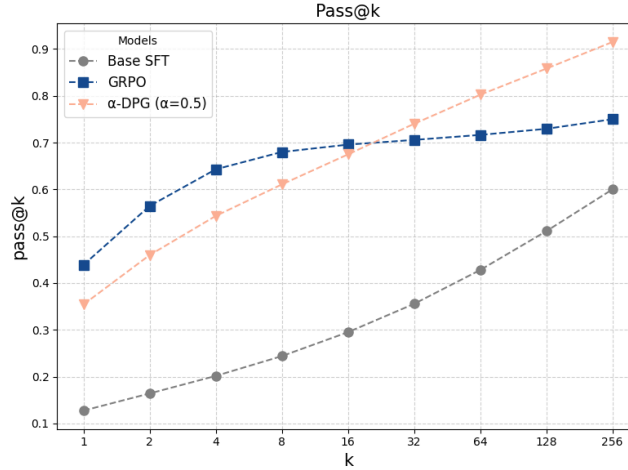


Figure 10: Pass@k curves for Kimina-Prover-Distill-1.7B trained on the Kimina-Prover-Promptset and evaluated on 200 held-out problems.

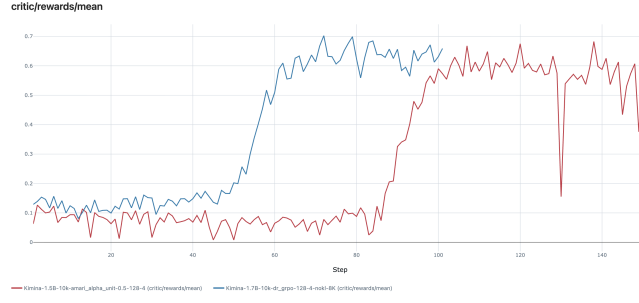


Figure 11: Training curves for Kimina-Prover-Distill-1.7B

G.3. Lean

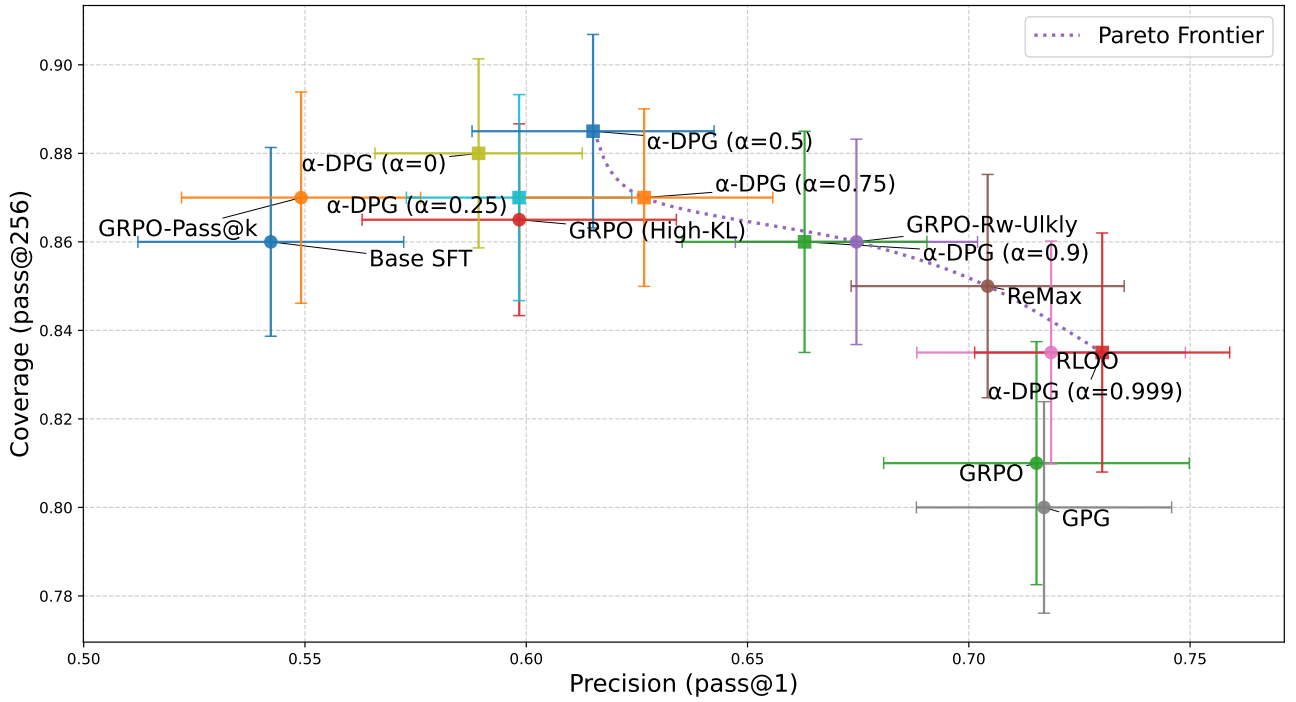


Figure 12: Estimates of models precision (pass@1) and coverage (pass@256) with bootstrap variance estimates ($n = 1000$). α -DPG models sit along a Pareto frontier.

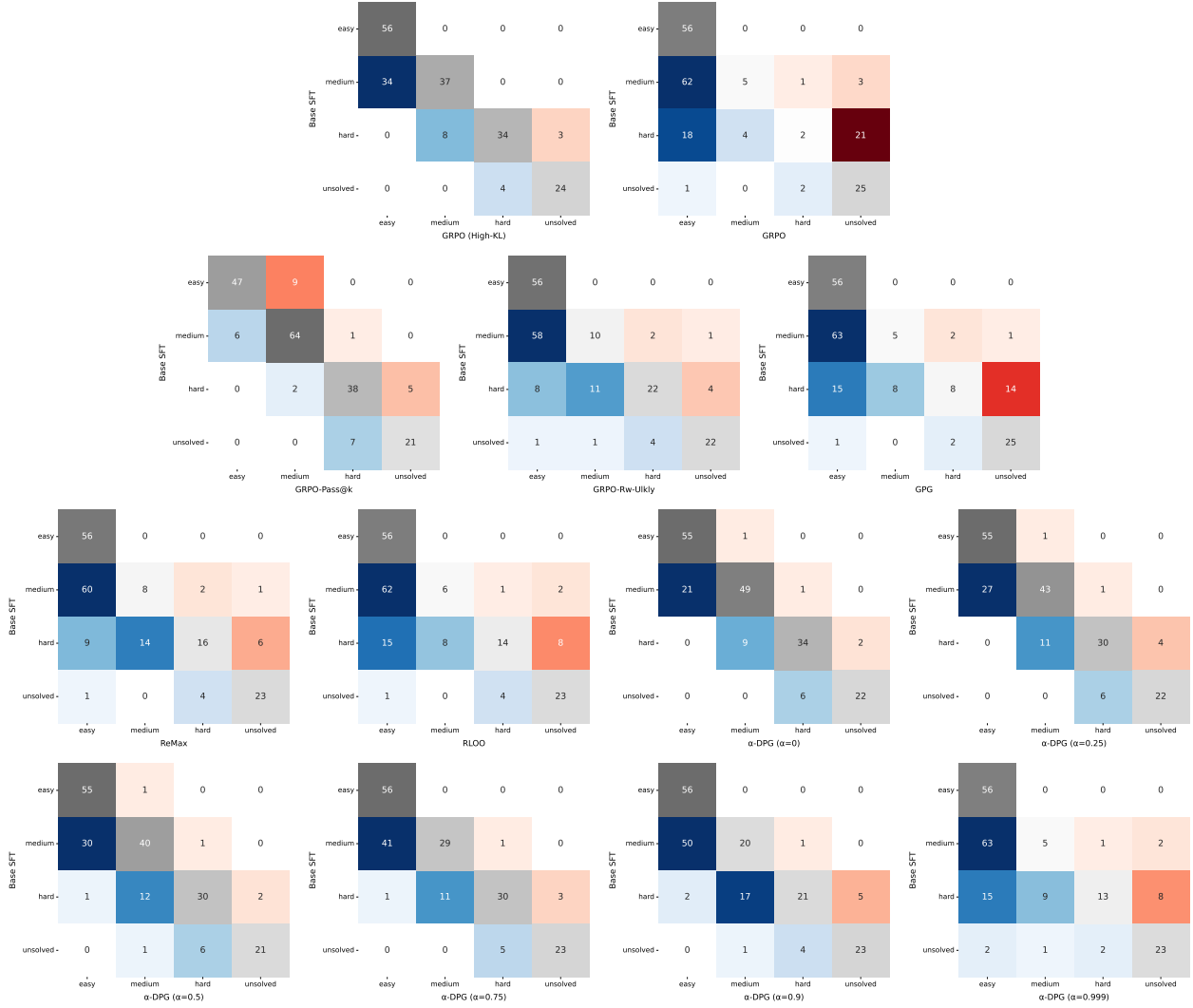


Figure 13: Problem Difficulty Transition Matrix from the Base-SFT to GRPO. The matrix shows the number of problems that transition from an initial difficulty classification under the base model (Base-SFT) (y-axis) to a final classification after post-training (x-axis).

	Base SFT	GRPO-Pass@k	GRPO	GRPO (High-KL)	GRPO-Rw-Uikly	ReMax	RLOO	GPG	α -DPG ($\alpha=0$)	α -DPG ($\alpha=0.25$)	α -DPG ($\alpha=0.5$)	α -DPG ($\alpha=0.75$)	α -DPG ($\alpha=0.9$)	α -DPG ($\alpha=0.999$)
Base SFT	54.2/86.0	-1.0	+5.0	-0.5	-0.0	+1.0	+2.5	+6.0	-2.0	-1.0	-2.5	-1.0	-0.0	+2.5
GRPO-Pass@k	+0.7	54.8/87.0	+6.0	+0.5	+1.0	+2.0	+3.5	+7.0	-1.0	-0.0	-1.5	-0.0	+1.0	+3.5
GRPO	+17.3	+16.6	71.5/81.0	-5.5	-5.0	-4.0	-2.5	+1.0	-7.0	-6.0	-7.5	-6.0	-5.0	-2.5
GRPO (High-KL)	+5.6	+4.9	-11.7	59.8/86.5	+0.5	+1.5	+3.0	+6.5	-1.5	-0.5	-2.0	-0.5	+0.5	+3.0
GRPO-Rw-Uikly	+13.3	+12.6	-4.1	+7.6	67.4/86.0	+1.0	+2.5	+6.0	-2.0	-1.0	-2.5	-1.0	-0.0	+2.5
ReMax	+16.2	+15.5	-1.1	+10.6	+3.0	70.4/85.0	+1.5	+5.0	-3.0	-2.0	-3.5	-2.0	-1.0	+1.5
RLOO	+17.7	+17.0	+0.3	+12.0	+4.4	+1.4	71.8/83.5	+3.5	-4.5	-3.5	-5.0	-3.5	-2.5	-0.0
GPG	+17.5	+16.8	+0.2	+11.9	+4.2	+1.3	-0.2	71.6/80.0	-8.0	-7.0	-8.5	-7.0	-6.0	-3.5
α -DPG ($\alpha=0$)	+4.7	+4.0	-12.6	-0.9	-8.6	-11.5	-13.0	-12.8	58.9/88.0	+1.0	-0.5	+1.0	+2.0	+4.5
α -DPG ($\alpha=0.25$)	+5.6	+4.9	-11.7	-0.0	-7.6	-10.6	-12.0	-11.9	+0.9	59.8/87.0	-1.5	+0.0	+1.0	+3.5
α -DPG ($\alpha=0.5$)	+7.3	+6.6	-10.0	+1.7	-6.0	-8.9	-10.4	-10.2	+2.6	+1.7	61.4/88.5	+1.5	+2.5	+5.0
α -DPG ($\alpha=0.75$)	+8.4	+7.8	-8.9	+2.8	-4.8	-7.8	-9.2	-9.0	+3.7	+2.8	+1.2	62.6/87.0	+1.0	+3.5
α -DPG ($\alpha=0.9$)	+12.1	+11.4	-5.3	+6.5	-1.2	-4.2	-5.6	-5.4	+7.4	+6.5	+4.8	+3.6	66.2/86.0	+2.5
α -DPG ($\alpha=0.999$)	+18.8	+18.1	+1.5	+13.2	+5.6	+2.6	+1.2	+1.3	+14.1	+13.2	+11.5	+10.4	+6.8	73.0/83.5

Table 4: Pairwise performance comparison. **Diagonal**: Absolute Pass@1 / Pass@256 scores (%). **Upper Triangle**: Pass@256 differences (Row – Column). **Lower Triangle**: Pass@1 differences (Row – Column). **Bold** indicates statistical significance ($p < 0.05$) via paired bootstrap.

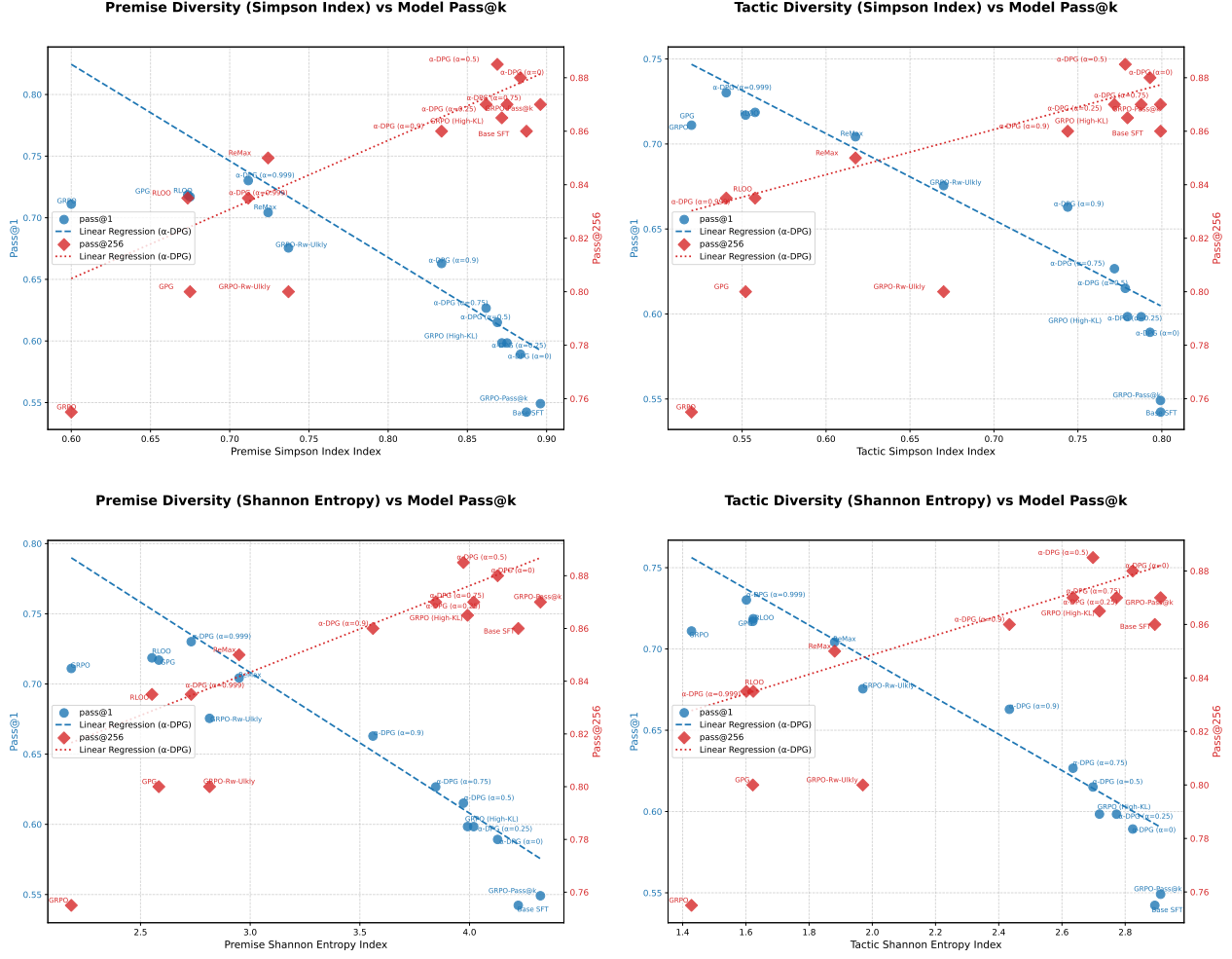


Figure 14: Diversity index vs Pass@k

Model	Premises SI	Tactics SI	Tactics Entropy	Premises Entropy	Pass@256	Pass@1
GRPO	0.6000	0.5198	1.4280	2.1844	0.7550	0.7111
α -DPG ($\alpha = 0.999$)	0.6248	0.5289	1.4605	2.2799	0.7800	0.7212
GRPO-Rw-Ulkly	0.7371	0.6700	1.9700	2.8147	0.8000	0.6755
α -DPG ($\alpha = 0.9$)	0.7653	0.7003	2.0941	2.9818	0.8100	0.6842
GRPO (High-KL)	0.8066	0.7356	2.3106	3.2714	0.8100	0.6361
α -DPG ($\alpha = 0.75$)	0.8093	0.7358	2.2919	3.3189	0.8350	0.6588
α -DPG ($\alpha = 0.5$)	0.8267	0.7420	2.3616	3.4591	0.8400	0.6404
α -DPG ($\alpha = 0.25$)	0.8368	0.7559	2.4511	3.5615	0.8600	0.6266
Base SFT	0.8442	0.7523	2.4216	3.5502	0.8150	0.5908
GRPO-Pass@k	0.8447	0.7545	2.4594	3.6700	0.8500	0.6003
α -DPG ($\alpha = 0.0$)	0.8851	0.7889	2.7843	4.1168	0.865	0.5772

Table 5: Diversity and performance metrics across models. For each problem statement, we evaluate 256 generated proof sequences. At every proof state, we compute the Simpson Index (SI) and Shannon entropy over both tactic choices and premise selections. The metrics are then aggregated across all problems to capture the overall diversity of candidate sequences. Higher diversity in tactics and premises generally correlates with improved pass@256 performance.

H. Formal complements on α -divergence

H.1. Smoothness of α -divergence, behaviour for $\alpha \rightarrow 0$ and $\alpha \rightarrow 1$

The fact that $D_{f_\alpha}(\pi, p)$ is a continuous function of α and that it converges to the forward $KL(p||\pi)$ for $\alpha \rightarrow 0$ and to the reverse $KL(\pi||p)$ for $\alpha \rightarrow 1$ — including cases where these KL-divergences may be infinite — is well-known in the literature, e.g. Cichocki & Amari [11].

In our specific situation, while typically the autoregressive policy π is full-support ($\pi(y) > 0, \forall y \in \mathcal{Y}$), the support $A := \{y : p(y) > 0\}$ of p is a proper subset of $\mathcal{Y}$. In such cases, while $KL(p||\pi)$ is (typically) finite,³ $KL(\pi||p)$ is infinite.

H.2. Comparing different policies, illustration

In order to provide a “non gradient” interpretation of what happens at the edges, it is instructive to compare $D_{f_\alpha}(\pi, p)$ with $D_{f_\alpha}(\pi', p)$ for different candidate policies π and π' .

Illustration We provide an illustration in Fig. 15, on a toy example with a small finite sample space $\mathcal{Y}$, with $A \subseteq \mathcal{Y}$, and with three policies. The policy π_1 is more “covering” than π_2 and π_3 , with a lower forward $KL(p||\pi)$, but it is less “focussed” on the valid region A than either π_2 or π_3 , which are both concentrated on A (with $\pi_2(A) = \pi_3(A) = 0.9$), but with different peaks. Despite the fact that the reverse $KL(\pi||p)$ is infinite for the three policies, the divergences, for α close to 1 — while large (and tending to infinity) — still show a clear order, with $D_{f_\alpha}(\pi_1, p)$ much higher than $D_{f_\alpha}(\pi_3, p)$, which in turn is slightly higher than $D_{f_\alpha}(\pi_2, p)$.

More in detail, we consider the discrete space $\mathcal{Y} = \{y_1, y_2, y_3\}$, over which the base model π_{base} has an (almost) uniform distribution $\pi_{base} = (0.33, 0.34, 0.33)$, and where the binary verifier $v(y)$ takes the values $(1, 0, 1)$, i.e. accepts y_1 and y_3 and rejects y_2 . This results in the target distribution:

$$p = (0.5, 0, 0.5).$$

We study three alternative policies, with distributions:

$$\pi_1 = \pi_{base} = (0.33, 0.34, 0.33), \quad \pi_2 = (0.8, 0.1, 0.1), \quad \pi_3 = (\epsilon, 0.1, 1 - \epsilon),$$

taking $\epsilon = 0.01$. π_1 is the more diverse/covering of the three, but has some significant mass on the invalid point y_2 , while π_2 and π_3 waste less mass on the invalid point, and are peaky on the first point and the third point respectively, even more so for π_3 .

The endpoints recover the forward and reverse KL divergences:

$$\lim_{\alpha \rightarrow 0} D_{f_\alpha}(\pi||p) = KL(p||\pi), \quad \lim_{\alpha \rightarrow 1} D_{f_\alpha}(\pi||p) = KL(\pi||p),$$

where we adopt the usual conventions that $KL(p||\pi) = +\infty$ whenever p charges a point on which π vanishes, and similarly for $KL(\pi||p)$.

In our setting, $p(y_2) = 0$ while each π_i assigns positive mass to y_2 , so $KL(\pi_i||p) = +\infty$ for $i = 1, 2, 3$.

Numerical values. Table 6 reports $D_{f_\alpha}(\pi_i||p)$ for a range of representative α values, including the limiting cases $\alpha = 0$ and $\alpha = 1$.

Curves as a function of α . Figure 15 shows the curves $\alpha \mapsto D_{f_\alpha}(\pi_i||p)$ for $i = 1, 2, 3$ on $\alpha \in (0, 1)$.

On the left of the plot, with α at 0, we see that π_1 , which is the more diverse/covering of the three relative to p , has the lowest value of forward KL, $KL(p, \pi)$, while π_3 , which is even more “peaky” than π_2 has a larger forward KL. On the right of the plot, with α tending to 1, the divergences all tend to infinity, but their order stabilizes, with the divergence of π_1 getting and staying much larger than both those of π_2 and π_3 . As we will see in Theorem 5 and equation 32 below, the relative behaviour of the policies for α close to 1 is determined predominantly by their relative values of $\pi(A)$, and here, with $\pi_2(A) = \pi_3(A) = 0.9 > \pi_1(A) = 0.66$, π_1 “loses” relative to π_2 and π_3 . For two policies, such as π_2, π_3 , which have exactly the same “valid” mass, their order is determined by a secondary term, the one appearing to the right of ‘+’ in equation 32.

³In some “pathological” situations, and for an infinite sample space $\mathcal{Y}$, $KL(p||q)$ can be infinite even when the supports coincide.

Distribution	$\alpha \rightarrow 0$	$\alpha = 0.1$	$\alpha = 0.5$	$\alpha = 0.7$	$\alpha = 0.9$	$\alpha = 0.99$	$\alpha \rightarrow 1$
$\pi_1 = (0.33, 0.34, 0.33)$	0.4155	0.4522	0.7504	1.2018	3.4666	34.0658	∞
$\pi_2 = (0.8, 0.1, 0.1)$	0.5697	0.5585	0.5758	0.6816	1.3252	10.3160	∞
$\pi_3 = (0.01, 0.1, 0.89)$	1.6677	1.4693	1.0488	1.1827	1.7766	10.5828	∞

Table 6: Values of $D_{f_\alpha}(\pi_i||p)$ for $p = (0.5, 0, 0.5)$ and π_1, π_2, π_3 at selected α values. For $\alpha \rightarrow 0$ and $\alpha \rightarrow 1$ we recover the forward and reverse KL divergences, respectively.

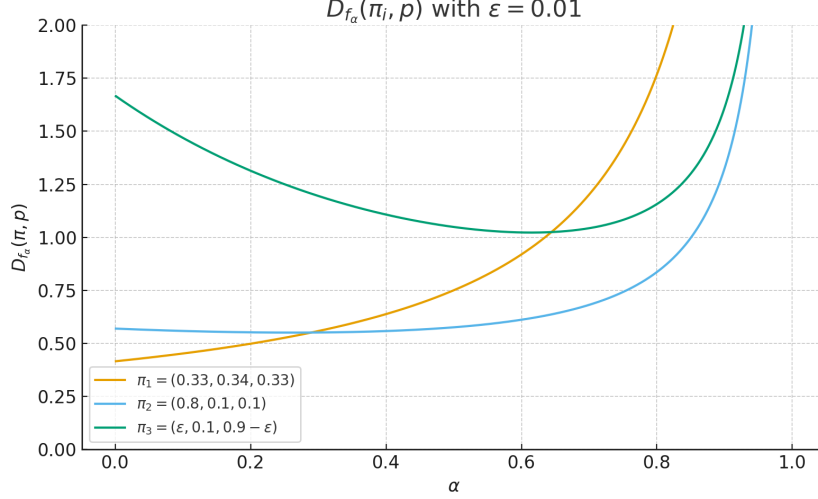


Figure 15: The divergence $D_{f_\alpha}(\pi, p)$ for $\pi_1 = (0.33, 0.34, 0.33)$, $\pi_2 = (0.8, 0.1, 0.1)$, and $\pi_3 = (\epsilon, 0.1, 0.9 - \epsilon)$ with $\epsilon = 0.01$.

H.3. A decomposition theorem for the α -divergence for targets with partial support

We now state a useful (and apparently novel) result, which permits a better understanding of what happens in the situation where the support A of the target p is strictly contained in the support of the model π . This result says that, with $\pi(A)$ the π mass of A , and with π_A is the “renormalization” of π to A , that is, $\pi_A(y) = \pi(y)/\pi(A)$ for $y \in A$, and for $\alpha \in (0, 1)$, we have the identity $D_{f_\alpha}(\pi, p) = \frac{1-\pi(A)^\alpha}{\alpha(1-\alpha)} + \pi(A)^\alpha D_{f_\alpha}(\pi_A, p)$.

This identity is especially interesting for the case of α tending to 1. In that case, with π full support and $\pi(A) < 1$, the support of π_A is equal to the support of p , and therefore $D_{f_\alpha}(\pi_A, p)$ tends to a finite value $KL(\pi_A, p)$, and the second term of the identity tends towards a finite value. On the other hand, the first term tends to infinity at a rate closer and closer to $\frac{1-\pi(A)}{1-\alpha}$, meaning that $\pi(A) > \pi'(A)$ implies that the divergence of π becomes and stays lower than the divergence of π' after a certain point α_0 .

In other words, when A is the subset of $\mathcal{Y}$ for which the binary reward $v(y)$ is equal to 1, and for α sufficiently close to 1, minimizing $D_{f_\alpha}(\pi_\theta, p)$ is essentially equivalent to maximizing $\mathbb{E}_{\pi_\theta} v(y)$, the same objective as pure REINFORCE.

Formal Result: the Support Decomposition of α -Divergence

Let π and p be probability distributions on a countable sample space Y . We consider the α -divergence defined by the generator function $f_\alpha(t) = \frac{t^\alpha - \alpha t - (1-\alpha)}{\alpha(\alpha-1)}$ for $\alpha \in (0, 1)$.

First, we establish the algebraic relationship between the divergence and the “Hellinger sum” (the discrete counterpart to the Hellinger integral [37], see also Appendix B of [35]).

Lemma 4 (Connection to Hellinger sum). *Let $\alpha \in (0, 1)$. Let $A = \text{supp}(p) \subseteq Y$. The α -divergence satisfies:*

$$D_{f_\alpha}(\pi, p) = \frac{1 - H_\alpha(\pi, p)}{\alpha(1 - \alpha)}, \quad (31)$$

where $H_\alpha(\pi, p) = \sum_{y \in Y} \pi(y)^\alpha p(y)^{1-\alpha}$ is the Hellinger sum. This holds even if $\text{supp}(\pi) \not\subseteq A$.

Proof. We use the extended definition of f -divergence [45] which includes the boundary term for the set where $p(y) = 0$ but $\pi(y) > 0$. Let $A^c = Y \setminus A$ and let $\epsilon = \pi(A^c)$ be the “leakage” mass.

$$D_{f_\alpha}(\pi, p) = \sum_{y \in A} p(y) f_\alpha \left(\frac{\pi(y)}{p(y)} \right) + f'_\alpha(\infty) \cdot \pi(A^c).$$

1. **The Boundary Term:** The term to the right uses the constant $f'_\alpha(\infty) = \lim_{t \rightarrow \infty} f_\alpha(t)/t$. For $\alpha < 1$, $t^{\alpha-1} \rightarrow 0$, so: $f'_\alpha(\infty) = \frac{-\alpha}{\alpha(\alpha-1)} = \frac{1}{1-\alpha}$. See also Table 1.

2. **The Sum on Support A:** Note that $\sum_{y \in A} \pi(y) = 1 - \epsilon$ and $\sum_{y \in A} p(y) = 1$.

$$\begin{aligned} \sum_{y \in A} p(y) f_\alpha \left(\frac{\pi(y)}{p(y)} \right) &= \frac{1}{\alpha(\alpha-1)} \left[\sum_{y \in A} \pi(y)^\alpha p(y)^{1-\alpha} - \alpha(1-\epsilon) - (1-\alpha)(1) \right] \\ &= \frac{1}{\alpha(\alpha-1)} [H_\alpha(\pi, p) - 1 + \alpha\epsilon]. \end{aligned}$$

3. **Combination:** Adding the boundary term: Boundary = $\frac{\epsilon}{1-\alpha} = \frac{-\alpha\epsilon}{\alpha(\alpha-1)}$. The terms involving ϵ cancel perfectly:

$$D_{f_\alpha}(\pi, p) = \frac{H_\alpha(\pi, p) - 1 + \alpha\epsilon - \alpha\epsilon}{\alpha(\alpha-1)} = \frac{1 - H_\alpha(\pi, p)}{\alpha(1-\alpha)}.$$

□

Using Lemma 4, we now derive the main decomposition theorem.

Theorem 5 (Support Decomposition). *Assume that $A = \text{supp}(p)$ is strictly included in $\text{supp}(\pi)$. Let π_A be the renormalization of π on A , i.e., $\pi_A(y) = \pi(y)/\pi(A)$ for $y \in A$, or, equivalently $\pi_A(y) = \pi(y|A)$. The divergence decomposes as:*

$$D_{f_\alpha}(\pi, p) = \underbrace{\frac{1 - \pi(A)^\alpha}{\alpha(1-\alpha)}}_{\text{Leakage Penalty}} + \underbrace{\pi(A)^\alpha D_{f_\alpha}(\pi_A, p)}_{\text{Shape Divergence}}. \quad (32)$$

Proof. Since $p(y) = 0$ for $y \notin A$, the Hellinger sum restricts to A . Substituting $\pi(y) = \pi(A)\pi_A(y)$:

$$H_\alpha(\pi, p) = \sum_{y \in A} (\pi(A)\pi_A(y))^\alpha p(y)^{1-\alpha} = \pi(A)^\alpha H_\alpha(\pi_A, p).$$

From Lemma 4, we have $H_\alpha(\pi_A, p) = 1 - \alpha(1-\alpha)D_{f_\alpha}(\pi_A, p)$. Substituting this back into the global divergence formula:

$$\begin{aligned} D_{f_\alpha}(\pi, p) &= \frac{1 - \pi(A)^\alpha H_\alpha(\pi_A, p)}{\alpha(1-\alpha)} \\ &= \frac{1 - \pi(A)^\alpha [1 - \alpha(1-\alpha)D_{f_\alpha}(\pi_A, p)]}{\alpha(1-\alpha)} \\ &= \frac{1 - \pi(A)^\alpha}{\alpha(1-\alpha)} + \pi(A)^\alpha D_{f_\alpha}(\pi_A, p). \end{aligned}$$

□

Consequences for Fixed Target p

We analyze the case where π has full support and p has strictly partial support A .

Remark 1 (Limit $\alpha \rightarrow 1$: The Strong Constraint). As $\alpha \rightarrow 1$, D_{f_α} converges to the Reverse KL divergence $D_{KL}(\pi||p) = +\infty$. The Mass Penalty term diverges:

$$\lim_{\alpha \rightarrow 1} \frac{1 - \pi(A)^\alpha}{\alpha(1-\alpha)} = +\infty \quad (\text{if } \pi(A) < 1).$$

This acts as a strong constraint, heavily penalizing support leakage.

As for the Shape divergence term $\pi(A)^\alpha D_{f_\alpha}(\pi_A, p)$, it remains finite, and is dominated by the first term.

Remark 2 (Limit $\alpha \rightarrow 0$: The Weak Constraint). As $\alpha \rightarrow 0$, D_{f_α} converges to the Forward KL divergence $D_{KL}(p||\pi)$. The Mass Penalty term remains finite:

$$\lim_{\alpha \rightarrow 0} \frac{1 - \pi(A)^\alpha}{\alpha(1 - \alpha)} = -\ln(\pi(A)).$$

This acts as a soft constraint (“Surprise Penalty”), allowing a trade-off between coverage ($\pi(A)$) and conditional shape matching ($D_{KL}(p||\pi_A)$). The Shape divergence term $\pi(A)^\alpha D_{f_\alpha}(\pi_A, p)$ also remains finite and converges to $D_{f_\alpha}(\pi_A, p)$.