
Provably Mitigating Overoptimization in RLHF: Your SFT Loss is Implicitly an Adversarial Regularizer

Zhihan Liu^{1*} Miao Lu^{2*} Shenao Zhang¹ Boyi Liu³
Hongyi Guo¹ Yingxiang Yang³ Jose Blanchet² Zhaoran Wang¹

¹Northwestern University ²Stanford University ³ByteDance Inc.
{zhihanliu2027, shenaozhang2028, hongyiguo2025}@u.northwestern.edu
{miaolu, jose.blanchet}@stanford.edu, zhaoranwang@gmail.com
{boyi.liu01, yingxiang.yang}@bytedance.com

Abstract

Aligning generative models with human preference via RLHF typically suffers from overoptimization, where an imperfectly learned reward model can misguide the generative model to output undesired responses. We investigate this problem in a principled manner by identifying the source of the misalignment as a form of distributional shift and uncertainty in learning human preferences. To mitigate overoptimization, we first propose a theoretical algorithm that chooses the best policy for an adversarially chosen reward model; one that simultaneously minimizes the maximum likelihood estimation of the loss and a reward penalty term. The penalty term is introduced to prevent the policy from choosing actions with spurious high proxy rewards, resulting in provable sample efficiency of the algorithm under a *partial coverage* style condition. Moving from theory to practice, the proposed algorithm further enjoys an equivalent but surprisingly easy-to-implement reformulation. Using the equivalence between reward models and the corresponding optimal policy, the algorithm features a simple objective that combines: (i) a preference optimization loss that directly aligns the policy with human preference, and (ii) a supervised learning loss that explicitly imitates the policy with a (suitable) baseline distribution. In the context of aligning large language models (LLM), this objective fuses the direct preference optimization (DPO) loss with the supervised fine-tuning (SFT) loss to help mitigate the overoptimization towards undesired responses, for which we name the algorithm Regularized Preference Optimization (RPO). Experiments of aligning LLMs demonstrate the improved performance of RPO compared with DPO baselines. Our work sheds light on the interplay between preference optimization and SFT in tuning LLMs with both theoretical guarantees and empirical evidence.

1 Introduction

A key step in building state-of-the-art LLMs is Reinforcement Learning from Human Feedback (RLHF) [12, 87], which aligns pretrained LLMs with human preferences using human assessment data, making the model more helpful, truthful, and harmless [42, 10]. Typically, RLHF first learns a reward model from data (pair-wise comparisons of responses) to quantify the human preferences of LLM outputs. Then it fine-tunes the LLM to maximize the learned reward using RL techniques.

In this pipeline, a crucial challenge is *reward overoptimization* or *reward hacking* [40, 62, 25]. Since the reward model is learned from finite data, it might not be perfectly aligned with the underlying human preference. Optimizing the LLM towards such an imperfectly learned and potentially overfitted

*Equal contribution.

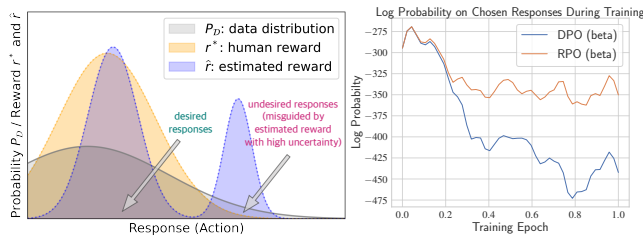


Figure 1: *Left*: Reward overoptimization due to the distributional shift and uncertainty in reward. *Right*: Overoptimization causes the probability of outputting preferred responses in the preference data to decrease substantially using original DPO proposed by [46]. Our algorithm (RPO) significantly alleviates this decrease. See more discussions in Section 6.

reward model leads to performance degeneration and a substantial decrease in the probability of choosing the preferred responses in the data [26, 45]. Given the importance of RLHF and the outlined challenge, a crucial research question is: *How to mitigate reward overoptimization in RLHF in a principled and efficient manner for better alignment?*

To answer the question, we model RLHF as an offline contextual bandit [42] and ascribe overoptimization to distributional shifts and reward uncertainty. Intuitively, when fine-tuning an LLM, the response (action) distribution of the tuned LLM could deviate from that of the training data. For the out-of-distribution responses, which are dissimilar with (or not well covered by) the responses in the data, the high inherent uncertainty of underlying human preferences could make the learned reward model misleading for out-of-distribution responses. In this situation, reward overoptimization can occur because the LLM is fine-tuned towards maximizing a reward model with defective out-of-distribution prediction, giving a potential consequence that the LLM responses are favored by the learned reward but less preferred by a human [86]. We illustrate these types of distributional shift and reward uncertainty issues inherent to overoptimization in Figure 1.

In this paper, we propose a new RLHF algorithm to mitigate reward overoptimization. From a high level, our theoretical algorithm seeks the best LLM for an *adversarially* chosen reward model that minimizes the sum of its maximum likelihood estimation loss and its own expected reward value. Intuitively, since the reward value is also minimized when minimizing the sum, it can automatically prevent the misleadingly high reward caused by the uncertainty inherent in having access to finite data. Furthermore, we show that the theoretical algorithm enjoys an easy implementation: it simply adopts a supervised fine-tuning (SFT) loss as a regularizer during training. By explicitly regularizing the LLM to imitate high-quality responses (e.g., preferred responses in dataset), the algorithm can effectively mitigate the issue of overoptimization. We establish theoretical guarantees and conduct experiments to demonstrate our findings, which we summarize next.

1.1 Our Contributions and Related Works

We summarize our contributions in three areas as follows.

A theoretical algorithm under general function approximation. Our first contribution is a new theoretical algorithm (Algorithm 1). It features an *unconstrained maximin* problem, outputting the optimal policy (LLM) against an adversarially chosen reward model that minimizes the summation of: (a) the MLE loss for estimating the underlying reward; and (b) a reward expected value term as a penalty that aims to prevent spuriously high reward estimation caused by data uncertainty and insufficient coverage. Algorithm 1 is compatible with general function approximations of the reward model, meaning that we do not impose any specific structural form to the hypothesis class of reward, demonstrating its generality, especially in language modeling.

In this regime of reward class, we establish the finite-sample suboptimality gap of Algorithm 1 as $\tilde{O}(C_{\text{coverage}}^2 \sqrt{\mathcal{N}_{\mathcal{R}}/N})$ when competing with any LLM in terms of the underlying true human reward (Theorem 5.3). Here N is the number of human comparison data, $\mathcal{N}_{\mathcal{R}}$ is the complexity of the reward model class $\mathcal{R}$, and C_{coverage} characterizes the coverage of the preference dataset with respect to the response distribution of the LLM to compete (please see Assumption 5.2 for details). This indicates that, as long as the training data well cover the LLM π to compete, the algorithm is guaranteed to align an LLM to output responses as good as π in terms of human reward, without suffering from overoptimization caused by distributional shifts and inherent uncertainty in human preference.

An easy-to-implement practical objective. Moving towards practice, we show that the objective of Algorithm 1 adopts a surprisingly simple and equivalent form for its use in practice. Specifically, with mild regularity conditions, we prove that the *maximin* objective (Algorithm 1) is equivalent to the

corresponding *minimax* objective, which is further reduced to a single minimization problem for the reward model since its inner problem adopts a closed form solution. Inspired from recent progress in RLHF that explores reward-model-free methods to align LLMs [46], we further re-parameterize the reward model via its corresponding KL-regularized optimal policy. Then the minimization objective of the reward modeling naturally translates to a target for directly aligning the LLM, which we call Regularized Preference Optimization (RPO; Algorithm 2). The objective of RPO features a simple weighted combination of two losses:

$$\text{RPO objective} = \text{Preference optimization loss} + \text{Imitation (SFT) loss}.$$

Here the *Preference optimization loss* coincides with the DPO [46] objective, tending to optimize the LLM towards maximizing the underlying true reward. The *Imitation (SFT) loss* explicitly supervises the LLM to mimic the responses from a proper distribution well covered by the dataset. The choice of the distribution is guided and justified by our theory of Algorithm 1, but can also be flexibly adapted in practice, e.g., the preferred response in the dataset, or the responses of the initial model.

We highlight that the *Imitation (SFT) loss* serves as an important term to mitigate overoptimization. Even though the original DPO objective has already involved a KL regularization between the tuned LLM and the initial LLM, is not enough to prevent overoptimization. As we elaborate in Section 4, the KL-regularization weight of the DPO objective could only control the scale of the gradient per training example, while the RPO objective can further modify the gradient direction. Calling back to the theoretical Algorithm 1, such a modification of gradient direction originates from the reward penalty in the adversarial objective for the reward model. This modification, as we expose in our theoretical analysis, helps to mitigate overoptimization. *Thus, incorporating SFT loss in RLHF gives you a regularizer that provably mitigates overoptimization.*

Empirical evaluations. Following the training setup of two series of released chat models Zephyr-7b-beta (trained on the Ultrafeedback dataset [16] by DPO) and Zephyr-7b-gemma (trained on the Argilla-DPO-Mix-7K dataset [3] by DPO) [64], we implement RPO for the beta series and gemma series respectively to show that: (i) RPO is a flexible plug-in module and can be applied to different reference models. (ii) RPO can alleviate the overoptimization issue. (iii) RPO consistently achieves better alignment performance than DPO in in-data distribution. (iv) RPO can also achieve consistently better performance in standard LLM benchmarks like MT-bench and AlpacaEval 2.0, which shows its potential of mitigating overoptimization for better alignment performance, justifying our theory.

Related works. Due to space limitation, we refer the readers to Appendix A for a detailed discussion.

2 Preliminaries of RLHF

In this section, we introduce the mathematical framework of studying RLHF for aligning LLMs. We adopt the framework of offline contextual bandits [42], where we identify the context space $\mathcal{X}$ as the space of prompts and the action space $\mathcal{A}$ as the space of responses. An LLM, defined as a policy $\pi(\cdot|\cdot) : \mathcal{X} \mapsto \Delta(\mathcal{A})$, takes a prompt $x \in \mathcal{X}$ as input and output a response $a \in \mathcal{A}$ from $a \sim \pi(\cdot|x)$.

Preference model. Given any reward function $r : \mathcal{X} \times \mathcal{A} \mapsto \mathbb{R}$ belonging to certain reward class $\mathcal{R}$ that represents the “human’s rating” of LLM responses given prompts, we consider the Bradley-Terry model [9] of human preference. That is, given a prompt $x \in \mathcal{X}$ and two responses $a^1, a^0 \in \mathcal{A}$, the probability of a^1 being preferred to a^0 (denoted by $y = 1$, otherwise $y = 0$) is given by

$$\mathbb{P}_r(y = 1|x, a^1, a^0) = \frac{\exp(r(x, a^1))}{\exp(r(x, a^1)) + \exp(r(x, a^0))} = \sigma(r(x, a^1) - r(x, a^0)), \quad (2.1)$$

where $\sigma(z) = 1/(1 + \exp(-z))$ is the sigmoid function. For simplicity of future discussion, we explicitly write out the dependence of the preference probability $\mathbb{P}_r(\cdot)$ on the reward model $r \in \mathcal{R}$. In the section of theory, i.e., Section 5, we specify the assumptions on the reward model class $\mathcal{R}$.

Learning protocol. Typically, the RLHF pipeline starts from certain reference policy π^{ref} obtained from pretraining. Then RLHF aligns the LLM based on certain human preference data. In this work, we consider offline RLHF setup, where the LLM is aligned using a fixed offline preference dataset $\mathcal{D}$. It consists of N i.i.d. tuples in the form of $\mathcal{D} = \{(x_i, a_i^1, a_i^0, y_i)\}_{i=1}^N$. Here the prompt x_i and the responses a_i^1, a_i^0 are distributed according to: $(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot)$, and conditioning on (x_i, a_i^1, a_i^0) , y_i is distributed according to (2.1) for an underlying true (but unknown) reward model $r^* \in \mathcal{R}$.

Algorithm 1 Theoretical Algorithm: Maximin Objective

- 1: **Input:** Preference dataset $\mathcal{D}$, parameters $\beta, \eta > 0$, reference policy π^{ref} , baseline policy π^{base} .
 - 2: **Output:** Policy $\hat{\pi}$ given by (3.2) with the cross-entropy loss function $\mathcal{L}_{\mathcal{D}}$ defined in (3.1)..
-

Performance metric. The target of RLHF is to align an LLM, or equivalently, to learn a policy π , so as to maximize the expected true reward r^* . Thus, we define the value function of any policy π as

$$J(\pi) = \mathbb{E}_{x \sim d_0, a \sim \pi(\cdot|x)} [r^*(x, a)]. \quad (2.2)$$

Here we allow the prompt distribution $d_0(\cdot)$ to be different from that of the offline dataset distribution $\mu_{\mathcal{D}}(\cdot)$, but is assumed to be known. In the meanwhile, we consider the policies that share the same support as the reference policy π^{ref} [71], that is, we take a policy class Π as

$$\Pi = \left\{ \pi : \mathcal{X} \mapsto \Delta(\mathcal{A}) \mid \text{Supp}(\pi(\cdot|x)) \subseteq \text{Supp}(\pi^{\text{ref}}(\cdot|x)), \forall x \in \mathcal{X} \right\}. \quad (2.3)$$

The performance gap of a learned policy $\hat{\pi} \in \Pi$ w.r.t. any other policy $\pi \in \Pi$ is measured as

$$\text{Gap}^{\pi}(\hat{\pi}) = J(\pi) - J(\hat{\pi}), \text{ given policy } \pi. \quad (2.4)$$

The goal is to propose a sample-efficient and also implementation-friendly algorithm to learn a policy $\hat{\pi} \in \Pi$ able to compete with any given policy $\pi \in \Pi$ in terms of $\text{Gap}^{\pi}(\hat{\pi}) \leq \varepsilon$, with sample complexity polynomial in $1/\varepsilon$ and logarithmic in the complexity of $\mathcal{R}$.

3 A Theory-motivated Objective

Our method seeks to find the best policy $\hat{\pi}$ against an *adversarially* chosen reward model $\hat{r}_{\text{adv}}$ that minimizes a weighted sum of its expected value and the maximum likelihood estimation (MLE) loss. Intuitively, such a reward model can prevent the overoptimization issue by taking its own value into account when minimizing the MLE loss. Since the reward value is also minimized when minimizing the sum, this method prevents the misleadingly high reward caused by the uncertainty due to finite data. Formally, given two hyperparameters $\beta, \eta > 0$ and a “baseline policy” π^{base} , we define

$$T_{\beta, \eta}^{\text{adv}}(\pi) = \min_{r \in \mathcal{R}} \left\{ \eta \mathbb{E}_{\substack{x \sim d_0, a^1 \sim \pi(\cdot|x), \\ a^0 \sim \pi^{\text{base}}(\cdot|x)}} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\},$$

where the loss function $\mathcal{L}_{\mathcal{D}}(\cdot)$ is the average negative log-likelihood function of the BT model (2.1) (and here it becomes the cross-entropy loss) over the preference dataset $\mathcal{D}$, defined as

$$\mathcal{L}_{\mathcal{D}}(r) = -\widehat{\mathbb{E}}_{\mathcal{D}} \left[y_i \log(\sigma(r(x_i, a_i^1) - r(x_i, a_i^0))) + (1 - y_i) \log(\sigma(r(x_i, a_i^0) - r(x_i, a_i^1))) \right]. \quad (3.1)$$

As we can see, $T_{\beta, \eta}^{\text{adv}}(\pi)$ is the minimum value of a weighted sum of the MLE loss and the expected reward value of π , but with two important modifications that we explain in the following.

Firstly, we subtract another expected reward of certain policy π^{base} . This is because the BT model (2.1) essentially only uses the reward differences to define the preference probabilities. As a result, the data can only reveal information of the differences between the true reward r^* of different responses [78]. Accordingly, we subtract such a baseline expected reward value to match this observation. The choice of the baseline policy is discussed in the theory part (Section 5) and experiments (Section 6).

Secondly, we subtract a KL divergence between π and π^{ref} from the expected reward, weighted by the coefficient $\beta > 0$. Such a term is for practical considerations that would be explained in Sections 4 and 5.2. We note that the KL regularized reward is commonly adopted in RLHF practice to ensure the learned policy is not far away from the reference policy [42, 71].

Finally, the overall algorithm design (Algorithm 1) is to output the policy that maximizes $T_{\beta, \eta}^{\text{adv}}(\pi)$, i.e., $\hat{\pi} \in \arg\max_{\pi \in \Pi} T_{\beta, \eta}^{\text{adv}}(\pi)$, which gives the following theoretical target:

$$\hat{\pi} \in \operatorname{argmax}_{\pi \in \Pi} \min_{r \in \mathcal{R}} \left\{ \eta \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\}. \quad (3.2)$$

Given the form of (3.2), we name it the *maximin* objective in the sequel. Upon seeing (3.2), one might be arguing that such a theory-motivated objective seems hard to implement in practice. Nevertheless, in the coming Section 4, we demonstrate that the maximin objective (3.2) adopts an easy-to-implement equivalent form, allowing us to design a practical algorithm for aligning LLMs.

4 An Equivalent and Implementation-friendly Objective

In this section, we propose another *minimax*-style objective that is equivalent to the maximin objective (3.2). Based on the minimax objective, we propose a new LLM aligning algorithm called Regularized Preference Optimization (RPO). It draws inspirations from the reparametrization technique originated in Direct Preference Optimization (DPO) [46] and goes beyond to further address the overoptimization issue in offline RLHF by incorporating an SFT loss as an explicit adversarial regularizer.

An equivalent minimax objective. If the reward model class $\mathcal{R}$ satisfies certain regularity conditions, which we discuss in detail in Section 5.2, the minimax theorem holds: solving the *maximin* objective (3.2) is *equivalent* to solving a *minimax* target, given by

$$\min_{r \in \mathcal{R}} \max_{\pi \in \Pi} \left\{ \eta \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\}. \quad (4.1)$$

Such a minimax formulation (4.1) is the starting point of our practical algorithm. The magic of (4.1) is that the inner maximization problem adopts a closed form solution, which further simplifies such an objective. To see this, note that given any reward model $r \in \mathcal{R}$, the inner problem is equivalent to

$$\max_{\pi \in \Pi} \left\{ \mathbb{E}_{x \sim d_0, a \sim \pi(\cdot|x)} \left[r(x, a) - \beta \cdot \text{KL}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] \right\}. \quad (4.2)$$

It has been well established that the policy that maximizes the KL-regularized expected reward (4.2) has a closed form solution. Due to its importance, we present it as the following lemma.

Lemma 4.1 (Oracle optimal KL-regularized policy). *Given any reward model $r \in \mathcal{R}$, the optimal policy π_r to the maximization problem (4.2) is given by*

$$\pi_r(\cdot|x) = \frac{1}{Z_r(x)} \cdot \pi^{\text{ref}}(\cdot|x) \cdot \exp(\beta^{-1} r(x, \cdot)), \quad Z_r(x) = \int_{a \in \mathcal{A}} \exp(\beta^{-1} r(x, a)) d\pi^{\text{ref}}(a|x),$$

and correspondingly the optimal value of (4.2) is given by (4.2) = $\mathbb{E}_{x \sim d_0} [\beta \cdot \log(Z_r(x))]$.

Specifically, by Lemma 4.1, we can solve the inner maximization problem in (4.1) and obtain that

$$(4.1) = \min_{r \in \mathcal{R}} \left\{ \eta \mathbb{E}_{x \sim d_0, a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[-r(x, a^0) + \beta \cdot \log(Z_r(x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\}.$$

Furthermore, from Lemma 4.1, one immediately see that given any reward model $r \in \mathcal{R}$, we can reparameterize it via its corresponding optimal KL-regularized policy π_r [46], that is,

$$r(x, \cdot) = \beta \cdot \log \left(\frac{\pi_r(\cdot|x)}{\pi^{\text{ref}}(\cdot|x)} \right) + \beta \cdot \log(Z_r(x)). \quad (4.3)$$

Taking (4.3) back into (4.1), we are able to further simplify it as

$$(4.1) = \min_{r \in \mathcal{R}} \left\{ \eta \mathbb{E}_{x \sim d_0, a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[-\beta \cdot \log(\pi_r(a^0|x)) \right] + \mathcal{L}_{\mathcal{D}} \left(\beta \cdot \log \left(\frac{\pi_r(\cdot|x)}{\pi^{\text{ref}}(\cdot|x)} \right) \right) \right\}. \quad (4.4)$$

Thanks to the KL-regularization term in the original minimax objective (4.1) (or equivalently, the maximin objective (3.2)), we have the following theorem. It theoretically shows that the policy $\pi_{\hat{r}}$ associated with the reward model $\hat{r}$ solving (4.4) also solves the maximin target (3.2) of the theoretical algorithm (Algorithm 1) that enjoys finite-sample convergence guarantees. (Please see Section 5.2 for a formal statement and proof of Theorem 4.2).

Algorithm 2 Practical Algorithm: Regularized Preference Optimization (RPO)

- 1: **Input:** Preference dataset $\mathcal{D}$, parameters $\beta, \eta > 0$, reference policy π^{ref} , baseline policy π^{base} .
 - 2: **Output:** Policy $\pi_{\hat{\theta}}$ obtained by optimizing objective (4.5).
-

Theorem 4.2 (Equivalence between *maximin* and *minimax* algorithm (informal)). *Under certain regularity assumptions on $\mathcal{R}$ and given $\eta, \beta > 0$, solving the minimax objective (4.1) via (4.4), i.e.,*

$$\hat{r} = \operatorname{argmin}_{r \in \mathcal{R}} \left\{ \eta \mathbb{E}_{x \sim d_0, a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[-\beta \cdot \log(\pi_r(a^0|x)) \right] + \mathcal{L}_{\mathcal{D}} \left(\beta \cdot \log \left(\frac{\pi_r(\cdot|x)}{\pi^{\text{ref}}(\cdot|x)} \right) \right) \right\},$$

then the corresponding optimal KL-regularized policy $\pi_{\hat{r}}$ also solves the maximin objective (3.2).

Regularized Preference Optimization. Target (4.4) gives a quite simple objective to use in practice! Since (4.4) depends on $r \in \mathcal{R}$ only through its corresponding optimal policy π_r , one can formulate a minimization objective over a parameterized policy π_{θ} , i.e., the LLM to be aligned, and directly optimize the parameters $\theta \in \Theta$. More formally, the new RLHF objective becomes

$$\min_{\theta \in \Theta} \left\{ \underbrace{\mathcal{L}_{\text{RPO}}(\theta) := \eta \beta \cdot \mathbb{E}_{x \sim d_0, a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[-\log(\pi_{\theta}(a^0|x)) \right]}_{\text{Imitation (SFT) loss}} + \underbrace{\mathcal{L}_{\mathcal{D}} \left(\beta \cdot \log \left(\frac{\pi_{\theta}(\cdot|x)}{\pi^{\text{ref}}(\cdot|x)} \right) \right)}_{\text{Preference opt. loss}} \right\}. \quad (4.5)$$

In (4.5), the second term coincides with the objective of DPO algorithm [46] which optimizes the policy towards maximizing the underlying true reward, and the first term stands for a regularization term weighted by $\eta \cdot \beta$ which *explicitly* regularizes the policy to imitate the baseline policy. Therefore, we name the resulting algorithm as Regularized Preference Optimization (RPO). We summarize it abstractly in Algorithm 2. As for DPO, implementing RPO does not require to maintain a reward model r . Thus it is computationally more friendly compared to reward-based algorithms.

How does RPO improve DPO? We illustrate the effect of the imitation loss by analyzing the gradient of the RPO target $\mathcal{L}_{\text{RPO}}(\theta)$ in (4.5). Notice that by (4.5) we have

$$\nabla_{\theta} \mathcal{L}_{\text{RPO}}(\theta) = \underbrace{\eta \beta \cdot \mathbb{E}_{x \sim d_0, a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[-\nabla_{\theta} \log(\pi_{\theta}(a^0|x)) \right]}_{\text{increase the alignment with the baseline policy}} + \underbrace{\nabla_{\theta} \mathcal{L}_{\text{DPO}}(\theta)}_{\text{decrease the DPO Loss}},$$

where the derivative of the DPO loss $\nabla_{\theta} \mathcal{L}_{\text{DPO}}(\theta)$ is given by the following,

$$\nabla_{\theta} \mathcal{L}_{\text{DPO}}(\theta) = -\widehat{\mathbb{E}}_{\mathcal{D}} \left[\underbrace{\beta \cdot \sigma(\hat{r}_{\theta}(x, a_{\text{rej}}) - \hat{r}_{\theta}(x, a_{\text{cho}}))}_{\text{gradient weight}} \cdot (\nabla_{\theta} \log \pi_{\theta}(a_{\text{cho}}|x) - \nabla_{\theta} \log \pi_{\theta}(a_{\text{rej}}|x)) \right].$$

For simplicity we denote $\hat{r}_{\theta}(x, a) = \beta \cdot \log(\pi_{\theta}(x, a)) / \log(\pi^{\text{ref}}(x, a))$, a_{cho} for the chosen response and a_{rej} for the rejected response. Intuitively, RPO (4.5) modifies the **gradient direction** of DPO to ensure the alignment with the baseline policy π^{base} , and the hyper-parameter η controls the power of alignment. In comparison, the hyper-parameter β in DPO only controls the **gradient weight** when increasing the likelihood of a_{cho} and decreasing the likelihood a_{rej} . In this perspective, the hyper-parameter β only changes the scale of the gradient instead of the direction. By introducing η , we stabilize the training and reduce the side-effect of uncertain labels in data to prevent overoptimization.

5 Theoretical Analysis

In this section, we establish theoretical analysis for Algorithms 1 and 2. We take the space of prompts and responses as compact subsets $\mathcal{X} \subseteq \mathbb{R}^{d_{\mathcal{X}}}$ and $\mathcal{A} \subseteq \mathbb{R}^{d_{\mathcal{A}}}$. We take the policy class Π as (2.3).

5.1 Establishing the Sample Complexity of Maximin Objective (Algorithm 1)

Assumption 5.1 (True reward model). *We assume that the true reward model $r^* \in \mathcal{R}$, and for any $r \in \mathcal{R}$ and $(x, a) \in \mathcal{X} \times \mathcal{A}$, it holds that $r(x, a) \in [0, R]$.*

Assumption 5.2 (Partial coverage coefficient [78]). *Given a policy $\pi \in \Pi$, the coverage coefficient of the offline dataset distribution $\mu_{\mathcal{D}}$ w.r.t. reward model class $\mathcal{R}$, policy π , and the baseline policy π^{base} , denoted by $C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}})$, is defined as*

$$\max \left\{ 0, \sup_{r \in \mathcal{R}} \frac{\mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} [(r^*(x, a^1) - r^*(x, a^0)) - (r(x, a^1) - r(x, a^0))] \right\}.$$

We assume that $C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}}) < +\infty$ for the policy π to compete. We remark that the quantity $C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}})$ is upper bounded by the density ratio $\|d_0 \otimes \pi \otimes \pi^{\text{base}} / \mu_{\mathcal{D}}\|_{\infty}$.

Assumption 5.1 is standard in sample complexity analysis [85, 78, 74]. Assumption 5.2 characterizes how well the dataset $\mathcal{D}$ covers the policy π to compete. To achieve provable sample efficiency, we only require that $\mathcal{D}$ covers the target policy π , a weak partial coverage style assumption for theoretical analysis. To illustrate it, when calling back to Figure 1, the data distribution therein well covers those nearly optimal responses under r^* , but does not sufficiently cover the responses with low r^* .

Under such a partial coverage data condition, however, human preference of responses $a \in \mathcal{A}$ that are not well covered by the dataset $\mathcal{D}$ can be poorly estimated, misleading the policy $\hat{\pi}$ to behave suboptimally if it is overoptimized (recall Figure 1). Fortunately, the following theorem shows that Algorithm 1 provably mitigates the overoptimization issue and achieves a finite-sample convergence of the suboptimality gap (2.4) competing with π . Proof is in Appendix D.

Theorem 5.3 (Suboptimality of Algorithm 1). *Taking the policy class Π as (2.3), supposing that Assumptions 5.1 and 5.2 hold, and assuming that the reward model class $\mathcal{R}$ has a finite ε -epsilon covering number under $\|\cdot\|_{\infty}$ -norm $\mathcal{N}_{\varepsilon}(\mathcal{R}, \|\cdot\|_{\infty}) < +\infty$ with $\varepsilon = (6 \cdot (1 + e^R) \cdot N)^{-1}$. Setting*

$$\eta = (1 + \exp(R))^{-2} \cdot \sqrt{24 \log(\mathcal{N}_{\varepsilon}(\mathcal{R}, \|\cdot\|_{\infty}) / \delta) / N}, \quad \beta = 1 / \sqrt{N}$$

in Algorithm 1. Then the output policy $\hat{\pi}$ of Algorithm 1 satisfies that with probability at least $1 - \delta$,

$$\text{Gap}^{\pi}(\hat{\pi}) \leq \frac{\sqrt{6}(1 + \exp(R))^2 ((C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}}))^2 + 1)\iota + 4\mathbb{E}_{x \sim d_0} [\text{KL}(\pi(\cdot|x) \| \pi^{\text{ref}}(\cdot|x))]}{4\sqrt{N}},$$

where $\iota = \sqrt{\log(\mathcal{N}_{\varepsilon}(\mathcal{R}, \|\cdot\|_{\infty}) / \delta)}$ with $\varepsilon = (6 \cdot (1 + e^R) \cdot N)^{-1}$. Here, N denotes the number of preference pairs in $\mathcal{D}$, R denotes the upper bound of the reward models, and the partial coverage coefficient $C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}})$ is defined in Assumption 5.2.

Remark 5.4 (Choice of the baseline policy). *As is indicated by Assumption 5.2, the least requirement is that π^{base} can be covered by the offline data distribution. E.g., we can take π^{base} as the distribution of the preferred responses in the data. In this case, the SFT loss in RPO explicitly regularizes the LLM to imitate the preferred responses. We choose this type of baseline policy in our experiments.*

5.2 Equivalence between Maximin and Minimax Objectives

Now we formally show that the theoretical target (maximin objective (3.2)) and the target for practical algorithm design (minimax objective (4.1)) are equivalent under certain regularity conditions. This can naturally extend the sample complexity of Algorithm 1 (Section 5.1) to that of minimax-based algorithms in Section 4, providing the theoretical guarantee for our practical algorithm design (RPO).

First, for notational simplicity, we denote the optimization target we investigate in Sections 3 and 4 as

$$\phi(\pi, r) := \eta \cdot \mathbb{E}_{\substack{x \sim d_0, a^1 \sim \pi(\cdot|x) \\ a^0 \sim \pi^{\text{base}}(\cdot|x)}} [r(x, a^1) - r(x, a^0) - \beta \cdot D_{\text{KL}}(\pi(\cdot|x) \| \pi^{\text{ref}}(\cdot|x))] + \mathcal{L}_{\mathcal{D}}(r), \quad (5.1)$$

for any $(\pi, r) \in \Pi \times \mathcal{R}$. Our result relies on the following assumptions on the reward model class $\mathcal{R}$.

Assumption 5.5 (Regularity of reward model class). *We assume the following things on the reward model class $\mathcal{R}$: (i) the space $\mathcal{R}$ is a compact topological space; (ii) the function ϕ in (5.1) is convex-like on $\mathcal{R}$, that is, for any $r_1, r_2 \in \mathcal{R}$ and $\alpha \in [0, 1]$, there exists $r_3 \in \mathcal{R}$ such that*

$$\phi(\pi, r_3) \leq \alpha \cdot \phi(\pi, r_1) + (1 - \alpha) \cdot \phi(\pi, r_2), \quad \forall \pi \in \Pi, \quad (5.2)$$

We note if $\mathcal{R}$ is convex, e.g., a linear model class [85, 71, 86] or more general the Lipschitz continuous model class $\mathcal{R}$, we can directly obtain that the function $\phi(\pi, \cdot)$ is *convex* over $\mathcal{R}$ (since the dependence on $r \in \mathcal{R}$ is linear terms plus a convex loss $\mathcal{L}_{\mathcal{D}}$ of $r \in \mathcal{R}$), which implies the convex-like property (5.2). Under Assumption 5.5, it holds that (Lemma E.1)

$$\max_{\pi \in \Pi} \min_{r \in \mathcal{R}} \phi(\pi, r) = \min_{r \in \mathcal{R}} \max_{\pi \in \Pi} \phi(\pi, r). \quad (5.3)$$

Furthermore, thanks to the KL-divergence regularization in ϕ which intuitively makes ϕ “strongly concave” over the policy π , (5.3) can give us the following stronger result, proved in Appendix E.1.

Theorem 5.6 (Formal statement of Theorem 4.2). *For the policy class Π defined in (2.3) and the reward model class $\mathcal{R}$ satisfying Assumption 5.5, consider the following policy defined as*

$$\pi_{\hat{r}} \in \operatorname{argmax}_{\pi \in \Pi} \phi(\hat{r}, \pi), \quad \text{where} \quad \hat{r} \in \operatorname{argmin}_{r \in \mathcal{R}} \max_{\pi \in \Pi} \phi(\pi, r). \quad (5.4)$$

Then the policy $\pi_{\hat{r}}$ also satisfies the maximin objective (3.2) of Algorithm 1, that is,

$$\pi_{\hat{r}} \in \operatorname{argmax}_{\pi \in \Pi} \min_{r \in \mathcal{R}} \phi(\pi, r).$$

Theorem 5.6 shows that the optimal KL-regularized policy associated with the reward model solving the minimax objective (3.2) also solves the maximin objective (i.e., objective (4.1) of Algorithm 1). This further allows us to extend our theoretical guarantee of Algorithm 1 (Section 5.1) to that of minimax-based algorithms, justifying our practical algorithm design in Section 4.

Corollary 5.7 (Suboptimality of minimax-based algorithm). *Take the policy class Π in (2.3) and the reward model class satisfying Assumption 5.5. Given any given policy π to compete, if Assumption 5.2 holds for π , then under the same choice of η and β as in Theorem 5.3, the policy $\pi_{\hat{r}}$ defined in (5.4) satisfies that $\operatorname{Gap}^{\pi}(\pi_{\hat{r}}) \leq \tilde{O}(1/\sqrt{N})$ with probability at least $1 - \delta$.*

6 Experiments

In this section, we provide a detailed empirical analysis of RPO to highlight the following four key points: (1) RPO is a flexible plug-in module and can be applied to different reference models. (2) RPO can alleviate the overoptimization issue in the training phase by giving more trust to the chosen responses in the preference dataset. (3) As a justification of our theoretical analysis, RPO achieves better alignment performance than DPO in in-data distribution. (4) RPO can also achieve consistently better performance in LLM benchmarks like MT-bench [83] and AlpacaEval 2.0 [20], which shows the potential of mitigating overoptimization for better generalization performance. The code for the experiments can be found in <https://github.com/YSLIU627/Regularized-Preference-Optimization/tree/master>.

Experiment setup. To show that RPO is a flexible plug-in module regardless of the reference model, we follow the training setup for two well-studied series of released chat models with around 7 billion parameters trained by DPO: Zephyr-7b-beta and Zephyr-7b-gemma [64] to implement RPO in beta and gemma series. Mirrored by their training configurations, we introduce how we select the reference model and the preference dataset for our training on these two series as follows. For the beta series, we use `mistral-7b-sft-beta` as the reference model π^{ref} . `mistral-7b-sft-beta` is a fine-tuned version of `Mistral-7b-v0.1` on the distilled version of the UltraChat dataset [17], which contains approximately 200k examples of multi-turn dialogues generated by GPT-3.5-TURBO. For the training preference dataset, we use Ultrafeedback Dataset [16], which consists of approximately 60k prompts. For the gemma series, we use `zephyr-7b-gemma-sft-v0.1` as our reference model π^{ref} . `zephyr-7b-gemma-sft-v0.1` is a fine-tuned version of `gemma-7b` on the Deita dataset [35], which involves around 10k distilled SFT data. For the training preference dataset, we use Argilla-DPO-Mix-7K Dataset [3], which is a mixture of multiple distilled public preference datasets. For simplicity, we denote Ref. (beta) as the reference model, DPO (beta) as the model trained by DPO, RPO (beta) as the model trained by RPO, all for the beta series. We use the same notations for the gemma series.

Practical implementation. According to Algorithm 2 and as discussed in Remark 5.4, we implement RPO by adding an SFT loss (log probability of chosen responses in the preference dataset) to the original DPO loss. By comparing the evaluation performance on the test split of the training

dataset, we select the hyperparameter η as 0.005 for both RPO (beta) and RPO (gemma). During the training of DPO and RPO, We keep the remaining hyperparameters including β , batch size, and learning rate to be the same for a fair comparison. Please see Appendix F.1 for a detailed training configuration.

RPO alleviates overoptimization. As mentioned in the introduction part, DPO is observed to have a significant and continuous decrease in log probability on chosen responses [26, 45] during training and we regard it as the consequence of overoptimization. Implied by our theory, overoptimization could arise when the model maximizes its own proxy reward formed on the responses less covered by the data. Due to the overoptimization, the model tends to disprefer the chosen responses as they are away from the maximizers of the proxy reward despite that some chosen responses are highly preferred by humans. Consistent with our theoretical conclusion, we empirically find that RPO can indeed alleviate overoptimization in DPO. During the training phase of both beta and gemma series, we observe that the log probability given by the RPO-trained model is notably higher than that given by the DPO-trained model for the chosen responses, which are shown in Fig. 1 and 2.

RPO improves the alignment ability in in-data distribution. For the in-data distribution evaluation, we select the 200 prompts (which are not used in the selection of η) in the test split of the training dataset to let the reference model, DPO, and RPO generate the response respectively. We choose GPT-4 to annotate the preference in the response pairs. Though we instruct GPT-4 to give an annotation among win, lose, and tie (please see the full prompt in Appendix F.2), GPT-4 may still give undesired annotations. Hence, we filter all the undesired annotations and collect 150 examples for evaluation. We report the pairwise win rate among Ref., RPO, and DPO in Table 1 for both the beta and gemma series. To show a more illustrative comparison between DPO and RPO, we provide the barplot to report the number of pairwise examples annotated by GPT-4 in Fig. 3 and Fig. 4. We observe that for both beta and gemma series, RPO has a better performance than DPO in terms of both RPO/DPO-SFT and RPO-DPO win rates. The performance improvement matches our theoretical results in Corollary 5.7, which shows the credit of the alleviation of overoptimization.

Table 1: Pairwise win rate (left vs. right) among RPO-trained model, DPO-trained model, and the reference model. Annotated by GPT-4, evaluations of beta and gemma series are made on the 150 examples of the test split of the Ultrafeedback and the Argilla-DPO-Mix-7K dataset, respectively.

Win rate (%)	RPO (beta)	Ref. (beta)	DPO (beta)	Win rate (%)	RPO (gemma)	Ref. (gemma)	DPO (gemma)
RPO (beta)	50.0	79.0	56.0	RPO (gemma)	50.0	71.7	54.0
Ref. (beta)	21.0	50.0	22.7	Ref. (gemma)	28.3	50.0	32.7
DPO (beta)	44.0	77.3	50.0	DPO (gemma)	46.0	67.3	50.0

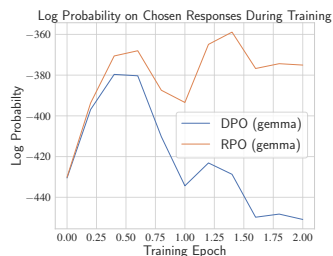


Figure 2: Log probability of the model for chosen responses during the training of RPO (gemma) and DPO (gemma).

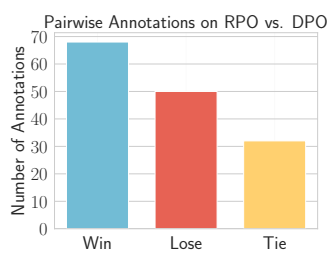


Figure 3: Pairwise annotations (by GPT-4) on RPO (beta) vs. DPO (beta) on the test split of the Ultrafeedback dataset.

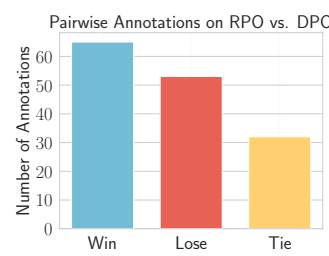


Figure 4: Pairwise annotations (by GPT-4) on RPO (gemma) vs. DPO (gemma) on the test split of the Argilla-DPO-Mix-7K dataset.

RPO consistently improves the benchmark performance. We further evaluate the reference model, RPO-trained model, DPO-trained model, and the officially released DPO-trained model for both beta and gemma series in two standard LLM chat benchmarks: MT-Bench and AlpacaEval 2.0. MT-Bench is a multi-turn benchmark that contains 160 questions across eight different domains of knowledge. The score for MT-Bench is evaluated by GPT-4 on a scale from 1 to 10. AlpacaEval 2.0 is a single-turn benchmark including 805 questions on different topics, mostly focused on helpfulness.

The metrics of AlpacaEval 2.0 are the win rate and Length-Control (LC) win rate compared with GPT-4 Preview (11/06), where the annotator is also GPT-4 Preview (11/06) and LC win rate is proposed to mitigate the length bias of GPT-4. The results are summarized in Table 2, which shows that RPO consistently exceeds the performance of all the competitors (DPO, Reference model, and the officially released model trained by DPO) on MT-Bench and AlpacaEval 2.0. We also provide additional results on the pairwise win rate for these two benchmarks in Appendix F.3 to illustrate the performance improvement. Finally, we remark that RPO is a flexible plug-in module and can steadily improve the benchmark performance without changing the original training configuration or accessing extra preference data. This also sheds light on the potential of mitigating overoptimization for better alignment and generalization performance.

Table 2: Results on MT-Bench scores and AlpacaEval 2.0. `zephyr-beta-7b` and `zephyr-gemma-7b` are the officially released models. win rates and Length-Control (LC) win rates in AlpacaEval 2.0 are evaluated by GPT-4 compared with GPT-4.

Model Name	MT-Bench Score	AlpacaEval 2.0		Model Name	MT-Bench Score	AlpacaEval 2.0	
		LC win rate (%)	win rate (%)			LC win rate (%)	win rate (%)
RPO (beta)	7.381	23.28	21.01	RPO (gemma)	7.916	15.51	13.85
Ref. (beta)	5.088	7.19	4.69	Ref. (gemma)	7.266	8.35	4.61
DPO (beta)	7.278	21.15	17.27	DPO (gemma)	7.688	15.36	13.69
<code>zephyr-beta-7b</code>	7.200	13.20	10.99	<code>zephyr-gemma-7b</code>	7.719	14.78	12.14

RPO also improves the math, reasoning, and coding abilities. In addition to the MT-Bench and AlpacaEval 2.0 benchmarks, we introduce more benchmarks on the math, reasoning, and coding tasks for evaluations of the RPO algorithm. Specifically, we choose the Grade School Math 8K (GSM8K) [14], AI2 Reasoning Challenge (ARC) [13], and Mostly Basic Python Programming (MBPP) [4] to measure math, reasoning, and coding abilities of the model trained by RPO, respectively. To due space limitation, we refer the readers to Appendix G for the setups and results of these experiments.

7 Conclusions

This paper proposes a new algorithm that provably mitigates reward overoptimization in RLHF. We establish its finite-sample convergence under a partial coverage style data condition, and provide an equivalent practical implementation, RPO. As a flexible plug-in module, RPO exhibits consistent improvement over the DPO baseline and effectively mitigates overoptimization. Future work includes extending our idea of algorithm design to online (iterative) RLHF where preference data are collected and updated iteratively during LLM fine-tuning. We give more detailed discussions in Appendix B.

Acknowledgement

Zhaoran Wang acknowledges National Science Foundation (Awards 2048075, 2008827, 2015568, 1934931), Simons Institute (Theory of Reinforcement Learning), Amazon, J.P. Morgan, and Two Sigma for their supports. The authors thank Junyan Zhang on valuable discussions on the equivalence between the min-max and max-min optimization.

References

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023. 16
- [2] Anthropic. Introducing claude. <https://www.anthropic.com/news/introducing-claude>, 2023. 16
- [3] argill. argilla-dpo-mix-7k. <https://huggingface.co/datasets/argilla/dpo-mix-7k>. 3, 8
- [4] Jacob Austin, Augustus Odena, Maxwell Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie Cai, Michael Terry, Quoc Le, et al. Program synthesis with large language models. *arXiv preprint arXiv:2108.07732*, 2021. 10

- [5] Mohammad Gheshlaghi Azar, Mark Rowland, Bilal Piot, Daniel Guo, Daniele Calandriello, Michal Valko, and Rémi Munos. A general theoretical paradigm to understand learning from human preferences. *arXiv preprint arXiv:2310.12036*, 2023. 16
- [6] Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*, 2022. 16
- [7] Viktor Bengs, Róbert Busa-Fekete, Adil El Mesaoudi-Paul, and Eyke Hüllermeier. Preference-based online learning with dueling bandits: A survey. *Journal of Machine Learning Research*, 22(7):1–108, 2021. 16
- [8] Jose Blanchet, Miao Lu, Tong Zhang, and Han Zhong. Double pessimism is provably efficient for distributionally robust offline reinforcement learning: Generic algorithm and robust partial coverage. *Advances in Neural Information Processing Systems*, 36, 2024. 16
- [9] Ralph Allan Bradley and Milton E Terry. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952. 3
- [10] Stephen Casper, Xander Davies, Claudia Shi, Thomas Krendl Gilbert, Jérémy Scheurer, Javier Rando, Rachel Freedman, Tomasz Korbak, David Lindner, Pedro Freire, et al. Open problems and fundamental limitations of reinforcement learning from human feedback. *arXiv preprint arXiv:2307.15217*, 2023. 1, 17
- [11] Xiaoyu Chen, Han Zhong, Zhuoran Yang, Zhaoran Wang, and Liwei Wang. Human-in-the-loop: Provably efficient preference-based reinforcement learning with general function approximation. In Kamalika Chaudhuri, Stefanie Jegelka, Le Song, Csaba Szepesvari, Gang Niu, and Sivan Sabato, editors, *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, pages 3773–3793. PMLR, 17–23 Jul 2022. 16
- [12] Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. Deep reinforcement learning from human preferences. *Advances in neural information processing systems*, 30, 2017. 1, 16
- [13] Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. Think you have solved question answering? try arc, the ai2 reasoning challenge. *arXiv preprint arXiv:1803.05457*, 2018. 10
- [14] Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, et al. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*, 2021. 10
- [15] Thomas Coste, Usman Anwar, Robert Kirk, and David Krueger. Reward model ensembles help mitigate overoptimization. *arXiv preprint arXiv:2310.02743*, 2023. 17
- [16] Ganqu Cui, Lifan Yuan, Ning Ding, Guanming Yao, Wei Zhu, Yuan Ni, Guotong Xie, Zhiyuan Liu, and Maosong Sun. Ultrafeedback: Boosting language models with high-quality feedback. *arXiv preprint arXiv:2310.01377*, 2023. 3, 8
- [17] Ning Ding, Yulin Chen, Bokai Xu, Yujia Qin, Zhi Zheng, Shengding Hu, Zhiyuan Liu, Maosong Sun, and Bowen Zhou. Enhancing chat language models by scaling high-quality instructional conversations. *arXiv preprint arXiv:2305.14233*, 2023. 8
- [18] Hanze Dong, Wei Xiong, Deepanshu Goyal, Yihan Zhang, Winnie Chow, Rui Pan, Shizhe Diao, Jipeng Zhang, Kashun Shum, and Tong Zhang. Raft: Reward ranked finetuning for generative foundation model alignment. *arXiv preprint arXiv:2304.06767*, 2023. 16
- [19] Yihan Du, Anna Winnicki, Gal Dalal, Shie Mannor, and R Srikant. Exploration-driven policy optimization in rlhf: Theoretical insights on efficient data utilization. *arXiv preprint arXiv:2402.10342*, 2024. 16

- [20] Yann Dubois, Balázs Galambosi, Percy Liang, and Tatsunori B Hashimoto. Length-controlled alpacaeval: A simple way to debias automatic evaluators. *arXiv preprint arXiv:2404.04475*, 2024. 8
- [21] Jacob Eisenstein, Chirag Nagpal, Alekh Agarwal, Ahmad Beirami, Alex D’Amour, DJ Dvijotham, Adam Fisch, Katherine Heller, Stephen Pfohl, Deepak Ramachandran, et al. Helping or herding? reward model ensembles mitigate but do not eliminate reward hacking. *arXiv preprint arXiv:2312.09244*, 2023. 17
- [22] Logan Engstrom, Andrew Ilyas, Shibani Santurkar, Dimitris Tsipras, Firdaus Janoos, Larry Rudolph, and Aleksander Madry. Implementation matters in deep policy gradients: A case study on ppo and trpo. *arXiv preprint arXiv:2005.12729*, 2020. 16
- [23] Ky Fan. Minimax theorems. *Proceedings of the National Academy of Sciences*, 39(1):42–47, 1953. 25
- [24] Zuyue Fu, Zhengling Qi, Zhaoran Wang, Zhuoran Yang, Yanxun Xu, and Michael R Kosorok. Offline reinforcement learning with instrumental variables in confounded markov decision processes. *arXiv preprint arXiv:2209.08666*, 2022. 16
- [25] Leo Gao, John Schulman, and Jacob Hilton. Scaling laws for reward model overoptimization. In *International Conference on Machine Learning*, pages 10835–10866. PMLR, 2023. 1, 17
- [26] Jiwoo Hong, Noah Lee, and James Thorne. Orpo: Monolithic preference optimization without reference model. *arXiv preprint arXiv:2403.07691*, 2024. 2, 9, 16
- [27] Jian Hu, Xibin Wu, Weixun Wang, Xianyu, Dehao Zhang, and Yu Cao. Openrlhf: An easy-to-use, scalable and high-performance rlhf framework. *arXiv preprint arXiv:2405.11143*, 2024. 28
- [28] Haozhe Ji, Cheng Lu, Yilin Niu, Pei Ke, Hongning Wang, Jun Zhu, Jie Tang, and Minlie Huang. Towards efficient and exact optimization of language model alignment. *arXiv preprint arXiv:2402.00856*, 2024. 16
- [29] Ying Jin, Zhuoran Yang, and Zhaoran Wang. Is pessimism provably efficient for offline rl? In *International Conference on Machine Learning*, pages 5084–5096. PMLR, 2021. 16
- [30] Chenliang Li, Siliang Zeng, Zeyi Liao, Jiaxiang Li, Dongyeop Kang, Alfredo Garcia, and Mingyi Hong. Joint demonstration and preference learning improves policy alignment with human feedback. *arXiv preprint arXiv:2406.06874*, 2024. 16
- [31] Zihao Li, Zhuoran Yang, and Mengdi Wang. Reinforcement learning with human feedback: Learning dynamic choices via pessimism. *arXiv preprint arXiv:2305.18438*, 2023. 16
- [32] Xize Liang, Chao Chen, Jie Wang, Yue Wu, Zhihang Fu, Zhihao Shi, Feng Wu, and Jieping Ye. Robust preference optimization with provable noise tolerance for llms. *arXiv preprint arXiv:2404.04102*, 2024. 16
- [33] Luofeng Liao, Zuyue Fu, Zhuoran Yang, Yixin Wang, Mladen Kolar, and Zhaoran Wang. Instrumental variable value iteration for causal offline reinforcement learning. *arXiv preprint arXiv:2102.09907*, 2021. 16
- [34] Tianqi Liu, Yao Zhao, Rishabh Joshi, Misha Khalman, Mohammad Saleh, Peter J Liu, and Jialu Liu. Statistical rejection sampling improves preference optimization. *arXiv preprint arXiv:2309.06657*, 2023. 16, 32
- [35] Wei Liu, Weihao Zeng, Keqing He, Yong Jiang, and Junxian He. What makes good data for alignment? a comprehensive study of automatic data selection in instruction tuning. *arXiv preprint arXiv:2312.15685*, 2023. 8
- [36] Zhihan Liu, Miao Lu, Zhaoran Wang, Michael Jordan, and Zhuoran Yang. Welfare maximization in competitive equilibrium: Reinforcement learning for markov exchange economy. In *International Conference on Machine Learning*, pages 13870–13911. PMLR, 2022. 16

- [37] Zhihan Liu, Miao Lu, Wei Xiong, Han Zhong, Hao Hu, Shenao Zhang, Sirui Zheng, Zhuoran Yang, and Zhaoran Wang. Maximize to explore: One objective function fusing estimation, planning, and exploration. *Advances in Neural Information Processing Systems*, 36, 2024. [22](#)
- [38] Zhihan Liu, Yufeng Zhang, Zuyue Fu, Zhuoran Yang, and Zhaoran Wang. Learning from demonstration: Provably efficient adversarial policy imitation with linear function approximation. In *International conference on machine learning*, pages 14094–14138. PMLR, 2022. [16](#)
- [39] Miao Lu, Yifei Min, Zhaoran Wang, and Zhuoran Yang. Pessimism in the face of confounders: Provably efficient offline reinforcement learning in partially observable markov decision processes. *arXiv preprint arXiv:2205.13589*, 2022. [16](#)
- [40] Eric J Michaud, Adam Gleave, and Stuart Russell. Understanding learned reward functions. *arXiv preprint arXiv:2012.05862*, 2020. [1](#), [17](#)
- [41] Ted Moskowitz, Aaditya K Singh, DJ Strouse, Tuomas Sandholm, Ruslan Salakhutdinov, Anca D Dragan, and Stephen McAleer. Confronting reward model overoptimization with constrained rlhf. *arXiv preprint arXiv:2310.04373*, 2023. [17](#)
- [42] Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744, 2022. [1](#), [2](#), [3](#), [4](#), [16](#)
- [43] Aldo Pacchiano, Aaditya Saha, and Jonathan Lee. Dueling rl: reinforcement learning with trajectory preferences. *arXiv preprint arXiv:2111.04850*, 2021. [16](#)
- [44] Arka Pal, Deep Karkhanis, Samuel Dooley, Manley Roberts, Siddhartha Naidu, and Colin White. Smaug: Fixing failure modes of preference optimisation with dpo-positive. *arXiv preprint arXiv:2402.13228*, 2024. [16](#)
- [45] Rafael Rafailov, Joey Hejna, Ryan Park, and Chelsea Finn. From r to q^* : Your language model is secretly a q -function. *arXiv preprint arXiv:2404.12358*, 2024. [2](#), [9](#)
- [46] Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36, 2023. [2](#), [3](#), [5](#), [6](#), [16](#), [32](#)
- [47] Paria Rashidinejad, Banghua Zhu, Cong Ma, Jiantao Jiao, and Stuart Russell. Bridging offline reinforcement learning and imitation learning: A tale of pessimism. *Advances in Neural Information Processing Systems*, 34:11702–11716, 2021. [16](#)
- [48] Paria Rashidinejad, Hanlin Zhu, Kunhe Yang, Stuart Russell, and Jiantao Jiao. Optimal conservative offline rl with general function approximation via augmented lagrangian. *arXiv preprint arXiv:2211.00716*, 2022. [16](#)
- [49] Mathieu Rita, Florian Strub, Rahma Chaabouni, Paul Michel, Emmanuel Dupoux, and Olivier Pietquin. Countering reward over-optimization in llm with demonstration-guided reinforcement learning. *arXiv preprint arXiv:2404.19409*, 2024. [17](#)
- [50] Corby Rosset, Ching-An Cheng, Arindam Mitra, Michael Santacrose, Ahmed Awadallah, and Tengyang Xie. Direct nash optimization: Teaching language models to self-improve with general preferences. *arXiv preprint arXiv:2404.03715*, 2024. [16](#)
- [51] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017. [16](#)
- [52] Tianhao Shen, Renren Jin, Yufei Huang, Chuang Liu, Weilong Dong, Zishan Guo, Xinwei Wu, Yan Liu, and Deyi Xiong. Large language model alignment: A survey. *arXiv preprint arXiv:2309.15025*, 2023. [16](#)
- [53] Laixi Shi and Yuejie Chi. Distributionally robust model-based offline reinforcement learning with near-optimal sample complexity. *Journal of Machine Learning Research*, 25(200):1–91, 2024. [16](#)

- [54] Laixi Shi, Gen Li, Yuting Wei, Yuxin Chen, and Yuejie Chi. Pessimistic q-learning for offline reinforcement learning: Towards optimal sample complexity. In *International conference on machine learning*, pages 19967–20025. PMLR, 2022. 16
- [55] Hao Sun. Reinforcement learning in the era of llms: What is essential? what is needed? an rl perspective on rlhf, prompting, and beyond. *arXiv preprint arXiv:2310.06147*, 2023. 16
- [56] Hao Sun, Alex James Chan, Nabeel Seedat, Alihan Hüyük, and Mihaela van der Schaar. When is off-policy evaluation (reward modeling) useful in contextual bandits? a data-centric perspective. *Journal of Data-centric Machine Learning Research*, 2024. 17
- [57] Hao Sun and Mihaela van der Schaar. Inverse-rlignment: Inverse reinforcement learning from demonstrations for llm alignment. *arXiv preprint arXiv:2405.15624*, 2024. 16
- [58] Fahim Tajwar, Anikait Singh, Archit Sharma, Rafael Rafailov, Jeff Schneider, Tengyang Xie, Stefano Ermon, Chelsea Finn, and Aviral Kumar. Preference fine-tuning of llms should leverage suboptimal, on-policy data. *arXiv preprint arXiv:2404.14367*, 2024. 16, 17
- [59] Yunhao Tang, Zhaohan Daniel Guo, Zeyu Zheng, Daniele Calandriello, Rémi Munos, Mark Rowland, Pierre Harvey Richemond, Michal Valko, Bernardo Ávila Pires, and Bilal Piot. Generalized preference optimization: A unified approach to offline alignment. *arXiv preprint arXiv:2402.05749*, 2024. 16
- [60] Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B Hashimoto. Stanford alpaca: An instruction-following llama model, 2023. 18
- [61] Gemini Team, Rohan Anil, Sebastian Borgeaud, Yonghui Wu, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan Schalkwyk, Andrew M Dai, Anja Hauth, et al. Gemini: a family of highly capable multimodal models. *arXiv preprint arXiv:2312.11805*, 2023. 16
- [62] Jeremy Tien, Jerry Zhi-Yang He, Zackory Erickson, Anca D Dragan, and Daniel S Brown. Causal confusion and reward misidentification in preference-based reward learning. *arXiv preprint arXiv:2204.06601*, 2022. 1, 17
- [63] Lewis Tunstall, Edward Beeching, Nathan Lambert, Nazneen Rajani, Shengyi Huang, Kashif Rasul, Alexander M. Rush, and Thomas Wolf. The alignment handbook. <https://github.com/huggingface/alignment-handbook>, 2023. 26
- [64] Lewis Tunstall, Edward Beeching, Nathan Lambert, Nazneen Rajani, Kashif Rasul, Younes Belkada, Shengyi Huang, Leandro von Werra, Clémentine Fourrier, Nathan Habib, et al. Zephyr: Direct distillation of lm alignment. *arXiv preprint arXiv:2310.16944*, 2023. 3, 8, 32
- [65] Masatoshi Uehara and Wen Sun. Pessimistic model-based offline reinforcement learning under partial coverage. In *International Conference on Learning Representations*. 16
- [66] Peiyi Wang, Lei Li, Liang Chen, Dawei Zhu, Binghuai Lin, Yunbo Cao, Qi Liu, Tianyu Liu, and Zhifang Sui. Large language models are not fair evaluators. *arXiv preprint arXiv:2305.17926*, 2023. 26
- [67] Yuanhao Wang, Qinghua Liu, and Chi Jin. Is rlhf more difficult than standard rl? a theoretical perspective. *Advances in Neural Information Processing Systems*, 36, 2023. 16
- [68] Yue Wu, Zhiqing Sun, Huizhuo Yuan, Kaixuan Ji, Yiming Yang, and Quanquan Gu. Self-play preference optimization for language model alignment. *arXiv preprint arXiv:2405.00675*, 2024. 16
- [69] Tengyang Xie, Ching-An Cheng, Nan Jiang, Paul Mineiro, and Alekh Agarwal. Bellman-consistent pessimism for offline reinforcement learning. *Advances in neural information processing systems*, 34:6683–6694, 2021. 16
- [70] Tengyang Xie, Nan Jiang, Huan Wang, Caiming Xiong, and Yu Bai. Policy finetuning: Bridging sample-efficient offline and online reinforcement learning. *Advances in neural information processing systems*, 34:27395–27407, 2021. 16

- [71] Wei Xiong, Hanze Dong, Chenlu Ye, Han Zhong, Nan Jiang, and Tong Zhang. Gibbs sampling from human feedback: A provable kl-constrained framework for rlhf. *arXiv preprint arXiv:2312.11456*, 2023. 4, 8, 16, 17, 32
- [72] Wei Xiong, Han Zhong, Chengshuai Shi, Cong Shen, Liwei Wang, and Tong Zhang. Nearly minimax optimal offline reinforcement learning with linear function approximation: Single-agent mdp and markov game. *arXiv preprint arXiv:2205.15512*, 2022. 16
- [73] Shusheng Xu, Wei Fu, Jiaxuan Gao, Wenjie Ye, Weilin Liu, Zhiyu Mei, Guangju Wang, Chao Yu, and Yi Wu. Is dpo superior to ppo for llm alignment? a comprehensive study. *arXiv preprint arXiv:2404.10719*, 2024. 18
- [74] Chenlu Ye, Wei Xiong, Yuheng Zhang, Nan Jiang, and Tong Zhang. A theoretical analysis of nash learning from human feedback under general kl-regularized preference. *arXiv preprint arXiv:2402.07314*, 2024. 7, 16
- [75] Ming Yin and Yu-Xiang Wang. Towards instance-optimal offline reinforcement learning with pessimism. *Advances in neural information processing systems*, 34:4065–4078, 2021. 16
- [76] Yisong Yue, Josef Broder, Robert Kleinberg, and Thorsten Joachims. The k-armed dueling bandits problem. *Journal of Computer and System Sciences*, 78(5):1538–1556, 2012. 16
- [77] Wenhao Zhan, Baihe Huang, Audrey Huang, Nan Jiang, and Jason Lee. Offline reinforcement learning with realizability and single-policy concentrability. In *Conference on Learning Theory*, pages 2730–2775. PMLR, 2022. 16
- [78] Wenhao Zhan, Masatoshi Uehara, Nathan Kallus, Jason D Lee, and Wen Sun. Provable offline preference-based reinforcement learning. In *The Twelfth International Conference on Learning Representations*, 2023. 4, 7, 16, 17
- [79] Wenhao Zhan, Masatoshi Uehara, Wen Sun, and Jason D Lee. How to query human feedback efficiently in rl? *arXiv preprint arXiv:2305.18505*, 2023. 16
- [80] Ruiqi Zhang, Licong Lin, Yu Bai, and Song Mei. Negative preference optimization: From catastrophic collapse to effective unlearning. *arXiv preprint arXiv:2404.05868*, 2024. 16
- [81] Xiaoying Zhang, Jean-Francois Ton, Wei Shen, Hongning Wang, and Yang Liu. Overcoming reward overoptimization via adversarial policy optimization with lightweight uncertainty estimation. *arXiv preprint arXiv:2403.05171*, 2024. 17
- [82] Yao Zhao, Rishabh Joshi, Tianqi Liu, Misha Khalman, Mohammad Saleh, and Peter J Liu. Slic-hf: Sequence likelihood calibration with human feedback. *arXiv preprint arXiv:2305.10425*, 2023. 16
- [83] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in Neural Information Processing Systems*, 36, 2024. 8
- [84] Han Zhong, Guhao Feng, Wei Xiong, Li Zhao, Di He, Jiang Bian, and Liwei Wang. Dpo meets ppo: Reinforced token optimization for rlhf. *arXiv preprint arXiv:2404.18922*, 2024. 16
- [85] Banghua Zhu, Jiantao Jiao, and Michael I Jordan. Principled reinforcement learning with human feedback from pairwise or k -wise comparisons. *arXiv preprint arXiv:2301.11270*, 2023. 7, 8, 16, 17
- [86] Banghua Zhu, Michael I Jordan, and Jiantao Jiao. Iterative data smoothing: Mitigating reward overfitting and overoptimization in rlhf. *arXiv preprint arXiv:2401.16335*, 2024. 2, 8, 17
- [87] 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. *arXiv preprint arXiv:1909.08593*, 2019. 1, 16

A Related Works

In the following, we relate our work to recent lines of RLHF research on both theory and practice sides. We also review related works on reward hacking and overoptimization in RLHF.

RLHF: algorithm design. The technique of RLHF [12, 87, 42, 6, 18, 55] has recently demonstrated its great importance in building the state-of-the-art LLMs, including ChatGPT [1], Gemini [61], Claude [2]. In the RLHF pipeline, the LLM is fine-tuned towards maximizing a learned reward model for better alignment [52, 57] with human preference using RL algorithms such as Proximal Policy Optimization (PPO; [51]). Meanwhile, PPO-style algorithm is also known for its instability, sample-inefficiency, and especially, a high demand for proper hyperparameter tuning [22]. This thus casts prohibitive computational cost to make the most effectiveness of PPO-based RLHF methods to align LLMs, especially for the open-source community.

Given that, further research on RLHF has explored various alternatives to PPO-based methods, with the most popular approach being the direct preference optimization method [82, 46], which skips the reward learning phase and directly optimizes the LLM to align it with the human preference. Our practical implementation (RPO) also harnesses the wisdom of reward-LLM equivalence to avoid explicit reward learning followed by PPO training.

Besides the original DPO algorithm [46], ever since it popularizing the direct preference learning style method, variants of the direct preference learning approach are proposed, including but not limited to [34, 5, 71, 59, 28, 74, 44, 26, 50, 32, 80, 58, 68, 30]. Each of them aims to address further challenges of direct preference learning from varying perspectives. Specifically, the algorithm proposed by [44, 26] share similar algorithmic components as RPO proposed in this work. Both work consider SFT style regularization during preference optimization. However, theoretical understanding of how SFT loss can help alignment remains unknown. In contrast, we provide theoretical justifications to the SFT loss as an implicit adversarial regularizer that provably mitigates overoptimization in preference learning.

RLHF: theoretical investigation. Initiated from the literature of dueling bandits and dueling RL [76, 7, 43], recent success of RLHF in fine-tuning LLMs also motivates a long line of research to investigate the theoretical foundations of RLHF under different settings [11, 85, 78, 79, 67, 31, 71, 74, 19, 84], aiming to propose provably sample-efficient algorithms to learn a human-reward-maximizing policy from human preference signals. Our theoretical study of RLHF falls into the paradigm of offline learning from a pre-collected preference dataset, and is mostly related to the work of [85, 78, 31, 71, 74]. In this setup, the main challenge is to address the overoptimization issues due to human reward uncertainty and distributional shifts when only a fixed dataset is available. In the sequel, we compare our work with them in more detail.

Existing theoretical work on provably sample-efficient offline RLHF typically suffers from two drawbacks: they are either restricted to the linear function approximations setting [85, 71] which is far from the practical situations, or are generally unable to be implemented in the LLM experiments. Typically, to encompass the pessimistic principle in the face of uncertainty, the existing literature proposes to return the optimal policy against either an estimated reward model plus a structure-aware reward uncertainty penalty [71] or the most pessimistic reward model inside a confidence region [85, 78]. Both of these two types of method involve intractable components for implementation and needs for additional algorithmic design to approximate the theoretical algorithm in practice. In contrast, our theory works in the context of general function approximations while being friendly to be implemented. Finally, we remark that, while our study focuses on the standard Bradley-Terry model of human preference with general reward function approximations, the work of [74] further considers a general human preference model. But it remains unknown how their algorithms can be efficiently implemented in practice. It serves as an interesting direction to extend our technique to RLHF with general reward model and device new practical algorithms.

Finally, we mention that the algorithm design of RPO is also related to the “pessimism” principle in the standard offline RL literature. It proposes to maintain a pessimistic estimate of the policy values or constrain the policy not to take unseen actions in the data to handle the challenge of the insufficient coverage of the dataset, e.g., [29, 65, 69, 70, 47, 75, 72, 36, 54, 77, 39, 48, 53, 8, 38, 33, 24]. In contrast, we consider the offline RLHF problem and the techniques to obtain the objective of the

RPO algorithm (see Section 4) along with its sample complexity analysis are new and different from these works.

Reward hacking and overoptimization in RLHF for LLM. As is discussed in the introduction, the challenge of reward hacking or overoptimization may prevent the successful alignment of LLMs, degenerating the performance of an LLM because of maximizing an imperfect, overfitted, and misgeneralized proxy reward learned from the finite data [40, 62, 25, 10]. Efforts have been made to mitigate this fundamental issue through the perspective of theory, e.g., [85, 71, 86], and practice, e.g., [15, 21, 41, 81, 49, 56]. Our approach starts from the theoretical insights of handling inherent uncertainty in learning human preference from finite data, while being surprisingly easy to implement.

B Limitations and Future Works

One limitation of the current work is that we focus on the setting of offline RLHF where only a fixed preference dataset is available. Recent RLHF research has shown great potential of using iterative methods for LLM alignment with multiple rounds of preference data collection and tuning [71, 58].

Future works include extending our idea of theoretical algorithm design and analysis to the iterative RLHF setup where further preference data can be collected. Also, since our practical algorithm RPO is a plug-in module that effectively mitigates overoptimization and improves alignment performance, it serves as an exciting direction to combine it with explorative preference data collecting mechanism in iterative RLHF to further boost the performance of LLM alignment.

C Further Discussions

Discussions on Algorithm 1 and Theorem 5.3. We compare our theory with [71] and [78].

Remark C.1 (Comparison with [71]). *Another theoretical work on RLHF [71] explicitly models the KL-regularization between the target policy and the reference policy in the learning objective, referred to as the KL-regularized contextual bandit. This means that their metric becomes the KL-regularized expected reward. In contrast, here we put the KL-regularization as a component of our algorithm design, but we still keep the metric as the expected reward (2.2). Therefore our theory in Section 5.1 directly reveals how the learned policy performs in terms of the expected reward compared to any given target policy (which can be a stochastic policy).*

Remark C.2 (Comparison with [78]). *We remark that in the work of [78], they also mentioned a maximin object similar to (3.2) for offline preference-based RL as a complementary to their theoretical algorithm. However, the sample complexity of the maximin-style algorithm they presented is unknown, while we provide finite sample convergence result for Algorithm 1 in Section 5. Furthermore, our objective (3.2) features another KL-regularization term, which is essential for the proposal of our new practical algorithm design for aligning LLM in Section 4.*

Discussions on the partial coverage assumption (Assumption 5.2). A sufficient condition to make this partial coverage condition (Assumption 5.2) hold is that the distribution of the offline dataset, which is $\mu_{\mathcal{D}}$, can well cover the joint distribution of $(a^1, a^0) \sim (\pi, \pi^{\text{base}})$. Here to discuss focus on $\pi^{\text{base}} = \pi^{\text{chosen}}$ as we adopted in the experiment part.

First, we clarify that the offline dataset distribution $\mu_{\mathcal{D}}$ is not simply $(a^1, a^0) \sim (\pi^{\text{unchosen}}, \pi^{\text{chosen}})$, since according to our definition (see Section 2) whether a^1 or a^0 is chosen is random and is determined by $y \in \{0, 1\}$ obeying the BT model. Thus, $(a^1, a^0) \sim \mu_{\mathcal{D}}$ can be interpreted as a mixture of $(\pi^{\text{unchosen}}, \pi^{\text{chosen}})$ and $(\pi^{\text{chosen}}, \pi^{\text{unchosen}})$. This mixture probability would not be too small as long as the quality of (a^1, a^0) does not vary too much, i.e., both of them are possible to be chosen, which is the case in practice. As a result, in the offline data distribution $(a^1, a^0) \sim \mu_{\mathcal{D}}$, both a^1 and a^0 partly comes from the chosen distribution π^{chosen} .

Then in order for $\mu_{\mathcal{D}}$ to cover the joint distribution of $(a^1, a^0) \sim (\pi, \pi^{\text{base}})$, it suffices to argue that π^{chosen} can cover the target policy π , which is then reduced back to the traditional coverage condition. Thus our assumption essentially requires that π^{chosen} well covers and only needs to cover the target policy π . This coincides with the spirit of the minimal data assumption in offline RL theory, i.e., the so-called partial coverage condition.

On the relationship between observed chosen probability and reward overoptimization. First, we note that the actions and their chosen probabilities can be interpreted as a proxy of analyzing the underlying (estimated) reward model $\hat{r}$ due to the representation $\pi_{\hat{r}}(a|x) \propto \pi^{\text{ref}}(a|x) \exp(\beta^{-1} \hat{r}(x, a))$. Analyzing the (log) probabilities of the actions can be utilized to detect the mitigation of overoptimization, because according to the representation, an overestimated reward of a poor action would result in a higher probability of choosing this action, and would also cause a decay in the probability of choosing other better actions (since the probabilities are normalized to 1).

To further showcase the ability of RPO to address overoptimization (through the lense of probability), consider the following theoretical example with only one state and three actions [73] where we can track everything clearly. It has three actions a, b, c with $R^*(a) = 1, R^*(b) = 0.5, R^*(c) = 0$. The reference policy $\pi^{\text{ref}}(a) = \pi^{\text{ref}}(b) = 0.4, \pi^{\text{ref}}(c) = 0.1$, and the dataset consists of one data point $\mathcal{D} = (a, b, 1)$ (meaning action a is preferred in the data). Then an ideally solved DPO objective would be π_{DPO} as long as $\pi^{\text{DPO}}(b) = 0$, and the value of $\pi^{\text{DPO}}(a)$ can be arbitrarily chosen in $[0, 1]$. Thus a possible solution to DPO would be $\pi^{\text{DPO}}(a) = 0.5, \pi^{\text{DPO}}(b) = 0$, and by the normalizing condition $\pi^{\text{DPO}}(c) = 0.5$, which is undesirable since the action c has reward $R^*(c) = 0$. In contrast, solving the RPO objective would additionally require the maximization of $\pi_{\text{RPO}}(a)$ due to the SFT regularization term, and thus the solution is shifted towards $\pi_{\text{RPO}}(a) = 1, \pi_{\text{RPO}}(b) = \pi_{\text{RPO}}(c) = 0$, which is better than the DPO policy. Thus, RPO is able to prevent overoptimization towards poor actions that are less covered by the dataset (action c here), therefore resulting in a better policy.

About the relationships and distinctions between PTX loss in [60] and the SFT loss of RPO. The original PTX loss is an imitation loss calculated on the pretraining data. In contrast, the SFT loss in the RPO objective is an imitation loss calculated on the RLHF dataset. In more specific, our experiments use this SFT loss to imitate the chosen responses in the RLHF dataset. Thus the relationship is that they are both imitation loss which aims to mimic certain data distribution. The distinction is that they are calculated on different data sources. The SFT loss in the RPO objective naturally comes from our theoretical algorithm and provably serves as an important regularization term to mitigate overoptimization in offline RLHF.

About the computational complexity of the SFT loss gradient. According to the paragraph **Practical implementation** in Section 6, RPO adds an additional SFT loss (the log probability of the chosen labels in the preference dataset) on the original DPO loss, where we highlight that the SFT loss is actually an intermediate quantity in the calculation of the DPO loss. Hence, our proposed method does not incur any additional computation overhead compared with the vanilla DPO.

D Proofs for Sample Complexity Analysis

D.1 Proof of Theorem 5.3

Proof of Theorem 5.3. By definition, the suboptimality gap of $\hat{\pi}$ w.r.t. π is decomposed as following,

$$\begin{aligned}
 \text{Gap}^{\pi}(\hat{\pi}) &= \mathbb{E}_{x \sim d_0, a \sim \pi(\cdot|x)} [r^*(x, a)] - \mathbb{E}_{x \sim d_0, a \sim \hat{\pi}(\cdot|x)} [r^*(x, a)] \\
 &= \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{ref}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] \\
 &\quad - \eta^{-1} \cdot \min_{r \in \mathcal{R}} \left\{ \eta \cdot \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\} \\
 &\quad + \eta^{-1} \cdot \min_{r \in \mathcal{R}} \left\{ \eta \cdot \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\} \\
 &\quad - \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] \\
 &\quad + \beta \cdot \mathbb{E}_{x \sim d_0} \left[\text{KL}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) - \text{KL}(\hat{\pi}(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] \\
 &:= \text{Term (A)} + \text{Term (B)} + \text{Term (C)}, \tag{D.1}
 \end{aligned}$$

where in the above Term (A), Term (B), and Term (C) are abbreviations for

Term (A)

$$\begin{aligned}
 &= \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] \\
 &\quad - \eta^{-1} \cdot \min_{r \in \mathcal{R}} \left\{ \eta \cdot \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\},
 \end{aligned}$$

and

Term (B)

$$\begin{aligned}
 &= \eta^{-1} \cdot \min_{r \in \mathcal{R}} \left\{ \eta \cdot \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\} \\
 &\quad - \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right],
 \end{aligned}$$

and

$$\text{Term (C)} = \beta \cdot \mathbb{E}_{x \sim d_0} \left[\text{KL}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) - \text{KL}(\hat{\pi}(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right].$$

In the following, we analyze Term (A) and Term (B) respectively.

Upper bound Term (A). Notice that by the optimality of our choice of policy $\hat{\pi}$ in (3.2), we have

Term (A)

$$\begin{aligned}
&= \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\pi(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right] \quad (\text{D.2}) \\
&\quad - \eta^{-1} \cdot \min_{r \in \mathcal{R}} \left\{ \eta \cdot \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\} \\
&\leq \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{ref}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\pi(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right] \\
&\quad - \eta^{-1} \cdot \min_{r \in \mathcal{R}} \left\{ \eta \cdot \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\pi(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\} \\
&= \max_{r \in \mathcal{R}} \left\{ \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[(r^*(x, a^1) - r^*(x, a^0)) - (r(x, a^1) - r(x, a^0)) \right] - \eta^{-1} \cdot \mathcal{L}_{\mathcal{D}}(r) \right\},
\end{aligned}$$

where in the inequality we apply the optimality of the choice of policy $\hat{\pi}$ in (3.2).

Upper bound Term (B). For this term, we directly consider the following bound,

Term (B)

$$\begin{aligned}
&= \eta^{-1} \cdot \min_{r \in \mathcal{R}} \left\{ \eta \cdot \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{ref}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right] + \mathcal{L}_{\mathcal{D}}(r) \right\} \\
&\quad - \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right] \\
&\leq \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right] + \eta^{-1} \cdot \mathcal{L}_{\mathcal{D}}(r^*) \\
&\quad - \mathbb{E}_{x \sim d_0, a^1 \sim \hat{\pi}(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r^*(x, a^1) - r^*(x, a^0) - \beta \cdot \text{KL}(\hat{\pi}(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right] \\
&= \eta^{-1} \cdot \mathcal{L}_{\mathcal{D}}(r^*), \quad (\text{D.3})
\end{aligned}$$

where in the inequality we apply the fact that $r^* \in \mathcal{R}$ by Assumption 5.1.

Combining Term (A), Term (B), and Term (C). Now by (D.1), (D.2), and (D.3), we have that

$$\text{Gap}_{\beta}^{\pi}(\hat{\pi}) = \text{Term (A)} + \text{Term (B)} + \text{Term (C)} \quad (\text{D.4})$$

$$\begin{aligned}
&\leq \max_{r \in \mathcal{R}} \left\{ \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[(r^*(x, a^1) - r^*(x, a^0)) - (r(x, a^1) - r(x, a^0)) \right] + \eta^{-1} \cdot (\mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r)) \right\} \\
&\quad + \beta \cdot \mathbb{E}_{x \sim d_0} \left[\text{KL}(\pi(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) - \text{KL}(\hat{\pi}(\cdot|x) \| \pi^{\text{ref}}(\cdot|x)) \right].
\end{aligned}$$

In the following, we upper bound the right hand side of (D.4) via relating the MLE loss difference term to the reward difference term through a careful analysis of the preference model. On the one hand, we invoke Lemma D.1 to give an upper bound of the difference of the MLE loss as following, with probability at least $1 - \delta$ over random samples and $\varepsilon = (6 \cdot (1 + e^R) \cdot N)^{-1}$, for any reward model $r \in \mathcal{R}$, it holds that

$$\begin{aligned}
&\mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r) \\
&\leq -2 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[D_{\text{Hellinger}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \| \mathbb{P}_r(\cdot|x, a^1, a^0)) \right]
\end{aligned}$$

$$+ \frac{3}{N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right),$$

where we recall that we use the subscript r in $\mathbb{P}_r$ to emphasize the dependence of the probabilistic model on the reward model. Here $\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)$ denotes the ε -covering number of the reward model class and R is the upper bound on the reward functionss (Assumption 5.1). Now to facilitate the calculation, we lower bound the Hellinger distance by total variation (TV) distance as

$$D_{\text{Hellinger}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \geq D_{\text{TV}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)),$$

By the expression of the probability model $\mathbb{P}_r$, we can further write the TV distance above as

$$\begin{aligned} D_{\text{TV}}(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) &= \frac{1}{2} \cdot \left| \sigma(r^*(x, a^1) - r^*(x, a^0)) - \sigma(r(x, a^1) - r(x, a^0)) \right| \\ &\quad + \frac{1}{2} \cdot \left| \sigma(r^*(x, a^0) - r^*(x, a^1)) - \sigma(r(x, a^0) - r(x, a^1)) \right| \\ &= \left| \sigma(r^*(x, a^1) - r^*(x, a^0)) - \sigma(r(x, a^1) - r(x, a^0)) \right|, \end{aligned} \quad (\text{D.5})$$

where in the second equality we use the fact that $\sigma(-z) = 1 - \sigma(z)$. Now by Lemma D.2 and the condition that $r(x, a) \in [0, R]$ for any $(x, a, r) \in \mathcal{X} \times \mathcal{A} \times \mathcal{R}$ (Assumption 5.1), we know that

$$\begin{aligned} &\left| \sigma(r^*(x, a^1) - r^*(x, a^0)) - \sigma(r(x, a^1) - r(x, a^0)) \right| \\ &\geq \kappa \cdot \left| (r^*(x, a^1) - r^*(x, a^0)) - (r(x, a^1) - r(x, a^0)) \right|, \end{aligned}$$

where $\kappa = 1/(1 + \exp(R))^2$. As a result, the difference of the MLE loss is upper bounded by

$$\begin{aligned} \mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r) &\leq -2\kappa^2 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[\left| (r^*(x, a^1) - r^*(x, a^0)) - (r(x, a^1) - r(x, a^0)) \right|^2 \right] \\ &\quad + \frac{3}{N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right). \end{aligned} \quad (\text{D.6})$$

On the other hand, the reward difference term in (D.4), which is evaluated on actions from π and π^{base} , can be related to the reward difference evaluated on the data distribution $\mu_{\mathcal{D}}$ via Assumption 5.2, i.e.,

$$\begin{aligned} &\mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[(r^*(x, a^1) - r^*(x, a^0)) - (r(x, a^1) - r(x, a^0)) \right] \\ &\leq C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}}) \sqrt{\mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}} \left[\left| (r^*(x, a^1) - r^*(x, a^0)) - (r(x, a^1) - r(x, a^0)) \right|^2 \right]}. \end{aligned} \quad (\text{D.7})$$

Finally, combining (D.6), (D.7), and (D.4), denoting

$$\Delta_r := \sqrt{\mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}} \left[\left| (r^*(x, a^1) - r^*(x, a^0)) - (r(x, a^1) - r(x, a^0)) \right|^2 \right]},$$

we have that

$$\begin{aligned} \text{Gap}^\pi(\hat{\pi}) &\leq \max_{r \in \mathcal{R}} \left\{ C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}}) \cdot \Delta_r - 2\eta^{-1}\kappa^2 \cdot \Delta_r^2 \right\} + \frac{3}{\eta N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right) \\ &\quad + \beta \cdot \mathbb{E}_{x \sim d_0} \left[\text{KL}(\pi(\cdot|x) \|\pi^{\text{ref}}(\cdot|x)) - \text{KL}(\hat{\pi}(\cdot|x) \|\pi^{\text{ref}}(\cdot|x)) \right] \\ &\leq \frac{(C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}}))^2 \eta}{8\kappa^2} + \frac{3}{\eta N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right) \\ &\quad + \beta \cdot \mathbb{E}_{x \sim d_0} \left[\text{KL}(\pi(\cdot|x) \|\pi^{\text{ref}}(\cdot|x)) \right], \end{aligned}$$

where in the second inequality we use that fact that $az - bz^2 \leq a^2/(4b)$ for any $z \in \mathbb{R}$ and that KL-divergence is non-negative. Consequently, with the choice of

$$\eta = 2\sqrt{6} \cdot \sqrt{\frac{\log(\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)/\delta)}{N}}, \quad \beta = \frac{1}{\sqrt{N}}, \quad \kappa = \frac{1}{(1 + \exp(R))^2},$$

we conclude that with probability at least $1 - \delta$ and $\varepsilon = (6 \cdot (1 + e^R) \cdot N)^{-1}$,

$$\begin{aligned} & \text{Gap}^\pi(\hat{\pi}) \\ & \leq \frac{\sqrt{6}(1 + \exp(R))^2 \left((C_{\mu_{\mathcal{D}}}(\mathcal{R}; \pi, \pi^{\text{base}}))^2 + 1 \right) \iota + 4\mathbb{E}_{x \sim d_0} \left[\text{KL}(\pi(\cdot|x) \|\pi^{\text{ref}}(\cdot|x)) \right]}{4\sqrt{N}}, \end{aligned}$$

where we denote $\iota = \sqrt{\log(\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)/\delta)}$ with $\varepsilon = (6 \cdot (1 + e^R) \cdot N)^{-1}$. This finishes the proof of Theorem 5.3. $\square$

D.2 Technical Lemmas

Lemma D.1 (Uniform concentration). *Consider the MLE loss (3.1) and define the approximation error as $\varepsilon = (6 \cdot (1 + e^R) \cdot N)^{-1}$ where R is the upper bound on the reward functions (Assumption 5.2). Suppose that the reward model class $\mathcal{R}$ has a finite ε -covering number $\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty) < \infty$. Then for any $\delta < 1/e$ it holds with probability at least $1 - \delta$ that*

$$\begin{aligned} & \mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r) \\ & \leq -2 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[D_{\text{Hellinger}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \right] \\ & \quad + \frac{3}{N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right). \end{aligned}$$

Proof of Lemma D.1. For notational simplicity, we use $\mathcal{C}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)$ to denote an ε -cover of the reward model class $\mathcal{R}$ under the $\|\cdot\|_\infty$ -norm. It holds that $\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty) = |\mathcal{C}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)|$. First we invoke Proposition 5.3 of [37] to obtain a uniform concentration over the finite set of ε -cover $\mathcal{C}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)$. Specifically, with probability at least $1 - \delta$, for any $r \in \mathcal{C}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)$,

$$\begin{aligned} & \mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r) \\ & \leq -2 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[D_{\text{Hellinger}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \right] \\ & \quad + \frac{2}{N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right). \end{aligned} \tag{D.8}$$

Now for any reward model $r \in \mathcal{R}$, we take a $r^\dagger \in \mathcal{C}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)$ satisfying $\|r - r^\dagger\|_\infty \leq \varepsilon$. We have

$$\begin{aligned} & \mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r) \\ & = \mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r^\dagger) + \mathcal{L}_{\mathcal{D}}(r^\dagger) - \mathcal{L}_{\mathcal{D}}(r) \\ & \leq -2 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[D_{\text{Hellinger}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_{r^\dagger}(\cdot|x, a^1, a^0)) \right] \\ & \quad + \frac{2}{N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right) + \mathcal{L}_{\mathcal{D}}(r^\dagger) - \mathcal{L}_{\mathcal{D}}(r) \\ & \leq -2 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[D_{\text{Hellinger}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \right] \\ & \quad + \frac{2}{N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right) + \mathcal{L}_{\mathcal{D}}(r^\dagger) - \mathcal{L}_{\mathcal{D}}(r) \\ & \quad + 4 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[D_{\text{Hellinger}}^2(\mathbb{P}_{r^\dagger}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \right], \end{aligned} \tag{D.9}$$

where in the first inequality we use (D.8) for $r^\dagger$ and in the second inequality we utilize the triangular inequality for Hellinger distance. Therefore, it remains to upper bound the approximation error induced by $r^\dagger$. On the one hand, by the definition of $\mathcal{L}_{\mathcal{D}}$ in (3.1), we have that

$$\mathcal{L}_{\mathcal{D}}(r^\dagger) - \mathcal{L}_{\mathcal{D}}(r)$$

$$\begin{aligned}
&= \frac{1}{N} \sum_{i=1}^N y_i \cdot \log \left(\frac{\sigma(r(x_i, a_i^1) - r(x_i, a_i^0))}{\sigma(r^\dagger(x_i, a_i^1) - r^\dagger(x_i, a_i^0))} \right) \\
&\quad + \frac{1}{N} \sum_{i=1}^N (1 - y_i) \cdot \log \left(\frac{\sigma(r(x_i, a_i^0) - r(x_i, a_i^1))}{\sigma(r^\dagger(x_i, a_i^0) - r^\dagger(x_i, a_i^1))} \right).
\end{aligned}$$

Use the inequality that $\log(x) \leq x - 1$, we can further upper bound $\mathcal{L}_{\mathcal{D}}(r^\dagger) - \mathcal{L}_{\mathcal{D}}(r)$ by

$$\begin{aligned}
&\mathcal{L}_{\mathcal{D}}(r^\dagger) - \mathcal{L}_{\mathcal{D}}(r) \\
&\leq \frac{1}{N} \sum_{i=1}^N y_i \cdot \frac{\sigma(r(x_i, a_i^1) - r(x_i, a_i^0)) - \sigma(r^\dagger(x_i, a_i^1) - r^\dagger(x_i, a_i^0))}{\sigma(r^\dagger(x_i, a_i^1) - r^\dagger(x_i, a_i^0))} \\
&\quad + \frac{1}{N} \sum_{i=1}^N (1 - y_i) \cdot \frac{\sigma(r(x_i, a_i^0) - r(x_i, a_i^1)) - \sigma(r^\dagger(x_i, a_i^0) - r^\dagger(x_i, a_i^1))}{\sigma(r^\dagger(x_i, a_i^0) - r^\dagger(x_i, a_i^1))}.
\end{aligned}$$

Now since $\|r^\dagger - r\|_\infty \leq \varepsilon$ and $r^\dagger \in [0, R]$, invoking Lemma D.2, we can derive that

$$\begin{aligned}
\mathcal{L}_{\mathcal{D}}(r^\dagger) - \mathcal{L}_{\mathcal{D}}(r) &\leq \frac{1}{N} \sum_{i=1}^N \frac{|(r(x_i, a_i^1) - r(x_i, a_i^0)) - (r^\dagger(x_i, a_i^1) - r^\dagger(x_i, a_i^0))|}{(1 + e^R)^{-1}} \\
&\quad + \frac{1}{N} \sum_{i=1}^N \frac{|(r(x_i, a_i^0) - r(x_i, a_i^1)) - (r^\dagger(x_i, a_i^0) - r^\dagger(x_i, a_i^1))|}{(1 + e^R)^{-1}} \\
&\leq 4 \cdot \|r^\dagger - r\|_\infty \cdot (1 + e^R) \leq 4\varepsilon \cdot (1 + e^R). \tag{D.10}
\end{aligned}$$

On the other hand, we upper bound the hellinger distance between $\mathbb{P}_r$ and $\mathbb{P}_{r^\dagger}$, for any $(x, a^1, a^0) \in \mathcal{X} \times \mathcal{A} \times \mathcal{A}$,

$$\begin{aligned}
&D_{\text{Hellinger}}^2(\mathbb{P}_{r^\dagger}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \\
&\leq D_{\text{TV}}(\mathbb{P}_{r^\dagger}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \\
&= \left| \sigma(r^\dagger(x, a^1) - r^\dagger(x, a^0)) - \sigma(r(x, a^1) - r(x, a^0)) \right| \\
&\leq \left| (r^\dagger(x, a^1) - r^\dagger(x, a^0)) - (r(x, a^1) - r(x, a^0)) \right| \\
&\leq 2 \cdot \|r^\dagger - r\|_\infty \leq 2\varepsilon, \tag{D.11}
\end{aligned}$$

where the first inequality uses the fact that $D_{\text{Hellinger}}^2 \leq D_{\text{TV}}$, the equality uses the same argument as (D.5), and the second inequality applies Lemma D.2. Finally, combining (D.9), (D.10), and (D.11), we conclude that

$$\begin{aligned}
\mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r) &\leq -2 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[D_{\text{Hellinger}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \right] \\
&\quad + \frac{2}{N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right) + 6\varepsilon \cdot (1 + e^R).
\end{aligned}$$

By taking the approximation error $\varepsilon = (6 \cdot (1 + e^R) \cdot N)^{-1}$, we conclude that for $\delta < e^{-1}$, with probability at least $1 - \delta$, for any $r \in \mathcal{R}$, it holds that

$$\begin{aligned}
&\mathcal{L}_{\mathcal{D}}(r^*) - \mathcal{L}_{\mathcal{D}}(r) \\
&\leq -2 \cdot \mathbb{E}_{(x, a^1, a^0) \sim \mu_{\mathcal{D}}(\cdot, \cdot, \cdot)} \left[D_{\text{Hellinger}}^2(\mathbb{P}_{r^*}(\cdot|x, a^1, a^0) \|\mathbb{P}_r(\cdot|x, a^1, a^0)) \right] \\
&\quad + \frac{3}{N} \cdot \log \left(\frac{\mathcal{N}_\varepsilon(\mathcal{R}, \|\cdot\|_\infty)}{\delta} \right).
\end{aligned}$$

This completes the proof of Lemma D.1. □

Lemma D.2 (Sigmoid function). *For any real numbers $z_1, z_2 \in [-R, R]$, it holds that*

$$\kappa \cdot |z_1 - z_2| \leq |\sigma(z_1) - \sigma(z_2)| \leq |z_1 - z_2|,$$

where the constant $\kappa = 1/(1 + \exp(R))^2$.

Proof of Lemma D.2. Since the sigmoid function $\sigma(\cdot)$ is differentiable, we know that for any $z_1, z_2 \in [-R, R]$, there exists some $\xi(z_1, z_2) \in [-R, R]$ such that

$$\sigma(z_1) - \sigma(z_2) = \sigma'(\xi(z_1, z_2)) \cdot (z_1 - z_2).$$

Notice that $\sigma'(z) = \sigma(z) \cdot (1 - \sigma(z))$, we can obtain that

$$\begin{aligned} 1 &\geq \sigma'(\xi(z_1, z_2)) = \sigma(\xi(z_1, z_2)) \cdot (1 - \sigma(\xi(z_1, z_2))) \\ &= \frac{1}{1 + \exp(\xi(z_1, z_2))} \cdot \left(1 - \frac{1}{1 + \exp(\xi(z_1, z_2))}\right) \\ &\geq \frac{1}{1 + \exp(R)} \cdot \left(1 - \frac{1}{1 + \exp(-R)}\right) \\ &= \frac{1}{(1 + \exp(R))^2}. \end{aligned}$$

This completes the proof of Lemma D.2. $\square$

E Proofs for Equivalence between Maximin and Minimax Objectives

E.1 Proof of Theorem 5.6

Proof of Theorem 5.6. Consider denoting an auxiliary policy $\hat{\pi}$ as

$$\hat{\pi} \in \operatorname{argmax}_{\pi \in \Pi} \min_{r \in \mathcal{R}} \phi(\pi, r). \quad (\text{E.1})$$

By the definition of $\hat{r}$ and $\hat{\pi}$, the duality gap of $(\hat{r}, \hat{\pi})$, defined as

$$\text{Dual}(\hat{r}, \hat{\pi}) := \max_{\pi \in \Pi} \phi(\pi, \hat{r}) - \min_{r \in \mathcal{R}} \phi(\hat{\pi}, r)$$

is zero. This is because the following deduction,

$$\begin{aligned} \text{Dual}(\hat{r}, \hat{\pi}) &= \left(\max_{\pi \in \Pi} \phi(\pi, \hat{r}) - \min_{r \in \mathcal{R}} \max_{\pi \in \Pi} \phi(\pi, r) \right) \\ &\quad + \left(\max_{\pi \in \Pi} \min_{r \in \mathcal{R}} \phi(\pi, r) - \min_{r \in \mathcal{R}} \phi(\hat{\pi}, r) \right) \\ &= 0, \end{aligned} \quad (\text{E.2})$$

where in the first equality we apply Lemma E.1 that the minimax objective and the maximin objective are equivalent, and the last equality applies the definition of $\hat{r}$ and $\hat{\pi}$ respectively. Note that we can rewrite the duality gap as following

$$\text{Dual}(\hat{r}, \hat{\pi}) = \left(\max_{\pi \in \Pi} \phi(\pi, \hat{r}) + \phi(\hat{\pi}, \hat{r}) \right) - \left(\phi(\hat{\pi}, \hat{r}) - \min_{r \in \mathcal{R}} \phi(\hat{\pi}, r) \right). \quad (\text{E.3})$$

Combining (E.2) and (E.3), we can conclude that

$$\max_{\pi \in \Pi} \phi(\pi, \hat{r}) = \phi(\hat{\pi}, \hat{r}) \Rightarrow \hat{\pi} \in \operatorname{argmax}_{\pi \in \Pi} \phi(\hat{r}, \pi). \quad (\text{E.4})$$

Now comparing what $\pi_{\hat{r}}$ and $\hat{\pi}$ satisfy in (5.4) and (E.4) respectively, invoking Lemma E.3 that the maximizer of $\phi(\cdot, r)$ given any $r \in \mathcal{R}$ is unique on the support of d_0 , we can conclude that

$$\pi_{\hat{r}}(\cdot|x) = \hat{\pi}(\cdot|x), \quad \forall x \in \text{Supp}(d_0). \quad (\text{E.5})$$

Therefore, by (E.1) and (E.5), and the fact that $\phi(\pi, r)$ depends on π only through its value on the support of d_0 , we can conclude that

$$\pi_{\hat{r}} \in \operatorname{argmax}_{\pi \in \Pi} \min_{r \in \mathcal{R}} \phi(\pi, r).$$

This finishes the proof of Theorem 5.6. $\square$

E.2 Auxiliary Lemmas

Lemma E.1 (Equivalence of maximin and minimax objectives). *For the policy class Π defined in (2.3) and the reward model class $\mathcal{R}$ satisfying Assumption 5.5, it holds that the maximin objective is equivalent to the minimax objective, i.e.,*

$$\max_{\pi \in \Pi} \min_{r \in \mathcal{R}} \phi(\pi, r) = \min_{r \in \mathcal{R}} \max_{\pi \in \Pi} \phi(\pi, r).$$

Proof of Lemma E.1. The foundation of this result is a minimax theorem given by [23] (Lemma E.2). In our setting, the policy class Π is a nonempty set, and the reward model class $\mathcal{R}$ is a nonempty compact Hausdorff space. Furthermore, by our choice of the policy class Π in (2.3), Π is a convex set. Meanwhile, the function ϕ is a concave function of $\pi \in \Pi$ since the dependence on π is linear terms plus a negative KL term (concave). Finally, by our assumption, the function ϕ is convex-like on the reward model class $\mathcal{R}$ and is also continuous on $\mathcal{R}$. As a result, all the conditions of Lemma E.2 are satisfied and the minimax theorem holds in our problem setup, finishing the proof of Lemma E.1. $\square$

Lemma E.2 (Minimax theorem [23]). *Let $\mathcal{X}$ be a nonempty set (not necessarily topologized) and $\mathcal{Y}$ be a nonempty compact topological space. Let $f : \mathcal{X} \times \mathcal{Y} \mapsto \mathbb{R}$ be lower semicontinuous on $\mathcal{Y}$. Suppose that f is concave-like on $\mathcal{X}$ and convex-like on $\mathcal{Y}$, i.e., for any $x_1, x_2 \in \mathcal{X}$, $\alpha \in [0, 1]$, there exists $x_3 \in \mathcal{X}$ such that*

$$f(x_3, \cdot) \geq \alpha \cdot f(x_1, \cdot) + (1 - \alpha) \cdot f(x_2, \cdot) \text{ on } \mathcal{Y},$$

and for any $y_1, y_2 \in \mathcal{Y}$, $\beta \in [0, 1]$, there exists $y_3 \in \mathcal{Y}$ such that

$$f(\cdot, y_3) \leq \beta \cdot f(\cdot, y_1) + (1 - \beta) \cdot f(\cdot, y_2) \text{ on } \mathcal{X}.$$

Then the following equation holds,

$$\max_{x \in \mathcal{X}} \min_{y \in \mathcal{Y}} f(x, y) = \min_{y \in \mathcal{Y}} \max_{x \in \mathcal{X}} f(x, y).$$

Lemma E.3 (Unique maximizer of ϕ). *Consider the function ϕ defined as*

$$\begin{aligned} \phi(\pi, r) := & \eta \cdot \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x), a^0 \sim \pi^{\text{base}}(\cdot|x)} \left[r(x, a^1) - r(x, a^0) - \beta \cdot D_{\text{KL}}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] \\ & + \mathcal{L}_{\mathcal{D}}(r). \end{aligned}$$

Then given any $r \in \mathcal{R}$, the maximizer of $\phi(\cdot, r)$ is unique on the support of d_0 .

Proof of Lemma E.3. Given any $r \in \mathcal{R}$, consider that

$$\begin{aligned} & \max_{\pi \in \Pi} \phi(\pi, r) \\ &= \eta \cdot \max_{\pi \in \Pi} \left\{ \mathbb{E}_{x \sim d_0, a^1 \sim \pi(\cdot|x)} \left[r(x, a^1) - \beta \cdot D_{\text{KL}}(\pi(\cdot|x) \parallel \pi^{\text{ref}}(\cdot|x)) \right] \right\} \\ &= \eta \cdot \max_{\pi \in \Pi} \left\{ C_r - \beta \cdot \mathbb{E}_{x \sim d_0} \left[D_{\text{KL}} \left(\pi(\cdot|x) \parallel \frac{\pi^{\text{ref}}(\cdot|x) \cdot \exp(\beta^{-1} \cdot r(x, \cdot))}{\int_{a' \in \mathcal{A}} d\pi^{\text{ref}}(a'|x) \cdot \exp(\beta^{-1} \cdot r(x, a'))} \right) \right] \right\}, \end{aligned}$$

where

$$C_r = \mathbb{E}_{x \sim d_0} \left[\beta \cdot \log \left(\int_{a \in \mathcal{A}} d\pi^{\text{ref}}(a|x) \cdot \exp(\beta^{-1} \cdot r(x, a)) \right) \right]$$

is a constant independent of π . Therefore, the maximizer of $\phi(\cdot, r)$ on the support of d_0 must equal to

$$\pi_r(\cdot|x) = \frac{\pi^{\text{ref}}(\cdot|x) \cdot \exp(\beta^{-1} \cdot r(x, \cdot))}{\int_{a' \in \mathcal{A}} d\pi^{\text{ref}}(a'|x) \cdot \exp(\beta^{-1} \cdot r(x, a'))},$$

which completes the proof of Lemma E.3. $\square$

F Additional Details on Experiments

F.1 Training Details

We train the gemma series models with 8 NVIDIA A6000 GPUs and the beta series models with 8 NVIDIA A100 GPUs, where they are all GPT-like models with around 7 billion parameters. It takes around three hours to train a beta series model and five hours to train a gemma one. Our codebase is adapted from the Alignment Handbook [63]. By comparing the validation loss on the test split (not used for later evaluation), we select the hyperparameter η of both RPO (beta) and RPO (gemma) to be 0.005. We list the remaining training configurations in Table 3, which are recommended by the Alignment Handbook.

Configuration	Beta Series	Gemma Series
learning rate	5.0e-7	5.0e-7
learning scheduler type	cosine	cosine
warmup ratio	1.0	1.0
batch size	128	128
gradient accumulation	2	16
batch size per device	8	1
training epoch	1	2
β	0.01	0.05
optimizer	adamw torch	adamw torch
seed	42	42
precision	bfloat16	bfloat16

Table 3: Training configurations for beta series and gemma series models in this paper.

F.2 Evaluation Details

GPT-4 evaluation on the test split. We use the following prompts to guide GPT-4 to annotate the preferences among win, lose, and tie (we denote them by A, B, and C, respectively).

Prompts: Please act as an impartial judge and evaluate the quality of the responses provided by two AI assistants to the user question displayed below. You should choose the assistant that follows the user’s instructions and answers the user’s question better. Your evaluation should consider factors such as the helpfulness, relevance, accuracy, depth, creativity, and level of detail of their responses. Begin your evaluation by comparing the two responses and provide a short explanation. Avoid any position biases and ensure that the order in which the responses were presented does not influence your decision. Do not allow the length of the responses to influence your evaluation. Do not favor certain names of the assistants. Be as objective as possible. After providing your explanation, output your final verdict by strictly following this format: `[[A]]` if assistant A is better, `[[B]]` if assistant B is better, and `[[C]]` for a tie. [Instruction] instruction [The Start of Assistant A’s Answer] {*answer A*} [The End of Assistant A’s Answer] [The Start of Assistant B’s Answer] {*answer B*} [The End of Assistant B’s Answer]

Here, we replace {*answer A*} and {*answer B*} with the answers of two models. Since GPT annotation has shown to prefer the answer in the first position [66], we randomly exchange the positions between two answers during the evaluation to ensure a fair comparison.

Benchmark evaluation. We use the default configuration for the evaluations on MT-Bench² and AlpacaEval 2.0³. By default, the annotator of MT-Bench is the *latest version* of GPT-4. The default annotator and the competitor model are both GPT-4 (Preview 11/06). We only need to manually import the proper chat template that formats the training dataset, which are shown as follows.

²https://github.com/lm-sys/FastChat/tree/main/fastchat/llm_judge

³https://github.com/tatsu-lab/alpaca_eval/tree/main

Chat Template for Beta Series: <|system|></s><|user|>
{instruction}</s>
<|assistant|>

Chat Template for Gemma Series: <bos> <lim_start|>user
{instruction}<lim_end|>
<lim_start|>assistant

F.3 Additional Results on Experiments

In this section, we provide the additional results to show the performance gain for RPO (beta) in MT-Bench and RPO (gemma) in AlpacaEval 2.0. We report the pairwise win rates in Tables 4, 5, and 6 to analyze their performance gaps, where all the annotation configurations are the same in Table 2. Results show that RPO still exceeds DPO in the metric of the pairwise win rates on the benchmarks for both beta series and gemma series.

win rate (%)	RPO (beta)	Ref. (beta)	DPO (beta)
RPO (beta)	50.00	83.75	57.81
Ref. (beta)	16.25	50.00	21.25
DPO (beta)	78.75	42.19	50.00

Table 4: Pairwise win rates (left vs. right) for beta series models on MT-Benchmark.

win rate (%)	RPO (beta)	Ref. (beta)	DPO (beta)
RPO(beta)	50.00	80.13	52.02
Ref.(beta)	19.87	50.00	20.61
DPO (beta)	47.98	79.39	50.00

Table 5: Pairwise win rates (left vs. right) for gemma series models on AlpacaEval 2.0.

win rate (%)	RPO (beta)	Ref. (beta)	DPO (beta)
RPO (beta)	50.00	64.93	51.33
Ref. (beta)	35.07	50.00	36.44
DPO (beta)	48.67	64.56	50.00

Table 6: Pairwise Length-Control (LC) win rates (left vs. right) for gemma series models on AlpacaEval 2.0.

G Experiments on Math, Reasoning, and Coding Tasks

G.1 Experimental Details

To provide a more comprehensive analysis of the trained LLM, we introduce more benchmarks on the math, reasoning, and coding tasks for evaluations. Specifically, we choose the Grade School Math 8K (GSM8K), AI2 Reasoning Challenge (ARC), and Mostly Basic Python Programming (MBPP) to measure math, reasoning, and coding abilities, respectively. In this section, we focus on the gemma series for the experiments. We do not use chain-of-thought or few shots in all the benchmarks. We compare the greedy decoding result (pass @1) on the MBPP benchmark.

Model Name	GSM8K (%)	ARC		MBPP (Pass @1)	
		Easy (%)	Challenge (%)	Normal (%)	Plus (%)
RPO	49.9	79.1	49.8	54.2	46.3
DPO	45.3	75.7	50.0	54.2	43.9
Ref.	45.4	75.0	45.8	50.3	44.2
zephyr-gemma-7b	47.3	77.6	48.6	54.5	44.7

Table 8: Results on GSM8K, ARC, and MBPP. Here, zephyr-gemma-7b is the officially released models trained by DPO and Ref. denotes the reference model zephyr-7b-gemma-sft used for our training. RPO and DPO are trained with the OpenRLHF codebase [27] and we average the SFT loss regularizer in RPO by the number of tokens of the chosen response. We do not use chain-of-thought or few shots in all the benchmarks. We compare the greedy decoding result (pass @1) for MBPP.

Here we use the OpenRLHF codebase [27] to implement a new variant of RPO, where the SFT loss regularizer is averaged by the number of tokens of the chosen labels, that is, $(\log \pi_{\theta}(a_{\text{cho}}|x))/|a_{\text{cho}}|$. Such a variant balances the weight of the averaged SFT loss regularizer between the shorter chosen response and the longer one. We set the coefficient for the SFT loss regularizer as 0.2. We use 8 NVIDIA A100 GPUs for the training and evaluation. The remaining hyperparameters are in Table 7.

Configuration	Gemma Series
learning rate	5.0e-7
learning scheduler type	cosine with a minimum learning rate
batch size	128
gradient accumulation	8
batch size per device	2
training epoch	2
β	0.5
optimizer	adamw torch
seed	42
precision	bfloat16

Table 7: Training configurations for DPO and RPO for the experiments in Appendix G.

G.2 Experimental Results

Table 8 demonstrates that our proposed method still outperforms or performs equally to the vanilla DPO on these benchmarks of math, reasoning, and coding, which verifies the effectiveness of our proposed method.

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes], [No], or [NA].
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]", it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

IMPORTANT, please:

- **Delete this instruction block, but keep the section heading "NeurIPS paper checklist",**
- **Keep the checklist subsection headings, questions/answers and guidelines below.**
- **Do not modify the questions and only use the provided macros for your answers.**

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: We support all the claims made in the abstract and the introduction sections by our theory and experiment sections.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: Please see the discussion of limitation of the work in Appendix [B](#).

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

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

Justification: We provide the complete and accurate proof in Appendix [D](#).

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

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

Justification: We provide all the detailed training and evaluation configurations in Appendix [F.1](#) and [F.2](#) for reproducibility. We also submit the codes in the supplementary for the purpose of reproducibility.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We submit the codes in the supplementary. All the datasets, reference models, and benchmarks used in this paper are open accessed.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.

- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

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

Justification: Please see the dataset choice and the detailed training configurations in Section 6 and Appendix F.1.

Guidelines:

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

7. Experiment Statistical Significance

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

Answer: [\[No\]](#)

Justification: Due to the budget of time and expense, we do not report the error bar, which is also common in many RLHF literature [46, 34, 71, 64]. However, we report all the training configurations and the random seed to ensure reproducibility.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

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

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

Justification: We provide the information in Appendix F.1.

Guidelines:

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

9. Code Of Ethics

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

Answer: [Yes]

Justification: We have reviewed and followed the code of ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: There is no societal impact of the work performed.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

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

Answer: [NA]

Justification: The scope of our work is not to publish a new released model but to analyze the overoptimization pheromone in RLHF both theoretically and empirically.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We respect all the licenses and terms of codes, models, and datasets used in this paper. We also properly cite their creators in this paper.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

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

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

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

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.