

PILAF: OPTIMAL HUMAN PREFERENCE SAMPLING FOR REWARD MODELING

Yunzhen Feng¹ Ariel Kwiatkowski^{2*} Kunhao Zheng^{2*} Julia Kempe^{1◇} Yaqi Duan^{1◇}

¹ NYU, ² Meta FAIR, * Joint second authors, ◇ Joint senior authors.

ABSTRACT

As large language models increasingly drive real-world applications, aligning them with human values becomes paramount. Reinforcement Learning from Human Feedback (RLHF) has emerged as a key technique, translating preference data into reward models when oracle human values remain inaccessible. In practice, RLHF mostly relies on approximate reward models, which may not consistently guide the policy toward maximizing the underlying human values. We propose Policy-Interpolated Learning for Aligned Feedback (PILAF), a novel response sampling strategy for preference labeling that explicitly aligns preference learning with maximizing the underlying oracle reward. PILAF is theoretically grounded, demonstrating optimality from both an optimization and a statistical perspective. The method is straightforward to implement and demonstrates strong performance in iterative and online RLHF settings where feedback curation is critical.

1 INTRODUCTION

Reinforcement Learning from Human Feedback (RLHF) (Ouyang et al., 2022) has revolutionized large language models (LLMs) by incorporating human preferences, enabling significant progress in applications such as conversational AI (Achiam et al., 2023), personalized tutoring (Limo et al., 2023), and content curation (Yue et al., 2024). At the core of RLHF is *reward modeling*, a critical process that translates human feedback—such as pairwise comparisons or rankings—into a measurable objective for model training. By formalizing human preferences, reward models then guide LLMs towards alignment through *policy optimization*.

While numerous studies have focused on improving language models (LMs) by optimizing fixed reward functions (Dong et al., 2023; Liu et al., 2024c) or leveraging pre-existing preference datasets (Ethayarajh et al., 2024; Azar et al., 2024; Xu et al., 2024), comparatively less attention has been paid to the critical challenge of collecting *effective* data for human-labeling in RLHF, to maximize its utility. This is an important problem, as the quality of preference data directly impacts the effectiveness of reward modeling and, consequently, the overall success of fine-tuning. This challenge is further compounded by the high cost of expert preference labeling (Lightman et al., 2023).

Preference data is usually generated by sampling response pairs $(\tilde{y}_i^a, \tilde{y}_i^b)$ to a prompt x_i from a policy, and presenting them to human labelers for preference annotation. It is commonly assumed that the annotation follows the Bradley-Terry (BT) model, under an *oracle reward*. Next, we use maximum likelihood estimation (MLE) on these preference data to train a reward model, which then serves as a measurable objective to optimize the policy (i.e. LLM) while staying close to a reference policy. In Direct Preference Optimization (DPO) (Rafailov et al., 2023), this pipeline is simplified by optimizing the policy with implicit reward modeling. However, all these pipelines give rise to a *misalignment of objectives*: RLHF (or DPO) should, in principle, train its policy to maximize the (inaccessible) *oracle objective* which combines the *oracle reward* from the BT model with reference regularization. In practice, RLHF relies on preference data through the MLE objective in reward modeling or through methods like DPO, which are *not* designed to guide policy optimization towards maximizing oracle rewards. Thus, reward optimization (either directly or implicitly via DPO) and (optimal) policy optimization are not inherently aligned, potentially leading to inefficiencies (Sec. 2).

In this work, we study this misalignment by examining the sampling scheme that generates response pairs $(\tilde{y}_i^a, \tilde{y}_i^b)$ for preference labeling, which is especially important when additional preference data is collected mid-RLHF training to mitigate the off-policy distributional shift, as is empirically

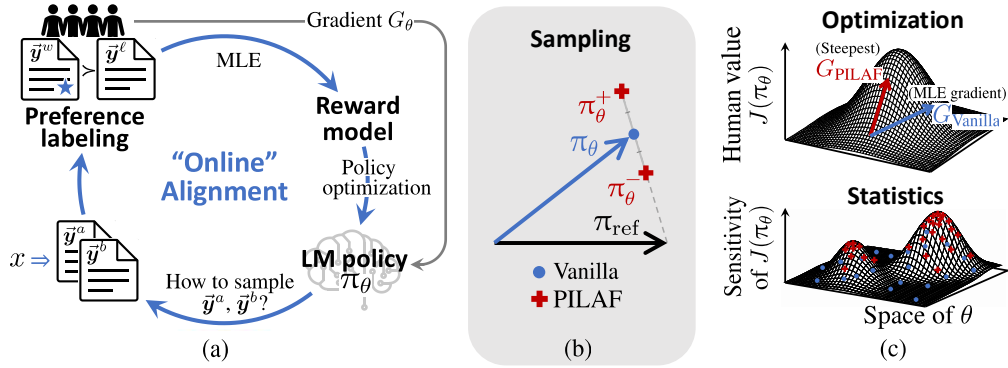


Figure 1: **Overview of our approach.** (a) We consider a full RLHF training setup, where a language model (LM) policy is iteratively refined through active data collection. Our goal is to develop an optimal response sampling method for preference labeling. (b) We introduce PILAF, which generates responses by interpolating between the current policy and a reference policy, balancing exploration and exploitation. (c) Our theoretical analysis shows that T-PILAF aligns the parameter gradient with the steepest direction for maximizing human values and achieves more favorable convergence in regions of high sensitivity.

standard (Touvron et al., 2023; Bai et al., 2022). We show that uniform sampling from the current policy, as is common, leads to misaligned gradients of the two objectives (reward model loss and true oracle objective).

To tackle this issue, we present *Theoretically Grounded Policy-Interpolated Learning for Aligned Feedback* (T-PILAF), a novel sampling method that aligns reward modeling with value optimization. Specifically, T-PILAF generates responses by interpolating the policy model and the reference model for a balanced exploration and exploitation. We provide rigorous theoretical analysis to show that for preference data generated with T-PILAF, the gradient of the MLE loss with respect to the policy network’s parameters is aligned with the policy gradient of the oracle objective in a first-order sense. This alignment enables the policy to optimize directly for the oracle value, achieving both alignment and efficiency. Furthermore, we separately show from a statistical perspective that T-PILAF aligns optimization with the steepest directions of the oracle objective. It thus makes the sampled preference pairs more informative, reducing variance and improving training stability.

We then present PILAF, a simple modification of our theoretical sampling scheme T-PILAF, which naturally lends itself to practical implementation. For clarity of exposition, we present our method in the context of DPO; however, PILAF can be adapted to a wide class of preference optimization methods.¹ See Figure 1 for an illustration of our setup, method, and the optimization and statistical principles underlying PILAF.

We conduct extensive experiments to validate PILAF’s effectiveness and robustness. As a stand-in for expensive human annotators, we use a well-trained reward model—Skywork-Llama-3.1-8B (Liu et al., 2024a)—as a proxy for the oracle reward. Throughout training, we query this model exclusively for preference labels, simulating human feedback. We then align the Llama-3.1-8B base model (Dubey et al., 2024) using these proxy-labeled preference data in two settings: iterative DPO (Xiong et al., 2024) and online DPO (Guo et al., 2024). In both scenarios, preference data is collected on-the-fly, either after each full training epoch in the iterative setting or after every training step in the online setting. Across all configurations, PILAF outperforms all the baselines, producing a policy with higher reward (as measured by the proxy) and a lower KL divergence from the reference model, reducing annotation and computation costs by over 40% in iterative DPO.

Our key contributions are as follows:

- (*Practical sampling algorithm*) We propose PILAF (Section 5), an efficient sampling algorithm for generating response pairs in the RLHF pipeline for improved sample efficiency and performance, derived from its theoretically grounded variant T-PILAF (Section 3).
- (*Theoretical optimality*) We provide theoretical guarantees for the efficiency of our approach from both optimization and statistical perspectives (Section 4).
- (*Empirical validation*) We validate PILAF in both iterative and online DPO settings (Section 6) and observe that it consistently outperforms baselines by achieving higher reward and lower KL divergence from the reference model. Moreover, PILAF achieves comparable performance at significantly reduced annotation and computational costs.

¹See Appendix G for the extension to PPO.

1.1 RELATED WORK

Existing Sampling Schemes. In academic papers, uniform vanilla sampling is the most commonly used approach, while methods such as best-of-N and worst-of-N have also been explored (Dong et al., 2024). Xie et al. (2024) propose sampling one response from the current policy model and another from a reference model, modifying the loss function to encourage optimistic behavior. Similarly, Zhang et al. (2024) sample one response from the current model but rank it alongside two offline responses from the reference model. Shi et al. (2024) present a formula similar to ours based on intuition, introducing several hyperparameters and analyzing convergence speed with DPO in a tabular setting. Liu et al. (2024d) train an ensemble of reward models to approximate a posterior distribution over possible rewards and use Thompson sampling to generate responses with exploration. In contrast to these works, we theoretically establish the principles of response generation for preference labeling, making minimal assumptions and simplifications while demonstrating the optimality of our approach. Our approach eliminates the need for hyperparameter tuning.

Policy Gradient. Our theoretical principle is closely related to the family of policy gradient methods (Williams, 1992; Sutton et al., 1999) in reinforcement learning, which optimize a policy π_θ by estimating and ascending the gradient of the expected return $\nabla_\theta J(\theta)$. Significant advancements have been made to improve the efficiency of these methods, including variance reduction techniques (Greensmith et al., 2004), off-policy gradient estimation (Degris et al., 2012), interpolating on-policy and off-policy updates (Gu et al., 2017), deterministic policy gradients (Silver et al., 2014), and three-way robust estimation approaches (Kallus & Uehara, 2020). Our study extends these principles to preference learning for LMs, aligning the MLE gradient with the oracle objective gradient by controlling the response sampling distribution, thereby improving learning efficiency.

A review of other RLHF literature, particularly on data selection for the preference dataset, is deferred to Appendix A.

2 PROBLEM SETUP AND MOTIVATION

Language Model (LM). At the core of RLHF is a language model that processes prompts $x \in \mathcal{X}$ and generates responses $\vec{y} \in \mathcal{Y}$. Each response is represented as a sequence of tokens $\vec{y} = (y_1, y_2, \dots, y_T)$. The primary goal of RLHF is to guide the model to generate responses that align with human preferences. This translates to designing a policy π (parameterized as a LM) that maps prompts to responses, maximizing a reward that reflects human preferences (with a KL regularization).

Preference Data. The oracle reward for human values is inherently inaccessible. Instead, the alignment process approximates the reward using a dataset of human-labeled preferences, $\mathcal{D} = \{(x_i, \vec{y}_i^w, \vec{y}_i^\ell)\}_{i=1}^n$, where each sample contains: (i) a prompt x_i , independently drawn from a distribution ρ , and (ii) a pair of responses $(\vec{y}_i^w, \vec{y}_i^\ell)$, where $\vec{y}_i^w$ is preferred over $\vec{y}_i^\ell$ in human labeling. The response pair $(\vec{y}_i^w, \vec{y}_i^\ell)$ is first generated from a joint distribution $\mu(\cdot | x)$ and then presented to human labelers for preference annotation. Human preferences are commonly modeled using the *Bradley-Terry (BT)* model, which assumes:

$$\mathbb{P}(\vec{y}^a \succ \vec{y}^b | x) = \sigma(r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)), \quad (1)$$

where $r^*(x, \vec{y})$ represents the (unknown) oracle reward of a response given a prompt, and $\sigma(z) = \{1 + \exp(-z)\}^{-1}$ is the sigmoid function, mapping differences in rewards to probabilities. We adopt the BT model throughout this paper.

Reward Modeling. The preference data, encoding human judgment, is then used to train a reward model, r_θ , which serves as a measurable objective for training the policy model. r_θ is trained by solving a MLE objective:

$$\min_{\theta} \hat{\mathcal{L}}(\theta) := -\frac{1}{n} \sum_{i=1}^n \log \sigma(r_\theta(x_i, \vec{y}_i^w) - r_\theta(x_i, \vec{y}_i^\ell)). \quad (2)$$

This empirical loss approximates the expected negative log-likelihood

$$\mathcal{L}(\theta) := \mathbb{E}_{x \sim \rho, (\vec{y}^a, \vec{y}^b) \sim \mu(\cdot | x)} \left[-\log \sigma(r_\theta(x, \vec{y}^w) - r_\theta(x, \vec{y}^\ell)) \right]. \quad (3)$$

Policy Optimization. To align a language model ϕ with human preferences, we optimize it to maximize the learned rewards r_θ while staying close to a reference policy π_{ref} . The objective is

$$\max_{\phi} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\phi}(\cdot | x)} [r_\theta(x, \vec{y})] - \beta D_{\text{KL}}(\pi_{\phi} \parallel \pi_{\text{ref}}). \quad (4)$$

It consists of two parts:

- (i) The *reward* term $\mathbb{E}_{x \sim \rho, \vec{y} \sim \pi(\cdot|x)}[r_\theta(x, \vec{y})]$ encourages the policy to generate high-quality responses.
- (ii) The *regularization* term $D_{\text{KL}}(\pi \parallel \pi_{\text{ref}})$ penalizes deviations from the reference policy π_{ref} and is defined as $\mathbb{E}_{x \sim \rho}[D_{\text{KL}}(\pi(\cdot|x) \parallel \pi_{\text{ref}}(\cdot|x))]$.

Here, β is a regularization parameter that balances the trade-off between reward maximization and adherence to the reference policy. We assume β is fixed and practitioner-specified.

Direct Preference Optimization. The above-described RLHF pipeline typically leverages the Proximal Policy Optimization (PPO) algorithm (Schulman et al., 2017) to perform policy optimization. This approach requires loading the policy network, reward model, reference model, and a value model onto the GPU during training, making it highly resource-intensive. To improve computational efficiency and practicality, Direct Preference Optimization (DPO) (Rafailov et al., 2023) has been proposed, enabling direct alignment without the need for a reward model or a value model.

A key insight of DPO is that any policy π_θ can be viewed as the optimal solution to problem equation 4 if the reward r_θ is

$$r_\theta(x, \vec{y}) := \beta \cdot \log \left(\frac{\pi_\theta(\vec{y} | x)}{\pi_{\text{ref}}(\vec{y} | x)} \right). \quad (5)$$

Thus, DPO can directly optimize the policy π_θ using $\hat{\mathcal{L}}(\theta)$ in Equation (2), where r_θ is replaced by π_θ as defined in Equation (5). This reformulation makes the objective dependent solely on θ , with the reward being implicitly learned through the policy itself. As a result, the optimization process becomes significantly more efficient.

Motivation. To fully align with human values, RLHF should, in principle, train the policy to maximize the oracle reward, r^* , as defined in the BT model. The corresponding oracle objective is then:

$$J(\pi) := \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi(\cdot|x)}[r^*(x, \vec{y})] - \beta D_{\text{KL}}(\pi \parallel \pi_{\text{ref}}).$$

Since direct access to r^* is unavailable, RLHF instead relies on preference data, either through MLE-based reward modeling or methods like DPO. However, these processes are not inherently designed to train the policy to directly maximize the oracle objective, $J(\pi)$.

In this work, we aim to design an optimal sampling distribution μ to realign DPO with the maximization of $J(\pi)$. Such a sampling strategy will improve the quality of the preference dataset, maximize the utility of limited data, and enhance both performance and efficiency. This focus is particularly crucial in scenarios where additional data is collected during mid-training—a key phase in the iterative fine-tuning of LMs (Touvron et al., 2023; Bai et al., 2022; Xiong et al., 2024; Guo et al., 2024). At this stage, a preliminary policy π_θ (distinct from π_{ref}) is already in place, but its performance may fall short of expectations. It is thus necessary to gather additional preference data, ideally on-policy data that target areas where the current policy shows room for improvement. An effective sampling design can significantly enhance the efficiency of leveraging human feedback in this process.

3 T-PILAF: THEORETICAL SAMPLING SCHEME

We now present T-PILAF - *theoretically grounded policy interpolation for aligned feedback* - our sampling scheme for generating responses in data collection². The scheme is shown (in Section 4) to be optimal from both optimization and statistical perspectives.

Consider we have an initial policy π_θ and aim to collect preference data to further refine its performance. We propose two complementary variants of policy π_θ : one that encourages exploration in regions more preferred by π_θ , reflecting an optimistic perspective, and another that focuses on areas less favored by π_θ , reflecting a conservative adjustment.

²The T in T-PILAF serves to distinguish the theoretical scheme from the derived, simplified, efficiently implementable PILAF.

Specifically, we define policies π_θ^+ and π_θ^- around π_θ as

$$\pi_\theta^+(\vec{y} | x) := \frac{1}{Z_\theta^+(x)} \pi_\theta(\vec{y} | x) \exp \{r_\theta(x, \vec{y})\}, \quad (6a)$$

$$\pi_\theta^-(\vec{y} | x) := \frac{1}{Z_\theta^-(x)} \pi_\theta(\vec{y} | x) \exp \{-r_\theta(x, \vec{y})\}, \quad (6b)$$

where the reward function r_θ is defined in equation 5. The partition function $Z_\theta^+(x)$ (or $Z_\theta^-(x)$) is given by $Z_\theta^+(x) := \int_{\mathcal{Y}} \pi_\theta(\vec{y} | x) \exp\{r_\theta(x, \vec{y})\} d\vec{y}$.

For any prompt $x \in \mathcal{X}$, our sampling procedure involves the following steps:

- (i) Draw a random variable ξ from Bernoulli($p_0(x)$), where $p_0(x) := Z_\theta^+(x) Z_\theta^-(x) / \{1 + Z_\theta^+(x) Z_\theta^-(x)\}$.
- (ii) If $\xi = 1$, independently draw responses $\vec{y}^a, \vec{y}^b \in \mathcal{Y}$ according to $\vec{y}^a \sim \pi_\theta^+(\cdot | x)$ and $\vec{y}^b \sim \pi_\theta^-(\cdot | x)$. If $\xi = 0$, draw responses as $\vec{y}^a, \vec{y}^b \sim \pi_\theta(\cdot | x)$.

In the next section, we will theoretically analyze T-PILAF. To account for the changes in sampling, we adopt a slightly modified loss function in the theoretical framework:

$$\hat{\mathcal{L}}(\theta) := -\frac{1}{n} \sum_{i=1}^n w(x_i) \cdot \log \sigma \left(r_\theta(x_i, \vec{y}_i^w) - r_\theta(x_i, \vec{y}_i^\ell) \right).$$

The newly introduced weight function w is defined as

$$w(x) := \{1 + Z_\theta^+(x) Z_\theta^-(x)\} / \bar{Z}_\theta, \quad (7)$$

where the normalization constant $\bar{Z}_\theta > 0$ is given by $\bar{Z}_\theta := 1 + \int_{\mathcal{X}} Z_\theta^+(x) Z_\theta^-(x) \rho(x) dx$. We also modify the population loss $\mathcal{L}$ in Equation (3) with the weight function.

4 THEORETICAL ANALYSIS

This section provides the theoretical grounding and analysis of our proposed sampling scheme from two perspectives. In the *optimization* analysis (Section 4.1) we show that T-PILAF *aligns two objectives (gradient alignment property)*: maximizing the likelihood function (Equation (3)) becomes equivalent to gradient ascent on the value function $J(\pi_\theta)$ (Equation (6)). Consequently, policy updates on π_θ move the parameters in the direction of steepest increase of J . T-PILAF thus provides the potential to accelerate training and improve generalization, compared to vanilla (uniform) sampling. In the *statistical* analysis (Section 4.2) we focus on statistical error and show that the asymptotic covariance of the estimated parameter $\hat{\theta}$ (inversely) aligns with the Hessian of the objective function J when sampling with T-PILAF. As a result, T-PILAF makes the sampled comparisons more informative, as they align with directions where J is most sensitive. The net outcome is reduced statistical variance of our method through tighter concentration of estimates in directions that matter most for performance.

4.1 OPTIMIZATION CONSIDERATIONS

We begin by analyzing the DPO algorithm from an optimization perspective. Theorem 4.1 below formally illustrates how T-PILAF ensures alignment between the MLE gradient, $\nabla_\theta \mathcal{L}(\theta)$, and the oracle objective gradient, $\nabla_\theta J(\pi_\theta)$.

Theorem 4.1 (Gradient structure in DPO training). *Using data collected from our proposed response sampling scheme T-PILAF, the gradient of $\mathcal{L}(\theta)$ satisfies*

$$\nabla_\theta \mathcal{L}(\theta) = -\frac{\beta}{\bar{Z}_\theta} \nabla_\theta J(\pi_\theta) + T_2,$$

where the constant $\bar{Z}_\theta$ is defined in equation 7, and the term T_2 is a second-order error.

The detailed proof of Theorem 4.1 is deferred to Appendix C.1. It involves calculation of explicit forms of the gradients $\nabla_\theta \mathcal{L}(\theta)$ and $\nabla_\theta J(\pi_\theta)$; the most notable technical contribution is showing how to leverage our sampling scheme to approximate the derivative σ' of the sigmoid function. By using T-PILAF sampling, we can transform a difference term of the form $\sigma(\Delta r^*) - \sigma(\Delta r_\theta)$ in $\nabla_\theta \mathcal{L}(\theta)$ into a linear difference $\Delta r^* - \Delta r_\theta$ in $\nabla_\theta J(\pi_\theta)$.

Theorem 4.1 establishes the *gradient alignment* property, demonstrating that minimizing the likelihood-based loss function $\mathcal{L}$ closely aligns with maximizing the oracle objective function J , with only a minor second-order error. It highlights how the proposed sampling scheme enables the DPO framework to effectively guide the policy toward optimizing the expected reward. Beyond DPO, in Appendix G, we show how the same principle can be applied to PPO-based RLHF algorithms to help improve the sampling.

4.2 STATISTICAL CONSIDERATIONS

From a statistical standpoint, we first examine the asymptotic distribution of the estimated parameter $\hat{\theta}$ when it (approximately) solves the optimization problem equation 2. In Theorem 4.2, we formally characterize the randomness or statistical error inherent in $\hat{\theta}$ under this idealized scenario. The detailed proof of Theorem 4.2 is provided in Appendix C.2.2.

Theorem 4.2. *Assume the reward model r^* in the BT model equation 1 satisfies $r^* = r_{\theta^*}$ for some parameter θ^* . Under mild regularity conditions, the estimate $\hat{\theta}$ asymptotically follows a Gaussian distribution $\sqrt{n}(\hat{\theta} - \theta^*) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \mathbf{\Omega})$ as $n \rightarrow \infty$. We have an estimate of the covariance matrix $\mathbf{\Omega}$: $\mathbf{\Omega} \preceq C_1 \cdot \mathbf{\Sigma}_*^{-1}$, where $C_1 > 0$ is a universal constant. When using T-PILAF, $\mathbf{\Sigma}_*$ is given by*

$$\mathbf{\Sigma}_* := \mathbb{E}_{x \sim \rho} \left[\text{Cov}_{\tilde{\mathbf{y}} \sim \pi^*(\cdot|x)} [\nabla_{\theta} r^*(x, \tilde{\mathbf{y}}) \mid x] \right]. \quad (8)$$

Next we analyze the performance of the output policy $\hat{\pi} = \pi_{\hat{\theta}}$ from Theorem 4.2 in terms of the expected value $J(\pi)$. In Theorem 4.3, we show that our proposed sampling method guarantees that the covariance of the statistical error in $\hat{\theta}$ aligns inversely with the Hessian of J at the optimal policy π^* . This alignment prioritizes convergence efficiency along directions where the Hessian has large eigenvalues, adapting to the geometry of the optimization landscape. It highlights the efficiency of our sampling scheme in reducing statistical error. For the detailed proof, see Appendix C.2.3.

Theorem 4.3. *The value function $J(\pi)$ we define in equation 6 satisfies $\nabla_{\theta} J(\pi^*) = \mathbf{0}$ and*

$$\nabla_{\theta}^2 J(\pi^*) = -\frac{1}{\beta} \mathbf{\Sigma}_* \quad (9)$$

for matrix $\mathbf{\Sigma}_*$ defined in equation 8. As a corollary, suppose $\mathbf{\Sigma}_*$ is nonsingular, then there exists a constant $C_2 > 0$ such that for any $\varepsilon > 0$,

$$\limsup_{n \rightarrow \infty} \mathbb{P} \left\{ J(\hat{\pi}) < J(\pi^*) - C_2 \cdot \frac{d(1+\varepsilon)}{n} \right\} \leq \mathbb{P} \{ \chi_d^2 > (1+\varepsilon)d \}. \quad (10)$$

Our proposed sampling distribution μ ensures that the output policy $\hat{\pi}$ performs predictably and reliably. The value gap $J(\pi^*) - J(\hat{\pi})$ asymptotically follows a chi-square distribution, irrespective of the problem instance details, such as the underlying reward model r^* . This *structure-invariant statistical efficiency* allows the method to achieve asymptotically efficient estimates without requiring explicit knowledge of the model structure.

We further derive a general lemma describing how μ affect the covariance in Appendix B. This result provides broader insights into what constitutes good preference data in RLHF.

5 PILAF ALGORITHM

We now demonstrate that the T-PILAF sampling scheme defined in Equation (6a) and (6b) can be naturally extended into an efficient empirical algorithm (PILAF).

The first challenge in implementing these definitions lies in calculating the normalizing factors $Z_{\theta}^+(x)$ and $Z_{\theta}^-(x)$, which can be computationally expensive for LLMs. To address this, we simplify the process by replacing them with 1.³ Consequently, the sampling process becomes straightforward: with probability 1/2, we sample using π_{θ} , and otherwise, we sample using π_{θ}^+ and π_{θ}^- .

The second challenge lies in sampling a response $\tilde{\mathbf{y}}$ from $\pi_{\theta}(\tilde{\mathbf{y}} \mid x) \exp \{ \pm r_{\theta}(x, \tilde{\mathbf{y}}) \}$ in an autoregressive way for next-token generation. We argue that the policy π_{θ}^+ (and π_{θ}^-) can be approximated in a token-wise manner:

$$\pi_{\theta}^+(\tilde{\mathbf{y}} \mid x) \approx \pi_{\theta}^+(y_1 \mid x) \pi_{\theta}^+(y_2 \mid x, y_1) \cdots \pi_{\theta}^+(y_t \mid x, y_{1:t-1}) \cdots \pi_{\theta}^+(y_T \mid x, y_{1:T-1}),$$

³When the regularization coefficient β is sufficiently small, the term $\exp\{r_{\theta}(x, \tilde{\mathbf{y}})\}$ in equation 6a stays close to 1 and has only a minor effect. Consequently, the partition function $Z_{\theta}^+(x)$ is approximately 1. A similar reasoning applies to $Z_{\theta}^-(x)$.

Table 1: A cost summary of PILAF and sampling methods from related works. *Best-of-N* method in Xiong et al. (2024) uses the oracle reward to score all candidate responses, then selects the highest- and lowest-scoring ones—instead of providing a preference label for only two responses. We restrict the oracle to providing only preference labels. Thus, we create a *Best-of-N* variant that uses the DPO internal reward for selection and then applies preference labeling, with an annotation cost of 2. We compare with this variant in the experiment.

METHOD	$\tilde{y}^a$	$\tilde{y}^b$	SAMPLING COST	ANNOTATION COST
<i>Vanilla</i> (RAFAILOV ET AL., 2023)	π_θ	π_θ	2	2
<i>Best-of-N</i> (XIONG ET AL., 2024), N=8	BEST OF π_θ	WORST OF π_θ	8	8*
<i>Best-of-N</i> (WITH DPO REWARD), N=8	BEST OF π_θ	WORST OF π_θ	8	2
<i>Hybrid</i> (XIE ET AL., 2024)	π_θ	π_{ref}	2	2
PILAF (OURS)	$\pi_\theta^+ / \pi_\theta$	$\pi_\theta^- / \pi_\theta$	3	2

where $\pi_\theta^+(y_t | x, y_{1:t-1}) = \frac{1}{Z(x, y_{1:t-1})} \pi_\theta(y_t | x, y_{1:t-1}) \left(\frac{\pi_\theta(y_t | x, y_{1:t-1})}{\pi_{\text{ref}}(y_t | x, y_{1:t-1})} \right)^\beta$, with $Z(x, y_{1:t-1})$ being a partition function. The substitution of r_θ uses the correspondence between the reward model r_θ and the policy π_θ in Equation (5), under the assumption that this correspondence holds for all truncations $y_{1:t-1}$. It gives us a direct per-token prediction rule:

$$\pi_\theta^+(\cdot | x, y_{1:t-1}) = \text{softmax}\left(\{(1 + \beta) \mathbf{h}_\theta - \beta \mathbf{h}_{\text{ref}}\}(x, y_{1:t-1})\right).$$

Here $\mathbf{h}_\theta$ and $\mathbf{h}_{\text{ref}}$ are the logits of the policies π_θ and π_{ref} , respectively. β is the regularization coefficient from the objective function $J(\pi)$ in Equation (6). Responses are then generated using standard decoding techniques, such as greedy decoding or nucleus sampling. Similarly, π_θ^- follows

$$\pi_\theta^-(\cdot | x, y_{1:t-1}) = \text{softmax}\left(\{(1 - \beta) \mathbf{h}_\theta + \beta \mathbf{h}_{\text{ref}}\}(x, y_{1:t-1})\right).$$

For a detailed, step-by-step proof, see Proposition 1 in Liu et al. (2024b).

We formalize our final algorithm in Algorithm 1. *Vanilla* DPO (Rafailov et al., 2023; Guo et al., 2024) employs a basic generation approach, sampling $\tilde{y}_i^a, \tilde{y}_i^b \sim \pi_\theta$ at Step 3. In contrast, instead of only sampling from π_θ , our sampling scheme interpolates and extrapolates the logits $\mathbf{h}_\theta$ and $\mathbf{h}_{\text{ref}}$ with coefficient β , enabling exploration of a wider response space to align learning from human preference with value optimization. The β here is the same parameter that controls the KL regularization in Equation (4), as set by the problem.

Cost analysis. We summarize sampling and annotation costs per preference pair for PILAF and related sampling schemes in Table 1. In *Vanilla* sampling (from π_θ), two generations and two annotations are required for human preference labeling, same to PILAF when the pair is sampled from π_θ , which happens half the time. With 50% probability, PILAF uses π_θ^+ and π_θ^- to generate, requiring two forward passes with π_θ and π_{ref} to generate one sample. Thus, on average, a preference pair sampled with PILAF requires a sampling cost of 3 forward passes (1.5 time the cost of *Vanilla*) with the same annotation cost. To compare, Xiong et al. (2024); Dong et al. (2024) perform *Best-of-N* sampling with $N = 8$, which generates and annotates all 8 responses, selecting the best and worst of them. Xie et al. (2024) use a *Hybrid* method that generates with π_θ and π_{ref} , thus matching the sampling and annotation costs of the *Vanilla* method. We empirically compare PILAF with these methods in the next section.

Algorithm 1 DPO with PILAF (ours).

input Prompt Dataset $\mathcal{D}_\rho$, preference oracle $\mathcal{O}$, $\pi_\theta, \pi_{\text{ref}}$.
1: **for** step $t = 1, \dots, T$ **do**
2: Sample n_t prompts $\{x_i\}_{i=1}^{n_t}$ from $\mathcal{D}_\rho$.
3: With probability 1/2, sample $\tilde{y}_i^a, \tilde{y}_i^b \sim \pi_\theta$; with probability 1/2, sample $\tilde{y}_i^a \sim \pi_\theta^+$ and $\tilde{y}_i^b \sim \pi_\theta^-$.
4: Query $\mathcal{O}$ to label $(x_i, \tilde{y}_i^a, \tilde{y}_i^b)$ into $(x_i, \tilde{y}_i^w, \tilde{y}_i^l)$.
5: Update π_{θ_t} with DPO loss using $\{(x_i, \tilde{y}_i^w, \tilde{y}_i^l)\}_{i=1}^{n_t}$.
6: **end for**

6 EXPERIMENTS

In this section, we empirically evaluate PILAF in both an iterative DPO setting (Section 6.1, following Xiong et al. (2024); Dong et al. (2024)) and an online DPO setting (Section 6.2, following Guo et al. (2024)) where the model undergoes multiple rounds of refinement through active data collection. Our findings indicate that, without requiring any hyper-parameter tuning, our sampling scheme stabilizes training, achieves higher reward scores, and maintains lower KL divergence from the reference model.

General Setup. We align the Llama-3.1-8B base model (Dubey et al., 2024) in terms of helpfulness and harmlessness using the HH-RLHF dataset (Bai et al., 2022), a widely-used benchmark dataset for alignment. It consists of 161k prompts in the training set. For response preference labeling, we

use a well-trained reward model to simulate human preferences by assigning preference to pairs of responses under the BT assumption in Equation (1). Specifically, we employ the Skywork-Reward-8B model (Liu et al., 2024a), a top-performing 8B model on RewardBench (Lambert et al., 2024), as our oracle $\mathcal{O}$. During training, interaction with this reward model is limited to providing two responses for comparison. We set $\beta = 0.1$ in all the experiments.

Supervised Fine-Tuning (SFT). To initialize training, following Rafailov et al. (2023), we first fine-tune the base model to obtain the SFT model as π_{ref} , which we fix as the reference model in all the experiments. We use the originally preferred responses from the HH-RLHF dataset as the SFT dataset and perform full-parameter tuning.

Evaluation. We present our results using the reward-KL curve, following Gao et al. (2023), with the reward evaluated by the oracle reward model $\mathcal{O}$. To monitor the impact of our sampling scheme on the optimization trajectory, we evaluate the model every 50 gradient steps during training. We use the entire testset of HH-RLHF (8.55K samples) to evaluate.

Baselines. We compare our sampling method against existing methods in Table 1. Since we treat the oracle $\mathcal{O}$ as a proxy for human labelers that can only provide pairwise preferences, all baselines are constrained to query the oracle with exactly two samples at a time. We thus adapt a *Best-of-N* variant that deploys the internal DPO reward to select the top and bottom candidates, which are then presented to the oracle for preference labeling, as listed in Table 1. We compare PILAF against the baselines: *Vanilla Sampling*, *Best-of-N Sampling* (with DPO reward), and *Hybrid Sampling* combined with a modified DPO loss (Xie et al., 2024).

Full experimental details can be found in Appendix F.

6.1 ITERATIVE DPO

Implementation. We first consider the iterative DPO framework (Xiong et al., 2024; Dong et al., 2024), in which preference data is collected in successive iterations rather than as a single fixed dataset. At the start of each iteration, a large dataset of responses is sampled using the current model, annotated for preferences, and then used to train the current model. Concretely, we set $n_t = |\mathcal{D}_\rho|$ in Algorithm 1, meaning that all prompts are used to generate new responses at each iteration. During the first iteration, when π_{ref} and π_θ are identical, PILAF reduces to *Vanilla Sampling*. Hence, we choose to focus our comparison on the second iteration. For consistency, we initialize all runs with the same policy model obtained at the end of the first iteration via *Vanilla Sampling*.

Results. Figure 2 presents the reward-KL curve for iterative DPO. PILAF significantly outperforms all the other methods: it achieves the end-point rewards of the baselines already around halfway through training, with around 40% less training time. This reduction directly translates to savings in both annotation and computational costs. We summarize the final performance in Table 2. PILAF produces a final policy with a high reward value and a modestly small KL divergence from the reference model, thereby achieving the highest overall objective J .

6.2 ONLINE DPO

Implementation. We further evaluate our sampling method in the online DPO setting (Guo et al., 2024), where new responses are generated and labeled at every training step, and these preference data are immediately used to update π_θ . This setting corresponds to the case where n_t (in Algorithm 1) is set to the training batch size, resulting in the most annotation-intensive and most actively on-policy alignment. By collecting and utilizing preference data on the fly for each batch, the policy is continuously refined using on-policy feedback throughout

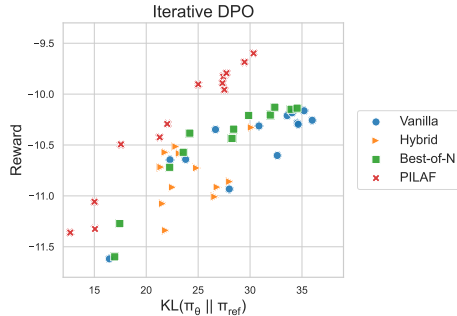


Figure 2: **Reward-KL curve for Iterative DPO.** All runs start from the same model obtained at the end of the first iteration via *Vanilla Sampling*. Each dot represents an evaluation performed every 50 training steps.

Table 2: **Results of Iterative DPO.** We report the average reward, KL divergence from the reference model, and objective J on the test set. Higher reward and J are better, while lower KL divergence is better. We use **boldface** to indicate the best result and underline to denote the second-best result.

METHOD	REWARD ($\uparrow$)	KL ($\downarrow$)	J ($\uparrow$)
<i>Vanilla</i>	-10.16	35.20	-13.68
<i>Best-of-N</i>	-10.13	32.38	-13.37
<i>Hybrid</i>	-10.51	22.86	<u>-12.80</u>
<i>PILAF</i>	-9.80	<u>25.01</u>	-12.30

the entire training process. Similar to Iterative DPO, we initialize all training runs with the same π_θ and focus on comparing the subsequent optimization. Further details are in Appendix F.

Results. Figure 3 demonstrates the effectiveness of PILAF in the pure online setting, and we summarize the final performance in Table 3. Compared with *Vanilla* and *Hybrid Sampling*, PILAF achieves a significantly better Reward-KL trade-off curve, attaining higher reward with lower KL. Although *Vanilla* eventually achieves roughly the same reward value as PILAF, it comes at the cost of a substantially higher KL. When compared with *Best-of-N*, PILAF traces a similar Reward-KL trajectory but ends with a higher reward and a better final objective after the same number of iterations, translating to lower sample complexity and reduced annotation and computational cost.

Robustness Analysis. Having established the effectiveness of PILAF, we further evaluate its robustness by testing whether it improves optimization and statistical convergence under challenging conditions, as predicted from our statistical theory in Section 4.2. Specifically, we replace the initial model with one that has overfit on a fixed off-policy dataset. This setup allows us to examine how different methods handle optimization starting from a poor initial point.

In Figure 4, we compare the performance of PILAF and *Vanilla Sampling* when both are initialized from an overfitted policy. We observe that *Vanilla Sampling* rapidly increases its KL divergence from the reference model while its reward improvement diminishes over time. In contrast, PILAF undergoes an early training phase with fluctuating KL values but ultimately attains a policy with higher reward and substantially lower KL divergence. We hypothesize that PILAF’s interpolation-based exploration design enables it to escape the suboptimal region of the loss landscape in which *Vanilla* remains. These results underscore PILAF’s effectiveness in more robustly optimizing overfitted (or even adversarially initialized) policies.

7 CONCLUSION

In this paper, we introduced Policy-Interpolated Learning for Aligned Feedback (PILAF), a novel sampling method designed to enhance response sampling for preference labeling. Theoretical analysis highlights PILAF’s superiority from both optimization and statistical perspectives, demonstrating its ability to stabilize training, accelerate convergence, and reduce variance. The method is straightforward to implement and requires no additional hyperparameter tuning. We empirically validated its performance in both iterative DPO and online DPO settings, where it consistently outperformed existing approaches. To achieve the same level of performance, PILAF consistently requires lower annotation costs, which can be substantial when annotations require experts in knowledge-intensive domains.

In future work, we hope to extend PILAF to other paradigms, such as KTO (Ethayarajh et al., 2024) and IPO (Azar et al., 2024). Due to resource constraints, our evaluations were conducted using 8B models and a reward model to simulate human feedback. Future studies involving larger-scale experiments and real human labeling would further generalize our findings.

Overall, this work takes an important step toward improving preference data curation in RLHF pipelines, laying the groundwork for more effective methods in alignment.

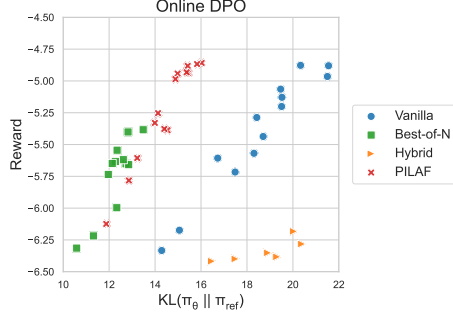


Figure 3: **Reward-KL curve for Online DPO.** Each dot represents an evaluation every 50 training steps.

Table 3: **Results of Online DPO.** We report the average reward, KL divergence from the reference model, and objective J on the testset.

METHOD	REWARD ($\uparrow$)	KL ($\downarrow$)	J ($\uparrow$)
<i>Vanilla</i>	-4.96	21.50	-7.11
<i>Best-of-N</i>	-5.54	12.35	-6.77
<i>Hybrid</i>	-6.42	16.46	-8.96
PILAF	-4.88	15.42	-6.42

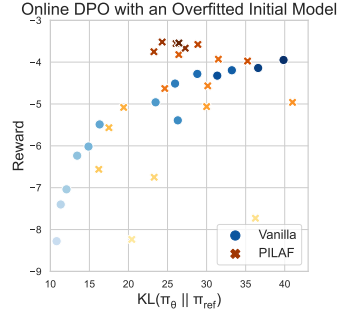


Figure 4: **Online DPO with an overfitted initial policy.** Each dot represents an evaluation performed every 50 training steps. Color saturation indicates the training step, with darker colors representing later steps.

REFERENCES

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- Mohammad Gheshlaghi Azar, Zhaohan Daniel Guo, Bilal Piot, Remi Munos, Mark Rowland, Michal Valko, and Daniele Calandriello. A general theoretical paradigm to understand learning from human preferences. In *International Conference on Artificial Intelligence and Statistics*, pp. 4447–4455. PMLR, 2024.
- 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.
- Shicong Cen, Jincheng Mei, Katayoon Goshvadi, Hanjun Dai, Tong Yang, Sherry Yang, Dale Schuurmans, Yuejie Chi, and Bo Dai. Value-incentivized preference optimization: A unified approach to online and offline rlhf. *arXiv preprint arXiv:2405.19320*, 2024.
- Nirjhar Das, Souradip Chakraborty, Aldo Pacchiano, and Sayak Ray Chowdhury. Active preference optimization for sample efficient rlhf. In *ICML 2024 Workshop on Theoretical Foundations of Foundation Models*, 2024.
- Thomas Degris, Martha White, and Richard S Sutton. Off-policy actor-critic. *arXiv preprint arXiv:1205.4839*, 2012.
- Hanze Dong, Wei Xiong, Deepanshu Goyal, Yihan Zhang, Winnie Chow, Rui Pan, Shizhe Diao, Jipeng Zhang, SHUM KaShun, and Tong Zhang. Raft: Reward ranked finetuning for generative foundation model alignment. *Transactions on Machine Learning Research*, 2023.
- Hanze Dong, Wei Xiong, Bo Pang, Haoxiang Wang, Han Zhao, Yingbo Zhou, Nan Jiang, Doyen Sahoo, Caiming Xiong, and Tong Zhang. Rlhf workflow: From reward modeling to online rlhf, 2024.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*, 2024.
- Kawin Ethayarajh, Winnie Xu, Niklas Muennighoff, Dan Jurafsky, and Douwe Kiela. Kto: Model alignment as prospect theoretic optimization. *arXiv preprint arXiv:2402.01306*, 2024.
- Leo Gao, John Schulman, and Jacob Hilton. Scaling laws for reward model overoptimization. In *International Conference on Machine Learning*, pp. 10835–10866, 2023.
- Evan Greensmith, Peter L Bartlett, and Jonathan Baxter. Variance reduction techniques for gradient estimates in reinforcement learning. *Journal of Machine Learning Research*, 5(9), 2004.
- Shixiang Shane Gu, Timothy Lillicrap, Richard E Turner, Zoubin Ghahramani, Bernhard Schölkopf, and Sergey Levine. Interpolated policy gradient: Merging on-policy and off-policy gradient estimation for deep reinforcement learning. *Advances in neural information processing systems*, 30, 2017.
- Shangmin Guo, Biao Zhang, Tianlin Liu, Tianqi Liu, Misha Khalman, Felipe Llinares, Alexandre Rame, Thomas Mesnard, Yao Zhao, Bilal Piot, et al. Direct language model alignment from online ai feedback. *arXiv preprint arXiv:2402.04792*, 2024.
- Jian Hu, Xibin Wu, Zilin Zhu, Xianyu, Weixun Wang, Dehao Zhang, and Yu Cao. Openrlhf: An easy-to-use, scalable and high-performance rlhf framework. *arXiv preprint arXiv:2405.11143*, 2024.
- Kaixuan Ji, Jiafan He, and Quanquan Gu. Reinforcement learning from human feedback with active queries. *arXiv preprint arXiv:2402.09401*, 2024.

- Nathan Kallus and Masatoshi Uehara. Statistically efficient off-policy policy gradients. In *International Conference on Machine Learning*, pp. 5089–5100. PMLR, 2020.
- Michael R Kosorok. *Introduction to empirical processes and semiparametric inference*, volume 61. Springer, 2008.
- Nathan Lambert, Valentina Pyatkin, Jacob Morrison, LJ Miranda, Bill Yuchen Lin, Khyathi Chandu, Nouha Dziri, Sachin Kumar, Tom Zick, Yejin Choi, Noah A. Smith, and Hannaneh Hajishirzi. Rewardbench: Evaluating reward models for language modeling. <https://huggingface.co/spaces/allenai/reward-bench>, 2024.
- Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step, 2023. URL <https://arxiv.org/abs/2305.20050>.
- Fernando Antonio Flores Limo, David Raul Hurtado Tiza, Maribel Mamani Roque, Edward Espinoza Herrera, José Patricio Muñoz Murillo, Jorge Jinchuñá Huallpa, Victor Andre Ariza Flores, Alejandro Guadalupe Rincón Castillo, Percy Fritz Puga Peña, Christian Paolo Martel Carranza, et al. Personalized tutoring: Chatgpt as a virtual tutor for personalized learning experiences. *Przestrzeń Społeczna (Social Space)*, 23(1):293–312, 2023.
- Chris Yuhao Liu, Liang Zeng, Jiakai Liu, Rui Yan, Jujie He, Chaojie Wang, Shuicheng Yan, Yang Liu, and Yahui Zhou. Skywork-reward: Bag of tricks for reward modeling in llms. *arXiv preprint arXiv:2410.18451*, 2024a.
- Tianlin Liu, Shangmin Guo, Leonardo Bianco, Daniele Calandriello, Quentin Berthet, Felipe Llinares, Jessica Hoffmann, Lucas Dixon, Michal Valko, and Mathieu Blondel. Decoding-time realignment of language models. *arXiv preprint arXiv:2402.02992*, 2024b.
- Tianqi Liu, Yao Zhao, Rishabh Joshi, Misha Khalman, Mohammad Saleh, Peter J Liu, and Jialu Liu. Statistical rejection sampling improves preference optimization. In *The Twelfth International Conference on Learning Representations*, 2024c.
- Zichen Liu, Changyu Chen, Chao Du, Wee Sun Lee, and Min Lin. Sample-efficient alignment for llms. *arXiv preprint arXiv:2411.01493*, 2024d.
- Viraj Mehta, Vikramjeet Das, Ojash Neopane, Yijia Dai, Ilija Bogunovic, Jeff Schneider, and Willie Neiswanger. Sample efficient reinforcement learning from human feedback via active exploration. *arXiv preprint arXiv:2312.00267*, 2023.
- William Muldrew, Peter Hayes, Mingtian Zhang, and David Barber. Active preference learning for large language models. In *Forty-first International Conference on Machine Learning*, 2024.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744, 2022.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36, 2023.
- Antoine Scheid, Etienne Boursier, Alain Durmus, Michael I Jordan, Pierre Ménard, Eric Moulines, and Michal Valko. Optimal design for reward modeling in rlhf. *arXiv preprint arXiv:2410.17055*, 2024.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- Ruizhe Shi, Runlong Zhou, and Simon S Du. The crucial role of samplers in online direct preference optimization. *arXiv preprint arXiv:2409.19605*, 2024.

- David Silver, Guy Lever, Nicolas Heess, Thomas Degris, Daan Wierstra, and Martin Riedmiller. Deterministic policy gradient algorithms. In *International conference on machine learning*, pp. 387–395. Pmlr, 2014.
- Richard S Sutton, David McAllester, Satinder Singh, and Yishay Mansour. Policy gradient methods for reinforcement learning with function approximation. *Advances in neural information processing systems*, 12, 1999.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.
- Ronald J Williams. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine learning*, 8:229–256, 1992.
- Tengyang Xie, Dylan J Foster, Akshay Krishnamurthy, Corby Rosset, Ahmed Awadallah, and Alexander Rakhlin. Exploratory preference optimization: Harnessing implicit q^* -approximation for sample-efficient rlhf. *arXiv preprint arXiv:2405.21046*, 2024.
- Wei Xiong, Hanze Dong, Chenlu Ye, Ziqi Wang, Han Zhong, Heng Ji, Nan Jiang, and Tong Zhang. Iterative preference learning from human feedback: Bridging theory and practice for rlhf under kl-constraint. In *Forty-first International Conference on Machine Learning*, 2024.
- Haoran Xu, Amr Sharaf, Yunmo Chen, Weiting Tan, Lingfeng Shen, Benjamin Van Durme, Kenton Murray, and Young Jin Kim. Contrastive preference optimization: Pushing the boundaries of llm performance in machine translation. In *Forty-first International Conference on Machine Learning*, 2024.
- Jing Xu, Andrew Lee, Sainbayar Sukhbaatar, and Jason Weston. Some things are more cringe than others: Preference optimization with the pairwise cringe loss. *arXiv preprint arXiv:2312.16682*, 2023.
- Zhenrui Yue, Honglei Zhuang, Aijun Bai, Kai Hui, Rolf Jagerman, Hansi Zeng, Zhen Qin, Dong Wang, Xuanhui Wang, and Michael Bendersky. Inference scaling for long-context retrieval augmented generation. *arXiv preprint arXiv:2410.04343*, 2024.
- Shenao Zhang, Donghan Yu, Hiteshi Sharma, Han Zhong, Zhihan Liu, Ziyi Yang, Shuohang Wang, Hany Hassan, and Zhaoran Wang. Self-exploring language models: Active preference elicitation for online alignment. *arXiv preprint arXiv:2405.19332*, 2024.

CONTENTS

A. Additional Literature Review	13
B. Additional Statistical Results	14
C. Proof of Main Results	15
C.1. Optimization Considerations: Proof of Theorem 4.1	15
C.1.1. Building Blocks	15
C.1.2. Derivation of Theorem 4.1	16
C.1.3. Proof of Claim equation 15	17
C.2. Statistical Considerations	18
C.2.1. Proof of Lemma B.1 (Theorem C.4)	18
C.2.2. Proof of Theorem 4.2	19
C.2.3. Proof of Theorem 4.3	20
D. Proof of Auxiliary Results	21
D.1. Proof of Auxiliary Results for Theorem 4.1	21
D.1.1. Gradients of Policy π_θ and Reward r_θ	21
D.1.2. Proof of Lemma C.2, Explicit Form of Gradient $\nabla_\theta J(\pi_\theta)$	22
D.1.3. Proof of Lemma C.3, Explicit Form of Gradient $\nabla_\theta \mathcal{L}(\theta)$	24
D.2. Proof of Auxiliary Results for Theorem 4.2	25
D.2.1. Proof of Condition equation 25	26
D.2.2. Proof of Lemma C.5, Explicit Form of Hessian $\nabla_\theta^2 \mathcal{L}(\theta^*)$	27
D.2.3. Proof of Lemma C.6, Asymptotic Distribution of Gradient $\nabla_\theta \widehat{\mathcal{L}}(\theta^*)$	27
D.3. Proof of Auxiliary Results for Theorem 4.3	28
D.3.1. Proof of Equation equation 9 from Theorem 4.3, Explicit Form of Hessian $\nabla_\theta^2 J(\pi^*)$	28
D.3.2. Proof of the Asymptotic Distribution in Equation equation 32	29
D.3.3. Proof of the Tail Bound in Equation equation 10	30
E. Supporting Theorem: Master Theorem for Z -Estimators	30
F. Experimental Details	31
G. Extension to Proximal Policy Optimization (PPO)	32

A ADDITIONAL LITERATURE REVIEW

RLHF. RLHF has emerged as a cornerstone methodology for aligning large language models with human values and preferences (Achiam et al., 2023). Early systems (Ouyang et al., 2022) turn human preference data into reward modeling to optimize model behavior accordingly. DPO has been proposed as a more efficient approach that directly trains LLMs on preference data. As LLMs evolve during training, continuing training on pre-generated preference data becomes suboptimal due to the distribution shift. Empirically, RLHF is applied iteratively—generating on-policy data at successive stages to enhance alignment and performance (Touvron et al., 2023; Bai et al., 2022). Similarly, researchers have introduced iterative DPO (Xiong et al., 2024; Xu et al., 2023) and online DPO (Guo et al., 2024) to fully leverage online preference labeling. Ultimately, the quality of preference data play a critical role in determining the effectiveness of the alignment.

Sampling in Frontier LLMs. Technical reports of Frontier LLMs briefly mention sampling techniques. For instance, Claude (Bai et al., 2022) utilizes models from different training steps to generate responses, while Llama-2 (Touvron et al., 2023) further use different temperatures for sampling. However, no further details are provided, leaving the development of a principled method an open challenge.

Data Selection. There is a line of research aimed at improving sample efficiency for preference labeling by selecting question and response pairs. Scheid et al. (2024) conceptualize this as a regret minimization problem, leveraging methods from linear dueling bandits. Das et al. (2024); Mehta et al. (2023); Muldrew et al. (2024); Ji et al. (2024) draw insights from active learning, using various uncertainty estimators to guide selection by prioritizing sample pairs with maximum uncertainty. These approaches focus directly on a dataset of questions and responses and are orthogonal to our work.

Other Changes in Response Sampling. Several works also modify the sampling design directly (Liu et al., 2024c; Dong et al., 2023), but with the goal of improving policy network optimization based on a reward model, rather than enhancing the reward modeling itself. Liu et al. (2024c) employ rejection sampling to approximate the response distribution induced by the reward model, thereby improving optimization. However, this approach requires access to the reward model and incurs higher computational and labeling costs. Similarly, Dong et al. (2023) use best-of-N sampling with the reward model to generate high-quality data for supervised fine-tuning (SFT). We consider these approaches orthogonal to our work.

Additionally, Cen et al. (2024) introduce a bonus term in the policy learning phase of online RLHF to promote exploration in response sampling, which aligns with the optimism principle.

B ADDITIONAL STATISTICAL RESULTS

In addition to our analysis of T-PILAF in Section 3, here we present a generalized version of Theorem 4.2 that applies to any response sampling distribution μ . While not directly tied to the main focus of this work, this broader result may be of independent interest to readers. The proof of Lemma B.1 is provided in Appendix C.2.1.

Lemma B.1. *For a general sampling distribution μ , the statement in Theorem 4.2 remains valid with the matrix $\Sigma_\star$ redefined as*

$$\Sigma_\star := \mathbb{E}_{x \sim \rho, (\vec{y}^a, \vec{y}^b) \sim \mu(\cdot|x)} \left[w(x) \cdot \text{Var}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} \mid x, \vec{y}^a, \vec{y}^b) \cdot \mathbf{g} \mathbf{g}^\top \right], \quad (11)$$

where the expectation is taken over the distribution

$$\bar{\mu}(\vec{y}^a, \vec{y}^b \mid x) := \frac{1}{2} \{ \mu(\vec{y}^a, \vec{y}^b \mid x) + \mu(\vec{y}^b, \vec{y}^a \mid x) \}. \quad (12a)$$

The variance term is specified as

$$\text{Var}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} \mid x, \vec{y}^a, \vec{y}^b) = \sigma(r^\star(x, \vec{y}^a) - r^\star(x, \vec{y}^b)) \sigma(r^\star(x, \vec{y}^b) - r^\star(x, \vec{y}^a)) \quad (12b)$$

and the gradient difference $\mathbf{g}$ is defined as

$$\mathbf{g} := \nabla_\theta r^\star(x, \vec{y}^a) - \nabla_\theta r^\star(x, \vec{y}^b). \quad (12c)$$

The general form of the matrix $\Sigma_\star$ offers valuable insights for designing a sampling scheme. To ensure $\Sigma_\star$ is well-conditioned (less singular), we must balance two key factors when selecting responses $\vec{y}^a$ and $\vec{y}^b$:

Large variance: The variance in definition equation 12b should be maximized. This occurs when $r^\star(x, \vec{y}^a) \approx r^\star(x, \vec{y}^b)$. Intuitively, preference feedback is most informative when annotators compare responses of similar quality.

Large gradient difference: The gradient difference $\mathbf{g}$ from definition equation 12c should also be large. This requires responses with significantly different gradients. Only then can the comparison provide a clear and meaningful direction for model training.

C PROOF OF MAIN RESULTS

This section provides the proofs of the main results from Section 4, covering both optimization and statistical aspects. In Appendix C.1, we prove Theorem 4.1, which establishes the gradient alignment property. For the statistical results, Appendix C.2 begins with the proofs of Lemma B.1 and Theorem 4.2, which derive the asymptotic distribution of the estimated parameter $\hat{\theta}$, and concludes with the proof of Theorem 4.3, analyzing the asymptotic behavior of the value gap $J(\pi^*) - J(\hat{\pi})$.

C.1 OPTIMIZATION CONSIDERATIONS: PROOF OF THEOREM 4.1

We begin by presenting a rigorous restatement of Theorem 4.1, formally detailed in Theorem C.1 below.

Theorem C.1 (Gradient structure in DPO training). *Consider the expected loss function $\mathcal{L}(\theta)$ during the DPO training phase. Using data collected from our proposed response sampling scheme μ , the gradient of $\mathcal{L}(\theta)$ satisfies*

$$\nabla_{\theta} \mathcal{L}(\theta) = -\frac{\beta}{\bar{Z}_{\theta}} \nabla_{\theta} J(\pi_{\theta}) + T_2,$$

where the constant $\bar{Z}_{\theta}$ is defined in equation 7, and the term T_2 represents a second-order error.

To control term T_2 , assume the following uniform bounds:

$$(i) \|r^*\|_{\infty} \leq R.$$

$$(ii) \text{ For any policy } \pi_{\theta} \in \Pi, \text{ the induced reward } r_{\theta} \text{ satisfies } \|r_{\theta}\|_{\infty} \leq R \text{ and } \sup_{x, \bar{\mathbf{y}}} \|\nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}})\|_2 \leq G.$$

Under these conditions, T_2 is bounded as

$$\|T_2\|_2 \leq C \cdot \mathbb{E}_{x \sim \rho, \bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b \sim \pi_{\theta}(\cdot|x)} \left[\left\{ (r^*(x, \bar{\mathbf{y}}^a) - r^*(x, \bar{\mathbf{y}}^b)) - (r_{\theta}(x, \bar{\mathbf{y}}^a) - r_{\theta}(x, \bar{\mathbf{y}}^b)) \right\}^2 \right],$$

where the constant C is given by $C = 0.1 (1 + e^{2R}) G / \bar{Z}_{\theta}$.

The proof of Theorem C.1 is structured into three sections. In Appendix C.1.1, we lay the foundation by presenting the key components, including the explicit expressions for the gradients $\nabla_{\theta} J(\pi_{\theta})$ and $\nabla_{\theta} \mathcal{L}(\theta)$, as well as for the sampling density $\bar{\mu}$. Then Appendix C.1.2 establishes the connection between $\nabla_{\theta} J(\pi_{\theta})$ and $\nabla_{\theta} \mathcal{L}(\theta)$ by leveraging these results, completing the proof of Theorem 4.1. Finally, in Appendix C.1.3, we provide a detailed derivation of the form of density function $\bar{\mu}$.

C.1.1 BUILDING BLOCKS

To establish Theorem 4.1, which uncovers the relationship between the gradients of the expected value $J(\pi_{\theta})$ and the negative log-likelihood function $\mathcal{L}(\theta)$, the first step is to derive explicit expressions for the gradients of both functions. The results are presented in Lemmas C.2 and C.3, with detailed proofs provided in Appendices D.1.2 and D.1.3, respectively.

Lemma C.2 (Gradient of value $J(\pi_{\theta})$). *For any π_{θ} in the parameterized policy class Π , the gradient of the expected value $J(\pi_{\theta})$ satisfies*

$$\begin{aligned} \nabla_{\theta} J(\pi_{\theta}) = \frac{1}{2\beta} \mathbb{E}_{x \sim \rho; \bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b \sim \pi_{\theta}(\cdot|x)} & \left[\left\{ (r^*(x, \bar{\mathbf{y}}^a) - r^*(x, \bar{\mathbf{y}}^b)) - (r_{\theta}(x, \bar{\mathbf{y}}^a) - r_{\theta}(x, \bar{\mathbf{y}}^b)) \right\} \right. \\ & \left. \cdot \left\{ \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^a) - \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^b) \right\} \right]. \quad (13) \end{aligned}$$

Lemma C.3 (Gradient of the loss function $\mathcal{L}(\theta)$). *For any π_{θ} in the parameterized policy class Π and any sampling distribution μ of the responses, the gradient of the negative log-likelihood function*

$\mathcal{L}(\theta)$ is given by

$$\begin{aligned} \nabla_{\theta} \mathcal{L}(\theta) = & -\mathbb{E}_{x \sim \rho; (\vec{y}^a, \vec{y}^b) \sim \bar{\mu}(\cdot|x)} \left[w(x) \cdot \left\{ \sigma(r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)) - \sigma(r_{\theta}(x, \vec{y}^a) - r_{\theta}(x, \vec{y}^b)) \right\} \right. \\ & \left. \cdot \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}^a) - \nabla_{\theta} r_{\theta}(x, \vec{y}^b) \right\} \right], \quad (14a) \end{aligned}$$

where the average density $\bar{\mu}$ is defined as

$$\bar{\mu}(\vec{y}^a, \vec{y}^b | x) := \frac{1}{2} \{ \mu(\vec{y}^a, \vec{y}^b | x) + \mu(\vec{y}^b, \vec{y}^a | x) \} \quad (14b)$$

as previously introduced in Equation (12a).

In Lemma C.3, we observe that the gradient $\nabla_{\theta} \mathcal{L}(\theta)$ is expressed as an expectation over the probability distribution $\bar{\mu}$. By applying the sampling scheme outlined in Section 3, we can derive a more detailed representation of $\nabla_{\theta} \mathcal{L}(\theta)$. This refined form will reveal its close relationship to the gradient $\nabla_{\theta} J(\pi_{\theta})$ given in expression equation 13.

Before moving forward, it is crucial for us to first derive the explicit form of $\bar{\mu}$. Specifically, we claim that the distribution $\bar{\mu}$ satisfies the following property

$$\frac{\bar{\mu}(\vec{y}^a, \vec{y}^b | x)}{\pi_{\theta}(\vec{y}^a | x) \pi_{\theta}(\vec{y}^b | x)} = \frac{1}{2 \{1 + Z_{\theta}^+(x) Z_{\theta}^-(x)\}} \cdot \frac{1}{\sigma'(r_{\theta}(x, \vec{y}^a) - r_{\theta}(x, \vec{y}^b))}, \quad (15)$$

where σ' denotes the derivative of the sigmoid function σ , given by

$$\sigma'(z) = \frac{1}{(1 + \exp(-z))(1 + \exp(z))} = \sigma(z) \sigma(-z) \quad \text{for any } z \in \mathbb{R}. \quad (16)$$

With these key components in place, we are now prepared to prove Theorem 4.1.

C.1.2 DERIVATION OF THEOREM 4.1

With the tools provided by Lemmas C.2 and C.3 and the sampling density expression in equation 15, we are now ready to prove Theorem 4.1.

We begin by applying Lemma C.3 and reformulating equation equation 14a as

$$\begin{aligned} \nabla_{\theta} \mathcal{L}(\theta) = & -\mathbb{E}_{x \sim \rho; \vec{y}^a, \vec{y}^b \sim \pi_{\theta}(\cdot|x)} \left[w(x) \cdot \frac{\bar{\mu}(\vec{y}^a, \vec{y}^b | x)}{\pi_{\theta}(\vec{y}^a | x) \pi_{\theta}(\vec{y}^b | x)} \right. \\ & \cdot \left\{ \sigma(r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)) - \sigma(r_{\theta}(x, \vec{y}^a) - r_{\theta}(x, \vec{y}^b)) \right\} \\ & \left. \cdot \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}^a) - \nabla_{\theta} r_{\theta}(x, \vec{y}^b) \right\} \right]. \quad (17) \end{aligned}$$

Substituting the density ratio from equation equation 15 into expression equation 17 and incorporating the weight function $w(x)$ defined in equation equation 7, we obtain

$$\begin{aligned} \nabla_{\theta} \mathcal{L}(\theta) = & -\frac{1}{2 Z_{\theta}} \mathbb{E}_{x \sim \rho; \vec{y}^a, \vec{y}^b \sim \pi_{\theta}(\cdot|x)} \left[\frac{\sigma(r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)) - \sigma(r_{\theta}(x, \vec{y}^a) - r_{\theta}(x, \vec{y}^b))}{\sigma'(r_{\theta}(x, \vec{y}^a) - r_{\theta}(x, \vec{y}^b))} \right. \\ & \left. \cdot \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}^a) - \nabla_{\theta} r_{\theta}(x, \vec{y}^b) \right\} \right]. \quad (18) \end{aligned}$$

Using the intuition that the first-order Taylor expansion

$$\frac{\sigma(z^*) - \sigma(z)}{\sigma'(z)} = (z^* - z) + \mathcal{O}((z^* - z)^2)$$

is valid when $z \rightarrow z^*$, with $z^* := r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)$ and $z := r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b)$, we find that

$$\begin{aligned} & \frac{\sigma(r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)) - \sigma(r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b))}{\sigma'(r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b))} \\ &= \left\{ (r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)) - (r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b)) \right\} + \text{second-order term.} \end{aligned}$$

Reformulating equation equation 18 in this context, we rewrite it as

$$\begin{aligned} \nabla_\theta \mathcal{L}(\phi) = & -\frac{1}{2\bar{Z}_\theta} \mathbb{E}_{x \sim p; \vec{y}^a, \vec{y}^b \sim \pi_\theta(\cdot|x)} \left[\left\{ (r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)) - (r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b)) \right\} \right. \\ & \left. \cdot \left\{ \nabla_\theta r_\theta(x, \vec{y}^a) - \nabla_\theta r_\theta(x, \vec{y}^b) \right\} \right] + T_2, \end{aligned} \quad (19)$$

where T_2 represents the second-order residual term related to the estimation error $r_\theta - r^*$. By applying Lemma C.2, we observe that the primary term in equation equation 19 aligns with the direction of $\nabla_\theta J(\pi_\theta)$, resulting in

$$\nabla_\theta \mathcal{L}(\phi) = -\frac{\beta}{\bar{Z}_\theta} \nabla_\theta J(\pi_\theta) + T_2. \quad (20)$$

Next, we proceed to control the second-order term T_2 . The conditions

$$\|r^*\|_\infty, \|r_\theta\|_\infty \leq R \quad \text{and} \quad \sup_{(x, \vec{y}) \in \mathcal{X} \times \mathcal{Y}} \|\nabla_\theta r_\theta(x, \vec{y})\|_2 \leq G,$$

lead to the bound

$$\left| \frac{\sigma(z^*) - \sigma(z)}{\sigma'(z)} - (z^* - z) \right| \leq 0.1 (1 + e^{2R}) \cdot (z^* - z)^2,$$

which in turn implies

$$\|T_2\|_2 \leq \frac{0.1 (1 + e^{2R}) G}{\bar{Z}_\theta} \mathbb{E}_{x \sim p; \vec{y}^a, \vec{y}^b \sim \pi_\theta(\cdot|x)} \quad (21)$$

$$\left[\left\{ (r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)) - (r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b)) \right\}^2 \right]. \quad (22)$$

Finally, combining equation equation 21 with equation equation 20, we conclude the proof of Theorem 4.1.

C.1.3 PROOF OF CLAIM EQUATION 15

The remaining step in the proof of Theorem 4.1 is to verify the expression for the density ratio in equation equation 15.

Based on the sampling scheme described in Section 3, we find that the sampling distribution for the response satisfies

$$\mu(\vec{y}^a, \vec{y}^b | x) = \{1 - p_0(x)\} \cdot \pi_\theta(\vec{y}^a | x) \pi_\theta(\vec{y}^b | x) + p_0(x) \cdot \pi_\theta^+(\vec{y}^a | x) \pi_\theta^-(\vec{y}^b | x), \quad (23)$$

where the probability $p_0(x)$ is defined as

$$p_0(x) = Z_\theta^+(x) Z_\theta^-(x) / \{1 + Z_\theta^+(x) Z_\theta^-(x)\}$$

and the policies π_θ^+ and π_θ^- are specified in equations equation 6a and equation 6b, respectively. This allows us to simplify equation equation 23 to

$$\mu(\vec{y}^a, \vec{y}^b | x) = \frac{\pi_\theta(\vec{y}^a | x) \pi_\theta(\vec{y}^b | x)}{1 + Z_\theta^+(x) Z_\theta^-(x)} \left\{ 1 + \exp \{r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b)\} \right\}.$$

Similarly, we derive an expression for $\mu(\vec{y}^b, \vec{y}^a | x)$. By averaging the two expressions, for $\mu(\vec{y}^a, \vec{y}^b | x)$ and $\mu(\vec{y}^b, \vec{y}^a | x)$, we obtain

$$\frac{\bar{\mu}(\vec{y}^a, \vec{y}^b | x)}{\pi_\theta(\vec{y}^a | x) \pi_\theta(\vec{y}^b | x)} = \frac{\pi_\theta(\vec{y}^a | x) \pi_\theta(\vec{y}^b | x)}{2 \{1 + Z_\theta^+(x) Z_\theta^-(x)\}} \left\{ 2 + \exp \{r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b)\} + \exp \{r_\theta(x, \vec{y}^b) - r_\theta(x, \vec{y}^a)\} \right\}.$$

Rewriting this expression using the formula for σ' in equation equation 16, we arrive at

$$\begin{aligned} & \{1 + Z_\theta^+(x) Z_\theta^-(x)\} \cdot \frac{\bar{\mu}(\vec{y}^a, \vec{y}^b | x)}{\pi_\theta(\vec{y}^a | x) \pi_\theta(\vec{y}^b | x)} \\ &= \frac{1}{2} \left\{ 1 + \exp \{r_\theta(x, \vec{y}^b) - r_\theta(x, \vec{y}^a)\} \right\} \left\{ 1 + \exp \{r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b)\} \right\} \\ &= \frac{1}{2 \sigma' (r_\theta(x, \vec{y}^a) - r_\theta(x, \vec{y}^b))}. \end{aligned}$$

Finally, rearranging terms, we recover equation equation 15, completing this part of the proof.

C.2 STATISTICAL CONSIDERATIONS

In this section, we present the proofs for Theorems 4.2 and 4.3 and Lemma B.1 from Section 4.2. We start with the proof of Lemma B.1 in Appendix C.2.1, with a rigorous restatement provided in Theorem C.4 below.

Theorem C.4. Assume the reward model r^* in the BT model equation 1 satisfies $r^* = r_{\theta^*}$ for some parameter θ^* . Assume that $\hat{\theta}$ minimizes the loss function $\hat{\mathcal{L}}(\theta)$ in the sense that $\sqrt{n} \nabla_\theta \hat{\mathcal{L}}(\hat{\theta}) \xrightarrow{P} \mathbf{0}$ and that $\hat{\theta} \xrightarrow{P} \theta^*$ as the sample size $n \rightarrow \infty$. Additionally, suppose the reward function $r_\theta(x, \vec{y})$, its gradient $\nabla_\theta r_\theta(x, \vec{y})$ and its Hessian $\nabla_\theta^2 r_\theta(x, \vec{y})$ are uniformly bounded and Lipchitz continuous with respect to θ , for all $(x, \vec{y}) \in \mathcal{X} \times \mathcal{Y}$.

Under these conditions, the estimate $\hat{\theta}$ asymptotically follows a Gaussian distribution

$$\sqrt{n} (\hat{\theta} - \theta^*) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \Omega) \quad \text{as } n \rightarrow \infty.$$

We have an estimate of the covariance matrix Ω :

$$\Omega \preceq \|w\|_\infty \cdot \Sigma_\star^{-1}.$$

For a general sampling scheme μ chosen, the matrix $\Sigma_\star$ is given by

$$\Sigma_\star := \mathbb{E}_{x \sim p, (\vec{y}^a, \vec{y}^b) \sim \bar{\mu}(\cdot | x)} \left[w(x) \cdot \text{Var}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} | x, \vec{y}^a, \vec{y}^b) \cdot \mathbf{g} \mathbf{g}^\top \right],$$

where the expectation is taken over the distribution

$$\bar{\mu}(\vec{y}^a, \vec{y}^b | x) := \frac{1}{2} \{ \mu(\vec{y}^a, \vec{y}^b | x) + \mu(\vec{y}^b, \vec{y}^a | x) \}.$$

The variance term is specified as

$$\text{Var}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} | x, \vec{y}^a, \vec{y}^b) = \sigma(r^*(x, \vec{y}^a) - r^*(x, \vec{y}^b)) \sigma(r^*(x, \vec{y}^b) - r^*(x, \vec{y}^a))$$

and the gradient difference $\mathbf{g}$ is defined as

$$\mathbf{g} := \nabla_\theta r^*(x, \vec{y}^a) - \nabla_\theta r^*(x, \vec{y}^b).$$

Theorem C.4 establishes the asymptotic distribution of the estimated parameter $\hat{\theta}$, which serves as the foundation for the subsequent results. Next, we show that Theorem 4.2 directly follows as a corollary of Theorem C.4, with the detailed derivation provided in Appendix C.2.2. Finally, in Appendix C.2.3, we prove Theorem 4.3, which describes the asymptotic behavior of the value gap $J(\pi^*) - J(\hat{\pi})$.

C.2.1 PROOF OF LEMMA B.1 (THEOREM C.4)

In this section, we analyze the asymptotic distribution of the estimated parameter $\hat{\theta}$ for a general sampling distribution μ . The parameter $\hat{\theta}$ is obtained by solving the optimization problem

$$\text{minimize}_\theta \quad \hat{\mathcal{L}}(\theta) := -\frac{1}{n} \sum_{i=1}^n w(x_i) \cdot \log \sigma(r_\theta(x_i, \vec{y}_i^w) - r_\theta(x_i, \vec{y}_i^\ell)).$$

We assume the optimization is performed to sufficient accuracy such that $\nabla_{\theta} \widehat{\mathcal{L}}(\widehat{\theta}) = o_p(n^{-\frac{1}{2}})$. Under this condition, $\widehat{\theta}$ qualifies as a Z -estimator. To study its asymptotic behavior, we use the master theorem for Z -estimators (Kosorok, 2008), the formal statement of which is provided in Theorem E.1 in Appendix E.

To apply the master theorem, we set $\Psi := \nabla_{\theta} \mathcal{L}$ and $\Psi_n := \nabla_{\theta} \widehat{\mathcal{L}}$ and verify the conditions. In particular, the smoothness condition equation 62 in Theorem E.1 translates to the following equation in our context:

$$\sqrt{n} \{ \nabla_{\theta} \widehat{\mathcal{L}}(\widehat{\theta}) - \nabla_{\theta} \mathcal{L}(\widehat{\theta}) \} - \sqrt{n} \{ \nabla_{\theta} \widehat{\mathcal{L}}(\theta^*) - \nabla_{\theta} \mathcal{L}(\theta^*) \} = o_p(1 + \sqrt{n} \|\widehat{\theta} - \theta^*\|_2). \quad (25)$$

This condition follows from the second-order smoothness of the reward function r_{θ} with respect to θ . A rigorous proof is provided in Appendix D.2.1.

We now provide the explicit form of the derivative $\dot{\Psi}_{\theta^*} = \nabla_{\theta}^2 \mathcal{L}(\theta^*)$, as captured in the following lemma. The proof of this result can be found in Appendix D.2.2.

Lemma C.5. *The Hessian matrix of the population loss $\mathcal{L}(\theta)$ at $\theta = \theta^*$ is*

$$\nabla_{\theta}^2 \mathcal{L}(\theta^*) = \Sigma_{\star}, \quad (26)$$

where the matrix $\Sigma_{\star}$ is defined in equation equation 11.

Next, we analyze the asymptotic behavior of the gradient $\nabla_{\theta} \widehat{\mathcal{L}}(\theta^*)$. The proof is deferred to Appendix D.2.3.

Lemma C.6. *The gradient of the empirical loss $\widehat{\mathcal{L}}(\theta)$ at $\theta = \theta^*$ satisfies*

$$\sqrt{n} (\nabla_{\theta} \widehat{\mathcal{L}}(\theta^*) - \nabla_{\theta} \mathcal{L}(\theta^*)) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \widetilde{\Omega}) \quad \text{as } n \rightarrow \infty, \quad (27a)$$

where the covariance matrix $\widetilde{\Omega} \in \mathbb{R}^{d \times d}$ is bounded as follows:

$$\widetilde{\Omega} \preceq \|w\|_{\infty} \cdot \Sigma_{\star}, \quad (27b)$$

with $\Sigma_{\star}$ defined in equation equation 11.

Combining these results, and assuming $\Sigma_{\star}$ is nonsingular, the master theorem (Theorem E.1) yields the asymptotic distribution of $\widehat{\theta}$:

$$\sqrt{n} (\widehat{\theta} - \theta^*) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \Sigma_{\star}^{-1} \widetilde{\Omega} \Sigma_{\star}^{-1}).$$

Furthermore, from the bound equation 27b, the covariance matrix $\Omega; := \Sigma_{\star}^{-1} \widetilde{\Omega} \Sigma_{\star}^{-1}$ satisfies

$$\Omega = \Sigma_{\star}^{-1} \widetilde{\Omega} \Sigma_{\star}^{-1} \preceq \|w\|_{\infty} \cdot \Sigma_{\star}^{-1}.$$

Therefore, we have established the asymptotic distribution of $\widehat{\theta}$, completing the proof of Lemma B.1.

C.2.2 PROOF OF THEOREM 4.2

Theorem 4.2 is a direct corollary of Lemma B.1, using our specific choice of sampling distribution μ . To establish this, we demonstrate how the general covariance matrix $\Sigma_{\star}$ in equation equation 11 simplifies to the form in equation equation 8 under our proposed sampling scheme.

To establish the result in this section, we impose the following regularity condition: There exists a constant $C \geq 1$ satisfying

$$\text{Var}_{r_{\theta}}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} \mid x, \vec{y}^a, \vec{y}^b) \leq C \cdot \text{Var}_{r^*}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} \mid x, \vec{y}^a, \vec{y}^b) \quad (28)$$

for any prompt $x \in \mathcal{X}$ and responses $\vec{y}^a, \vec{y}^b \in \mathcal{Y}$. Here $\text{Var}_{r_{\theta}}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} \mid x, \vec{y}^a, \vec{y}^b)$ denotes the conditional variance under the BT model equation 1, when the implicit reward function r^* is replaced by r_{θ} . The term $\text{Var}_{r^*}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} \mid x, \vec{y}^a, \vec{y}^b) \equiv \text{Var}(\mathbb{1}\{\vec{y}^a = \vec{y}^w\} \mid x, \vec{y}^a, \vec{y}^b)$ represents the conditional variance under the ground-truth BT model, where the reward function is given by r^* .

We begin by leveraging the property of the sampling distribution μ from equation equation 15 and the derivative σ' of the sigmoid function σ , given in equation equation 16. Specifically, we find that

$$\begin{aligned} \frac{\bar{\mu}(\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b | x)}{\pi_\theta(\bar{\mathbf{y}}^a | x) \pi_\theta(\bar{\mathbf{y}}^b | x)} &= \frac{1}{2 \{1 + Z_\theta^+(x) Z_\theta^-(x)\}} \cdot \frac{1}{\sigma(r_\theta(x, \bar{\mathbf{y}}^a) - r_\theta(x, \bar{\mathbf{y}}^b)) \sigma(r_\theta(x, \bar{\mathbf{y}}^b) - r_\theta(x, \bar{\mathbf{y}}^a))} \\ &= \frac{1}{2 \{1 + Z_\theta^+(x) Z_\theta^-(x)\}} \cdot \frac{1}{\text{Var}_{r_\theta}(\mathbb{1}\{\bar{\mathbf{y}}^a = \bar{\mathbf{y}}^b\} | x, \bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b)}. \end{aligned}$$

We then apply condition equation 28 and derive

$$\frac{\bar{\mu}(\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b | x)}{\pi_\theta(\bar{\mathbf{y}}^a | x) \pi_\theta(\bar{\mathbf{y}}^b | x)} \geq \frac{C^{-1}}{2 \{1 + Z_\theta^+(x) Z_\theta^-(x)\}} \cdot \frac{1}{\text{Var}_{r^*}(\mathbb{1}\{\bar{\mathbf{y}}^a = \bar{\mathbf{y}}^b\} | x, \bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b)}. \quad (29)$$

Next, substituting this result equation 29 into equation equation 11, alongside the weight function $w(x)$ from equation equation 7, we reform $\Sigma_\star$ as

$$\begin{aligned} \Sigma_\star &= \mathbb{E}_{x \sim \rho; \bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b \sim \pi_\theta(\cdot | x)} \left[\frac{\bar{\mu}(\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b | x)}{\pi_\theta(\bar{\mathbf{y}}^a | x) \pi_\theta(\bar{\mathbf{y}}^b | x)} \cdot w(x) \cdot \text{Var}(\mathbb{1}\{\bar{\mathbf{y}}^a = \bar{\mathbf{y}}^b\} | x, \bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \cdot \mathbf{g} \mathbf{g}^\top \right] \\ &\succeq \frac{1}{2C\bar{Z}_\theta} \mathbb{E}_{x \sim \rho; \bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b \sim \pi_\theta(\cdot | x)} [\mathbf{g} \mathbf{g}^\top]. \end{aligned} \quad (30)$$

The conditional expectation of $\mathbf{g} \mathbf{g}^\top$ simplifies as

$$\begin{aligned} &\mathbb{E}_{\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b \sim \pi_\theta(\cdot | x)} [\mathbf{g} \mathbf{g}^\top | x] \\ &= \mathbb{E}_{\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b \sim \pi_\theta(\cdot | x)} \left[\left\{ \nabla_\theta r^*(x, \bar{\mathbf{y}}^a) - \nabla_\theta r^*(x, \bar{\mathbf{y}}^b) \right\} \left\{ \nabla_\theta r^*(x, \bar{\mathbf{y}}^a) - \nabla_\theta r^*(x, \bar{\mathbf{y}}^b) \right\}^\top | x \right] \\ &= 2 \cdot \mathbb{E}_{\bar{\mathbf{y}} \sim \pi_\theta(\cdot | x)} \left[\nabla_\theta r^*(x, \bar{\mathbf{y}}) \nabla_\theta r^*(x, \bar{\mathbf{y}})^\top | x \right] - 2 \cdot \mathbb{E}_{\bar{\mathbf{y}} \sim \pi_\theta(\cdot | x)} [\nabla_\theta r^*(x, \bar{\mathbf{y}}) | x] \mathbb{E}_{\bar{\mathbf{y}} \sim \pi_\theta(\cdot | x)} [\nabla_\theta r^*(x, \bar{\mathbf{y}}) | x]^\top \\ &= 2 \cdot \text{Cov}_{\bar{\mathbf{y}} \sim \pi_\theta(\cdot | x)} [\nabla_\theta r^*(x, \bar{\mathbf{y}}) | x]. \end{aligned}$$

Substituting this result into equation equation 30, we arrive at the conclusion that

$$\Sigma_\star \succeq \frac{1}{C\bar{Z}_\phi} \mathbb{E}_{x \sim \rho} [\text{Cov}_{\bar{\mathbf{y}} \sim \pi^*(\cdot | x)} [\nabla_\theta r^*(x, \bar{\mathbf{y}}) | x]],$$

which matches the simplified form in equation equation 8 as stated in Theorem 4.2.

C.2.3 PROOF OF THEOREM 4.3

Gradient $\nabla_\theta J(\pi^*)$ and Hessian $\nabla_\theta^2 J(\pi^*)$: The equality $\nabla_\theta J(\pi^*) = 0$ follows directly from the gradient expression equation 39 for $\nabla_\theta J(\pi_\theta)$, evaluated at $\theta = \theta^*$ with $r_\theta = r^*$.

The proof of the Hessian result, $\nabla_\theta^2 J(\pi^*) = -(1/\beta) \cdot \Sigma_\star$, involves straightforward but technical differentiation of equation equation 39. For brevity, we defer this proof to Appendix D.3.1.

Asymptotic Distribution of Value Gap $J(\pi^*) - J(\hat{\pi})$: To understand the behavior of the value gap $J(\pi^*) - J(\hat{\pi})$, we start by applying a Taylor expansion of $J(\pi_\theta)$ around θ^* . This gives

$$J(\pi^*) - J(\hat{\pi}) = \nabla_\theta J(\pi^*)^\top (\theta^* - \hat{\theta}) - \frac{1}{2} (\theta^* - \hat{\theta})^\top \nabla_\theta^2 J(\pi^*) (\theta^* - \hat{\theta}) + o(\|\theta^* - \hat{\theta}\|_2^2).$$

By substituting $\nabla_\theta J(\pi^*) = \mathbf{0}$ (a direct result of the optimality of π^*), the linear term vanishes. Introducing the shorthand $\mathbf{H} := -\nabla_\theta^2 J(\pi^*) = (1/\beta) \cdot \Sigma_\star$, the expression simplifies to

$$J(\pi^*) - J(\hat{\pi}) = \frac{1}{2} (\hat{\theta} - \theta^*)^\top \mathbf{H} (\hat{\theta} - \theta^*) + o(\|\hat{\theta} - \theta^*\|_2^2). \quad (31)$$

When the sample size n is sufficiently large, $\hat{\theta}$ approaches θ^* , making the higher-order term negligible. Therefore, the value gap is dominated by the quadratic form.

From Theorem 4.2, we know the parameter estimate $\hat{\theta}$ satisfies

$$\sqrt{n} (\hat{\theta} - \theta^*) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \mathbf{\Omega}).$$

Substituting this result into the quadratic approximation of the value gap, we find that the scaled value gap has the asymptotic distribution

$$n \cdot \{J(\pi^*) - J(\hat{\pi})\} \xrightarrow{d} \frac{1}{2} \mathbf{z}^\top \boldsymbol{\Omega}^{\frac{1}{2}} \mathbf{H} \boldsymbol{\Omega}^{\frac{1}{2}} \mathbf{z} =: \mathbf{X} \quad \text{where } \mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}). \quad (32)$$

This approximation provides a clear intuition: the value gap is asymptotically driven by a weighted chi-squared-like term involving the covariance structure $\boldsymbol{\Omega}$ and the Hessian-like matrix $\mathbf{H}$.

To rigorously establish this result, we will apply Slutsky’s theorem. The full proof is presented in Appendix D.3.2.

Bounding the Chi-Square Distribution: To bound the random variable $\mathbf{X}$, we first leverage the estimate of the covariance matrix $\boldsymbol{\Omega}$ provided by Theorem 4.2:

$$\boldsymbol{\Omega} \preceq C \bar{Z}_\theta \|w\|_\infty \cdot \Sigma_\star^{-1},$$

where the constant C comes from condition equation 28. It follows that the matrix $\boldsymbol{\Omega}^{\frac{1}{2}} \mathbf{H} \boldsymbol{\Omega}^{\frac{1}{2}}$ appearing in equation equation 32 can be bounded as

$$\boldsymbol{\Omega}^{\frac{1}{2}} \mathbf{H} \boldsymbol{\Omega}^{\frac{1}{2}} \preceq C \|w\|_\infty \cdot \Sigma_\star^{-\frac{1}{2}} \mathbf{H} \Sigma_\star^{-\frac{1}{2}} = C \cdot \frac{\bar{Z}_\theta \|w\|_\infty}{\beta} \cdot \mathbf{I} = C \cdot \frac{1 + \|Z_\theta^+ Z_\theta^-\|_\infty}{\beta} \cdot \mathbf{I}.$$

Here the last equality uses the definition of the weight function w from equation equation 7. Substituting this bound into the quadratic form, we derive

$$\mathbf{X} = \frac{1}{2} \mathbf{z}^\top \boldsymbol{\Omega}^{\frac{1}{2}} \mathbf{H} \boldsymbol{\Omega}^{\frac{1}{2}} \mathbf{z} \leq C \cdot \frac{1 + \|Z_\theta^+ Z_\theta^-\|_\infty}{2\beta} \cdot \mathbf{z}^\top \mathbf{z},$$

where $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. Since $\mathbf{z}^\top \mathbf{z}$ follows a chi-square distribution with d degrees of freedom, $\mathbf{X}$ is stochastically dominated by a rescaled chi-square random variable

$$C \cdot \frac{1 + \|Z_\theta^+ Z_\theta^-\|_\infty}{2\beta} \cdot \chi_d^2.$$

Equivalently, we can express this dominance as

$$\limsup_{n \rightarrow \infty} \mathbb{P} \left\{ n \{J(\pi^*) - J(\hat{\pi})\} > C \cdot \frac{1 + \|Z_\theta^+ Z_\theta^-\|_\infty}{2\beta} \cdot t \right\} \leq \mathbb{P} \{ \chi_d^2 > t \} \quad \text{for any } t > 0. \quad (33)$$

This inequality, given in equation equation 33, corresponds to the first bound in equation equation 10.

The second inequality in equation equation 10 provides a precise tail bound for χ_d^2 . As its proof involves more technical details, we defer it to Appendix D.3.3.

D PROOF OF AUXILIARY RESULTS

This section provides proofs of auxiliary results supporting the main theorems and lemmas. In Appendix D.1, we present the auxiliary results required for Theorem 4.1. Appendix D.2 details the proofs of supporting results for Theorem 4.2. Finally, in Appendix D.3, we establish the auxiliary results necessary for Theorem 4.3.

D.1 PROOF OF AUXILIARY RESULTS FOR THEOREM 4.1

In this section, we provide the proofs of several auxiliary results that support the proof of Theorem 4.1. Specifically, Appendix D.1.1 presents the forms of the gradients of the policy π_θ and the reward r_θ , which serve as fundamental building blocks for deriving the lemmas. Appendix D.1.2 analyzes the gradient of the return function $J(\pi_\theta)$, as defined in equation equation 6. Appendix D.1.3 focuses on deriving expressions for the gradient of the negative log-likelihood function $\mathcal{L}(\theta)$.

D.1.1 GRADIENTS OF POLICY π_θ AND REWARD r_θ

In this part, we introduce results for the gradients of policy π_θ and reward r_θ with respect to parameter θ , which lay the foundation of our calculations.

Lemma D.1 (Gradients of policy π_θ and reward function r_θ). *The gradients of the policy π_θ and the reward function r_θ can be expressed in terms of each other as follows*

$$\nabla_\theta \pi_\theta(d\vec{y} | x) = \pi_\theta(d\vec{y} | x) \cdot \frac{1}{\beta} \left\{ \nabla_\theta r_\theta(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_\theta(\cdot | x)} [\nabla_\theta r_\theta(x, \vec{y}')] \right\}, \quad (34a)$$

$$\nabla_\theta r_\theta(x, \vec{y}) = \beta \cdot \frac{\nabla_\theta \pi_\theta(\vec{y} | x)}{\pi_\theta(\vec{y} | x)}. \quad (34b)$$

We now proceed to prove Lemma D.1.

To begin, recall our definition of the reward function r_θ as given in equation 5. It directly follows that

$$\nabla_\theta r_\theta(x, \vec{y}) = \beta \cdot \frac{\nabla_\theta \pi_\theta(\vec{y} | x)}{\pi_\theta(\vec{y} | x)}.$$

This result confirms equation 34b as stated in Lemma D.1.

Next, we express the policy $\pi_\theta(d\vec{y} | x)$ in terms of the reward function $r_\theta(x, \vec{y})$. By reformulating equation 5, we obtain

$$\pi_\theta(d\vec{y} | x) = \frac{1}{Z_\theta(x)} \pi_{\text{ref}}(d\vec{y} | x) \exp \left\{ \frac{1}{\beta} r_\theta(x, \vec{y}) \right\}, \quad (35a)$$

where $Z_\theta(x)$ is the partition function defined as

$$Z_\theta(x) = \int_{\mathcal{Y}} \pi_{\text{ref}}(d\vec{y} | x) \exp \left\{ \frac{1}{\beta} r_\theta(x, \vec{y}) \right\}. \quad (35b)$$

We then compute the gradient of $\pi_\theta(d\vec{y} | x)$ with respect to θ . Applying the chain rule, we get

$$\begin{aligned} \nabla_\theta \pi_\theta(d\vec{y} | x) &= \frac{1}{Z_\theta(x)} \pi_{\text{ref}}(d\vec{y} | x) \exp \left\{ \frac{1}{\beta} r_\theta(x, \vec{y}) \right\} \cdot \frac{1}{\beta} \nabla_\theta r_\theta(x, \vec{y}) \\ &\quad - \frac{1}{Z_\theta^2(x)} \pi_{\text{ref}}(d\vec{y} | x) \exp \left\{ \frac{1}{\beta} r_\theta(x, \vec{y}) \right\} \cdot \nabla_\theta Z_\theta(x). \end{aligned} \quad (36)$$

We need the gradient of the partition function $Z_\theta(x)$:

$$\begin{aligned} \nabla_\theta Z_\theta(x) &= \int_{\mathcal{Y}} \pi_{\text{ref}}(d\vec{y} | x) \exp \left\{ \frac{1}{\beta} r_\theta(x, \vec{y}) \right\} \cdot \frac{1}{\beta} \nabla_\theta r_\theta(x, \vec{y}) \\ &= Z_\theta(x) \cdot \int_{\mathcal{Y}} \pi_\theta(d\vec{y} | x) \cdot \frac{1}{\beta} \nabla_\theta r_\theta(x, \vec{y}) \\ &= Z_\theta(x) \cdot \frac{1}{\beta} \mathbb{E}_{\vec{y} \sim \pi_\theta(\cdot | x)} [\nabla_\theta r_\theta(x, \vec{y})]. \end{aligned} \quad (37)$$

Substituting equation 37 back into equation 36, we simplify the expression for the gradient of $\pi_\theta(d\vec{y} | x)$:

$$\nabla_\theta \pi_\theta(d\vec{y} | x) = \frac{1}{Z_\theta(x)} \pi_{\text{ref}}(d\vec{y} | x) \exp \left\{ \frac{1}{\beta} r_\theta(x, \vec{y}) \right\} \cdot \frac{1}{\beta} \left\{ \nabla_\theta r_\theta(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_\theta(\cdot | x)} [\nabla_\theta r_\theta(x, \vec{y}')] \right\}.$$

This matches equation 34a from Lemma D.1, thereby completing the proof.

D.1.2 PROOF OF LEMMA C.2

Equality equation 13 in Lemma C.2 can be derived as a consequence of a more detailed result. We state it in Lemma D.2.

Lemma D.2. *For a policy π_θ , the gradients with respect to the parameter θ of its expected return $\mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_\theta(\cdot | x)} [r^*(x, \vec{y})]$ and its KL divergence from a reference policy $D_{\text{KL}}(\pi_\theta \| \pi_{\text{ref}})$ are given*

by

$$\nabla_{\theta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} [r^*(x, \vec{y})] = \frac{1}{\beta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} \left[r^*(x, \vec{y}) \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_{\theta}(\cdot|x)} [\nabla_{\theta} r_{\theta}(x, \vec{y}')] \right\} \right], \quad (38a)$$

$$\nabla_{\theta} D_{KL}(\pi_{\theta} \parallel \pi_{\text{ref}}) = \frac{1}{\beta^2} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} \left[r_{\theta}(x, \vec{y}) \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_{\theta}(\cdot|x)} [\nabla_{\theta} r_{\theta}(x, \vec{y}')] \right\} \right]. \quad (38b)$$

Recall that the scalar value $J(\pi_{\theta})$ of the policy is defined as

$$J(\pi_{\theta}) = \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} [r^*(x, \vec{y})] - \beta D_{KL}(\pi_{\theta} \parallel \pi_{\text{ref}}).$$

Using Lemma D.2, we derive the gradient of $J(\pi_{\theta})$ as

$$\begin{aligned} \nabla_{\theta} J(\pi_{\theta}) &= \nabla_{\theta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} [r^*(x, \vec{y})] - \beta \nabla_{\theta} D_{KL}(\pi_{\theta} \parallel \pi_{\text{ref}}) \\ &= \frac{1}{\beta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} \left[\left\{ r^*(x, \vec{y}) - r_{\theta}(x, \vec{y}) \right\} \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_{\theta}(\cdot|x)} [\nabla_{\theta} r_{\theta}(x, \vec{y}')] \right\} \right]. \end{aligned} \quad (39)$$

We rewrite the expression in equation equation 39 in two equivalent forms by exchanging the roles of $\vec{y}^a$ and $\vec{y}^b$:

$$\nabla_{\theta} J(\pi_{\theta}) = \frac{1}{\beta} \mathbb{E}_{x \sim \rho, \vec{y}^a \sim \pi_{\theta}(\cdot|x)} \left[\left\{ r^*(x, \vec{y}^a) - r_{\theta}(x, \vec{y}^a) \right\} \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}^a) - \mathbb{E}_{\vec{y}^b \sim \pi_{\theta}(\cdot|x)} [\nabla_{\theta} r_{\theta}(x, \vec{y}^b)] \right\} \right], \quad (40a)$$

$$\nabla_{\theta} J(\pi_{\theta}) = \frac{1}{\beta} \mathbb{E}_{x \sim \rho, \vec{y}^b \sim \pi_{\theta}(\cdot|x)} \left[\left\{ r^*(x, \vec{y}^b) - r_{\theta}(x, \vec{y}^b) \right\} \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}^b) - \mathbb{E}_{\vec{y}^a \sim \pi_{\theta}(\cdot|x)} [\nabla_{\theta} r_{\theta}(x, \vec{y}^a)] \right\} \right]. \quad (40b)$$

By taking the average of the two equivalent formulations above, we obtain equality equation 13 and complete the proof of Lemma C.2.

We now proceed to prove Lemma D.2, tackling equalities equation 38a and equation 38b one by one.

Proof of Equality equation 38a from Lemma D.2: We begin by expressing the expected return as

$$\mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} [r^*(x, \vec{y})] = \mathbb{E}_{x \sim \rho} \left[\int_{\mathcal{Y}} r^*(x, \vec{y}) \pi_{\theta}(d\vec{y} | x) \right].$$

Taking the gradient of both sides with respect to θ , we have

$$\nabla_{\theta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} [r^*(x, \vec{y})] = \mathbb{E}_{x \sim \rho} \left[\int_{\mathcal{Y}} r^*(x, \vec{y}) \nabla_{\theta} \pi_{\theta}(d\vec{y} | x) \right]. \quad (41)$$

Using the expression for the policy gradient $\nabla_{\theta} \pi_{\theta}$ provided in Lemma D.1, the right-hand side of equation 41 simplifies to

$$\begin{aligned} \text{RHS of equation 41} &= \mathbb{E}_{x \sim \rho} \left[\int_{\mathcal{Y}} r^*(x, \vec{y}) \pi_{\theta}(d\vec{y} | x) \cdot \frac{1}{\beta} \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_{\theta}(\cdot|x)} [\nabla_{\theta} r_{\theta}(x, \vec{y}')] \right\} \right] \\ &= \frac{1}{\beta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot|x)} \left[r^*(x, \vec{y}) \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_{\theta}(\cdot|x)} [\nabla_{\theta} r_{\theta}(x, \vec{y}')] \right\} \right]. \end{aligned}$$

This completes the verification of equation equation 38a from Lemma C.2.

Proof of Equality equation 38b from Lemma D.2: Recall the definition of the KL divergence

$$D_{KL}(\pi_{\theta} \parallel \pi_{\text{ref}}) = \mathbb{E}_{x \sim \rho} \left[\int_{\mathcal{Y}} \pi_{\theta}(d\vec{y} | x) \log \left(\frac{\pi_{\theta}(\vec{y} | x)}{\pi_{\text{ref}}(\vec{y} | x)} \right) \right].$$

Applying the chain rule, we obtain

$$\nabla_{\theta} D_{\text{KL}}(\pi_{\theta} \parallel \pi_{\text{ref}}) = \mathbb{E}_{x \sim \rho} \left[\int_{\mathcal{Y}} \nabla_{\theta} \pi_{\theta}(d\vec{y} \mid x) \log \left(\frac{\pi_{\theta}(\vec{y} \mid x)}{\pi_{\text{ref}}(\vec{y} \mid x)} \right) \right] + \mathbb{E}_{x \sim \rho} \left[\int_{\mathcal{Y}} \nabla_{\theta} \pi_{\theta}(d\vec{y} \mid x) \right]. \quad (42)$$

Since the policy integrates to 1, i.e., $\int_{\mathcal{Y}} \pi_{\theta}(d\vec{y} \mid x) = 1$, it always holds that

$$\int_{\mathcal{Y}} \nabla_{\theta} \pi_{\theta}(d\vec{y} \mid x) = \nabla_{\theta} \int_{\mathcal{Y}} \pi_{\theta}(d\vec{y} \mid x) = 0, \quad (43)$$

i.e., the second term on the right-hand side of equation 42 is zero. Using the expression equation 35a, we take the logarithm

$$\log \left(\frac{\pi_{\theta}(\vec{y} \mid x)}{\pi_{\text{ref}}(\vec{y} \mid x)} \right) = \frac{1}{\beta} r_{\theta}(x, \vec{y}) - \log Z_{\theta}(x). \quad (44)$$

Combining equations equation 43 and equation 44, we get

$$\begin{aligned} & \int_{\mathcal{Y}} \nabla_{\theta} \pi_{\theta}(d\vec{y} \mid x) \log \left(\frac{\pi_{\theta}(\vec{y} \mid x)}{\pi_{\text{ref}}(\vec{y} \mid x)} \right) \\ &= \frac{1}{\beta} \int_{\mathcal{Y}} r_{\theta}(x, \vec{y}) \nabla_{\theta} \pi_{\theta}(d\vec{y} \mid x) - \log Z_{\theta}(x) \int_{\mathcal{Y}} \nabla_{\theta} \pi_{\theta}(d\vec{y} \mid x) \\ &= \frac{1}{\beta} \int_{\mathcal{Y}} r_{\theta}(x, \vec{y}) \nabla_{\theta} \pi_{\theta}(d\vec{y} \mid x). \end{aligned} \quad (45)$$

Now, similar to the proof of equation equation 38a, we derive

$$\begin{aligned} \text{RHS of equation 42} &= \frac{1}{\beta} \mathbb{E}_{x \sim \rho} \left[\int_{\mathcal{Y}} r_{\theta}(x, \vec{y}) \nabla_{\theta} \pi_{\theta}(d\vec{y} \mid x) \right] \\ &= \frac{1}{\beta^2} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_{\theta}(\cdot \mid x)} \left[r_{\theta}(x, \vec{y}) \left\{ \nabla_{\theta} r_{\theta}(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_{\theta}(\cdot \mid x)} [\nabla_{\theta} r_{\theta}(x, \vec{y}')] \right\} \right], \end{aligned}$$

which verifies equality equation 38b from Lemma D.2.

D.1.3 PROOF OF LEMMA C.3

In this section, we prove a full version of Lemma C.3 as stated in Lemma D.3 below. Equation equation 14a from Lemma C.3 follows directly as a straightforward corollary.

In Lemma D.3, we consider a general class of distributions parameterized by θ that models the binary preference $\mathbb{P}_{\theta}(\vec{y}^a \succ \vec{y}^b \mid x)$. The negative log-likelihood function is defined as

$$\mathcal{L}(\theta) = -\mathbb{E}_{x \sim \rho; (\vec{y}^a, \vec{y}^b) \sim \bar{\mu}(\cdot \mid x)} \left[w(x) \cdot \log \mathbb{P}_{\theta}(\vec{y}^a \succ \vec{y}^b \mid x) \right].$$

The Bradley-Terry (BT) model described in equation equation 1 and the corresponding loss function $\mathcal{L}(\theta)$ in equation equation 47 represent a special case of this general framework.

Lemma D.3 (Gradient of the loss function $\mathcal{L}(\theta)$, full version). *For a general distribution class $\{\mathbb{P}_{\theta}\}$, the gradient of $\mathcal{L}(\theta)$ with respect to θ is given by*

$$\begin{aligned} \nabla_{\theta} \mathcal{L}(\theta) &= -\mathbb{E}_{x \sim \rho; (\vec{y}^a, \vec{y}^b) \sim \bar{\mu}(\cdot \mid x)} \left[w(x) \cdot \left\{ \mathbb{P}(\vec{y}^a \succ \vec{y}^b \mid x) - \mathbb{P}_{\theta}(\vec{y}^a \succ \vec{y}^b \mid x) \right\} \right. \\ &\quad \left. \cdot \frac{\nabla_{\theta} \mathbb{P}_{\theta}(\vec{y}^a \succ \vec{y}^b \mid x)}{\mathbb{P}_{\theta}(\vec{y}^a \succ \vec{y}^b \mid x) \mathbb{P}_{\theta}(\vec{y}^b \succ \vec{y}^a \mid x)} \right], \end{aligned} \quad (46a)$$

where $\bar{\mu}$ is the average distribution defined in equation equation 14b. Specifically, for the Bradley-Terry (BT) model where

$$\mathbb{P}_{\theta}(\vec{y}^a \succ \vec{y}^b \mid x) = \sigma(r_{\theta}(x, \vec{y}^a) - r_{\theta}(x, \vec{y}^b)) = \left\{ 1 + \left(\frac{(\pi_{\theta}/\pi_{\text{ref}})(\vec{y}^b \mid x)}{(\pi_{\theta}/\pi_{\text{ref}})(\vec{y}^a \mid x)} \right)^{\beta} \right\}^{-1},$$

the gradient of $\mathcal{L}(\theta)$ becomes

$$\begin{aligned} \nabla_{\theta} \mathcal{L}(\theta) = & -\mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot | x)} \left[w(x) \cdot \left\{ \sigma(r^*(x, \bar{\mathbf{y}}^a) - r^*(x, \bar{\mathbf{y}}^b)) - \sigma(r_{\theta}(x, \bar{\mathbf{y}}^a) - r_{\theta}(x, \bar{\mathbf{y}}^b)) \right\} \right. \\ & \left. \cdot \left\{ \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^a) - \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^b) \right\} \right]. \quad (46b) \end{aligned}$$

For notational simplicity, we focus on the proof for the case where the weight function $w(x) = 1$. The results for a general weight function $w(x) > 0$ can be derived in a similar manner.

Recall that the negative log-likelihood function $\mathcal{L}(\theta)$ is defined as

$$\mathcal{L}(\theta) = \mathbb{E} \left[-\log \mathbb{P}_{\theta}(\bar{\mathbf{y}}^w \succ \bar{\mathbf{y}}^{\ell} | x) \right].$$

Based on the data generation mechanism, we can expand the expectation in $\mathcal{L}(\theta)$ as

$$\begin{aligned} \mathcal{L}(\theta) = & \mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \mu(\cdot | x)} \left[\mathbb{P}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x) \cdot \left\{ -\log \mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x) \right\} \right. \\ & \left. + \mathbb{P}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x) \cdot \left\{ -\log \mathbb{P}_{\theta}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x) \right\} \right]. \quad (47) \end{aligned}$$

Notice that we can exchange the roles of $\bar{\mathbf{y}}^a$ and $\bar{\mathbf{y}}^b$ in the expectation above. This means that we can equivalently express the expectation using the pair $(\bar{\mathbf{y}}^b, \bar{\mathbf{y}}^a) \sim \mu(\cdot | x)$. This symmetry allows us to replace μ in equation 47 with the average distribution $\bar{\mu}$ as defined in equation 14b.

Next, we take the gradient of the loss function $\mathcal{L}(\theta)$ with respect to the parameter θ and obtain

$$\begin{aligned} \nabla_{\theta} \mathcal{L}(\theta) = & \mathbb{E}_{x \sim \rho, (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot | x)} \left[\frac{\mathbb{P}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)}{\mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)} \cdot \left\{ -\nabla_{\theta} \mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x) \right\} \right. \\ & \left. + \frac{\mathbb{P}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x)}{\mathbb{P}_{\theta}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x)} \cdot \left\{ -\nabla_{\theta} \mathbb{P}_{\theta}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x) \right\} \right]. \end{aligned}$$

Note that $\mathbb{P}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x) = 1 - \mathbb{P}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)$ and $\mathbb{P}_{\theta}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x) = 1 - \mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)$. Using this, we can rewrite the gradient as

$$\nabla_{\theta} \mathcal{L}(\theta) = \mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot | x)} \left[\left\{ \frac{1 - \mathbb{P}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)}{1 - \mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)} - \frac{\mathbb{P}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)}{\mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)} \right\} \cdot \nabla_{\theta} \mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x) \right].$$

We simplify the expression further to obtain

$$\nabla_{\theta} \mathcal{L}(\theta) = \mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot | x)} \left[\left\{ \mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x) - \mathbb{P}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x) \right\} \cdot \frac{\nabla_{\theta} \mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)}{\mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x) \mathbb{P}_{\theta}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x)} \right].$$

This establishes equation 46a from Lemma C.3.

As for the Bradley-Terry (BT) model, we use the equality

$$\sigma'(z) = \frac{1}{(1 + \exp(-z))(1 + \exp(z))} = \sigma(z) \sigma(-z) \quad \text{for any } z \in \mathbb{R}$$

to derive the following expression

$$\frac{\nabla_{\theta} \mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x)}{\mathbb{P}_{\theta}(\bar{\mathbf{y}}^a \succ \bar{\mathbf{y}}^b | x) \mathbb{P}_{\theta}(\bar{\mathbf{y}}^b \succ \bar{\mathbf{y}}^a | x)} = \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^a) - \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^b). \quad (48)$$

By substituting this gradient expression from equation 48 into equation 46a, we directly obtain equation 46b, thereby completing the proof of Lemma C.3.

D.2 PROOF OF AUXILIARY RESULTS FOR THEOREM 4.2

In this section, we present the detailed proofs of the supporting lemmas used in the proof of Theorem 4.2. We begin in Appendix D.2.1 by establishing condition equation 25, which is crucial for the valid application of the master theorem for Z -estimators. Following this, in Appendix D.2.2,

we compute the Hessian matrix $\nabla_{\theta}^2 \mathcal{L}(\theta^*)$ explicitly. Finally, in Appendix D.2.3, we derive the asymptotic distribution of the gradient $\nabla_{\theta} \hat{\mathcal{L}}(\theta^*)$.

D.2.1 PROOF OF CONDITION EQUATION 25

We begin by rewriting the left-hand side of equation equation 25 as follows:

$$\begin{aligned} \Delta &:= \sqrt{n} \{ \nabla_{\theta} \hat{\mathcal{L}}(\hat{\theta}) - \nabla_{\theta} \hat{\mathcal{L}}(\hat{\theta}) \} - \sqrt{n} \{ \nabla_{\theta} \hat{\mathcal{L}}(\theta^*) - \nabla_{\theta} \mathcal{L}(\theta^*) \} \\ &= \sqrt{n} \{ \nabla_{\theta} \hat{\mathcal{L}}(\hat{\theta}) - \nabla_{\theta} \hat{\mathcal{L}}(\theta^*) \} - \sqrt{n} \{ \nabla_{\theta} \mathcal{L}(\hat{\theta}) - \nabla_{\theta} \mathcal{L}(\theta^*) \}. \end{aligned} \quad (49)$$

We then leverage the smoothness properties of the function r_{θ} , which guarantee the following approximations:

$$\nabla_{\theta} \hat{\mathcal{L}}(\hat{\theta}) - \nabla_{\theta} \hat{\mathcal{L}}(\theta^*) = \nabla_{\theta}^2 \hat{\mathcal{L}}(\theta^*) (\hat{\theta} - \theta^*) + o_p(\|\hat{\theta} - \theta^*\|_2), \quad (50a)$$

$$\nabla_{\theta} \mathcal{L}(\hat{\theta}) - \nabla_{\theta} \mathcal{L}(\theta^*) = \nabla_{\theta}^2 \mathcal{L}(\theta^*) (\hat{\theta} - \theta^*) + o_p(\|\hat{\theta} - \theta^*\|_2). \quad (50b)$$

Assuming these equalities equation 50a and equation 50b hold, we substitute them into equation equation 49, leading to

$$\begin{aligned} \Delta &= \sqrt{n} \{ \nabla_{\theta}^2 \hat{\mathcal{L}}(\theta^*) (\hat{\theta} - \theta^*) + o_p(\|\hat{\theta} - \theta^*\|_2) \} - \sqrt{n} \{ \nabla_{\theta}^2 \mathcal{L}(\theta^*) (\hat{\theta} - \theta^*) + o_p(\|\hat{\theta} - \theta^*\|_2) \} \\ &= \sqrt{n} \{ \nabla_{\theta}^2 \hat{\mathcal{L}}(\theta^*) - \nabla_{\theta}^2 \mathcal{L}(\theta^*) \} (\hat{\theta} - \theta^*) + o_p(1 + \sqrt{n} \|\hat{\theta} - \theta^*\|_2). \end{aligned} \quad (51)$$

Using the law of large numbers, we know that $\nabla_{\theta}^2 \hat{\mathcal{L}}(\theta^*) \xrightarrow{P} \nabla_{\theta}^2 \mathcal{L}(\theta^*)$, which implies

$$\sqrt{n} \{ \nabla_{\theta}^2 \hat{\mathcal{L}}(\theta^*) - \nabla_{\theta}^2 \mathcal{L}(\theta^*) \} (\hat{\theta} - \theta^*) = o_p(\sqrt{n} \|\hat{\theta} - \theta^*\|_2).$$

Therefore, we conclude that

$$\Delta = o_p(1 + \sqrt{n} \|\hat{\theta} - \theta^*\|_2)$$

as claimed in equation equation 25.

The only remaining task is to establish the validity of equalities equation 50a and equation 50b.

Proof of Equalities equation 50a and equation 50b: We express the loss function $\hat{\mathcal{L}}(\theta)$ in the form

$$\hat{\mathcal{L}}(\theta) := \frac{1}{n} \sum_{i=1}^n w(x_i) \cdot \ell_{\theta}(x_i, \vec{y}_i^w, \vec{y}_i^{\ell}),$$

where the function ℓ_{θ} is defined as

$$\ell_{\theta}(x, \vec{y}_1, \vec{y}_2) = -\log \sigma(r_{\theta}(x, \vec{y}_1) - r_{\theta}(x, \vec{y}_2)).$$

We then calculate the gradient $\nabla_{\theta} \ell_{\theta}$ and $\nabla_{\theta}^2 \ell_{\theta}$ as follows:

$$\begin{aligned} \nabla_{\theta} \ell_{\theta}(x, \vec{y}_1, \vec{y}_2) &= \sigma(r_{\theta}(x, \vec{y}_2) - r_{\theta}(x, \vec{y}_1)) \cdot \{ \nabla_{\theta} r_{\theta}(x, \vec{y}_2) - \nabla_{\theta} r_{\theta}(x, \vec{y}_1) \} \quad \text{and} \\ \nabla_{\theta}^2 \ell_{\theta}(x, \vec{y}_1, \vec{y}_2) &= \sigma'(r_{\theta}(x, \vec{y}_2) - r_{\theta}(x, \vec{y}_1)) \cdot \{ \nabla_{\theta} r_{\theta}(x, \vec{y}_2) - \nabla_{\theta} r_{\theta}(x, \vec{y}_1) \} \{ \nabla_{\theta} r_{\theta}(x, \vec{y}_2) - \nabla_{\theta} r_{\theta}(x, \vec{y}_1) \}^{\top} \\ &\quad + \sigma(r_{\theta}(x, \vec{y}_2) - r_{\theta}(x, \vec{y}_1)) \cdot \{ \nabla_{\theta}^2 r_{\theta}(x, \vec{y}_2) - \nabla_{\theta}^2 r_{\theta}(x, \vec{y}_1) \}. \end{aligned}$$

When the reward function $r_{\theta}(x, \vec{y})$, along with its gradient $\nabla_{\theta} r_{\theta}(x, \vec{y})$ and Hessian $\nabla_{\theta}^2 r_{\theta}(x, \vec{y})$, is uniformly bounded and Lipschitz continuous with respect to θ for all $(x, \vec{y}) \in \mathcal{X} \times \mathcal{Y}$, it guarantees that the Hessian of the loss function, $\nabla_{\theta}^2 \ell_{\theta}$, is also Lipschitz continuous. This holds with some constant $L > 0$ across all $(x, \vec{y}) \in \mathcal{X} \times \mathcal{Y}$, as demonstrated below:

$$\| \nabla_{\theta}^2 \ell_{\theta}(x, \vec{y}_1, \vec{y}_2) - \nabla_{\theta}^2 \ell_{\theta^*}(x, \vec{y}_1, \vec{y}_2) \|_2 \leq L \cdot \|\theta - \theta^*\|_2.$$

From this Lipschitz property, we deduce

$$\| \nabla_{\theta} \ell_{\theta}(x, \vec{y}_1, \vec{y}_2) - \nabla_{\theta} \ell_{\theta^*}(x, \vec{y}_1, \vec{y}_2) - \nabla_{\theta}^2 \ell_{\theta^*}(x, \vec{y}_1, \vec{y}_2) (\theta - \theta^*) \|_2 \leq \frac{L}{2} \cdot \|\theta - \theta^*\|_2^2$$

and further derive

$$\begin{aligned}\|\nabla_{\theta} \widehat{\mathcal{L}}(\theta) - \nabla_{\theta} \widehat{\mathcal{L}}(\theta^*) - \nabla_{\theta}^2 \widehat{\mathcal{L}}(\theta^*) (\theta - \theta^*)\|_2 &\leq \frac{L \|w\|_{\infty}}{2} \cdot \|\theta - \theta^*\|_2^2, \\ \|\nabla_{\theta} \mathcal{L}(\theta) - \nabla_{\theta} \mathcal{L}(\theta^*) - \nabla_{\theta}^2 \mathcal{L}(\theta^*) (\theta - \theta^*)\|_2 &\leq \frac{L \|w\|_{\infty}}{2} \cdot \|\theta - \theta^*\|_2^2.\end{aligned}$$

Finally, under the condition that $\widehat{\theta} \xrightarrow{P} \theta^*$, these results simplify to the expressions given in equations equation 50a and equation 50b, as previously claimed.

D.2.2 PROOF OF LEMMA C.5, EXPLICIT FORM OF HESSIAN $\nabla_{\theta}^2 \mathcal{L}(\theta^*)$

From equation equation 14a in Lemma C.3, we recall the explicit formula for the gradient $\nabla_{\theta} \mathcal{L}(\theta)$. Taking the derivative of both sides of equation equation 14a, we obtain

$$\begin{aligned}\nabla_{\theta}^2 \mathcal{L}(\theta) &= \mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot|x)} \left[w(x) \cdot \sigma'(r_{\theta}(x, \bar{\mathbf{y}}^a) - r_{\theta}(x, \bar{\mathbf{y}}^b)) \right. \\ &\quad \cdot \left\{ \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^a) - \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^b) \right\} \left\{ \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^a) - \nabla_{\theta} r_{\theta}(x, \bar{\mathbf{y}}^b) \right\}^{\top} \Big] \\ &\quad - \mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot|x)} \left[w(x) \cdot \left\{ \sigma(r^*(x, \bar{\mathbf{y}}^a) - r^*(x, \bar{\mathbf{y}}^b)) - \sigma(r_{\theta}(x, \bar{\mathbf{y}}^a) - r_{\theta}(x, \bar{\mathbf{y}}^b)) \right\} \right. \\ &\quad \cdot \left. \left\{ \nabla_{\theta}^2 r_{\theta}(x, \bar{\mathbf{y}}^a) - \nabla_{\theta}^2 r_{\theta}(x, \bar{\mathbf{y}}^b) \right\} \right].\end{aligned}\tag{52}$$

When we set $\theta = \theta^*$, it follows that $r_{\theta} = r^*$. This simplification eliminates the second term in expression equation 52, reducing the Hessian matrix to

$$\begin{aligned}\nabla_{\theta}^2 \mathcal{L}(\theta^*) &= \mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot|x)} \left[w(x) \cdot \sigma'(r^*(x, \bar{\mathbf{y}}^a) - r^*(x, \bar{\mathbf{y}}^b)) \right. \\ &\quad \cdot \left\{ \nabla_{\theta} r^*(x, \bar{\mathbf{y}}^a) - \nabla_{\theta} r^*(x, \bar{\mathbf{y}}^b) \right\} \left\{ \nabla_{\theta} r^*(x, \bar{\mathbf{y}}^a) - \nabla_{\theta} r^*(x, \bar{\mathbf{y}}^b) \right\}^{\top} \Big].\end{aligned}$$

Substituting the derivative σ' with its explicit form, $\sigma'(z) = \sigma(z) \sigma(-z)$ for any $z \in \mathbb{R}$, we refine the expression to

$$\nabla_{\theta}^2 \mathcal{L}(\theta^*) = \Sigma_{\star},$$

where the covariance matrix $\Sigma_{\star}$ is defined in equation equation 11. This completes the proof of expression equation 26 from Lemma C.5.

D.2.3 PROOF OF LEMMA C.6, ASYMPTOTIC DISTRIBUTION OF GRAIDENT $\nabla_{\theta} \widehat{\mathcal{L}}(\theta^*)$

In this section, we analyze the asymptotic distribution of the gradient $\nabla_{\theta} \widehat{\mathcal{L}}(\theta)$ at $\theta = \theta^*$, where the loss function $\widehat{\mathcal{L}}(\theta)$ is defined as

$$\widehat{\mathcal{L}}(\theta) = -\frac{1}{n} \sum_{i=1}^n w(x_i) \cdot \log \sigma(r_{\theta}(x_i, \bar{\mathbf{y}}_i^w) - r_{\theta}(x_i, \bar{\mathbf{y}}_i^{\ell})).$$

Using the definition of the sigmoid function σ , we calculate that

$$(\log \sigma(z))' = \sigma'(z)/\sigma(z) = \sigma(z) \sigma(-z)/\sigma(z) = \sigma(-z) \quad \text{for any real number } z \in \mathbb{R}.$$

This allows us to reformulate $\nabla_{\theta} \widehat{\mathcal{L}}(\theta)$ as the average of n i.i.d. vectors $\{\mathbf{u}_i\}_{i=1}^n$:

$$\nabla_{\theta} \widehat{\mathcal{L}}(\theta) = \frac{1}{n} \sum_{i=1}^n \mathbf{u}_i.\tag{53}$$

Here each vector $\mathbf{u}_i \in \mathbb{R}^d$ is defined as

$$\mathbf{u}_i := w(x_i) \cdot \sigma(r_{\theta}(x_i, \bar{\mathbf{y}}_i^{\ell}) - r_{\theta}(x_i, \bar{\mathbf{y}}_i^w)) \cdot \left\{ \nabla_{\theta} r_{\theta}(x_i, \bar{\mathbf{y}}_i^{\ell}) - \nabla_{\theta} r_{\theta}(x_i, \bar{\mathbf{y}}_i^w) \right\}.$$

At $\theta = \theta^*$, we denote $\mathbf{u}_i$ as $\mathbf{u}_i^*$ and $\mathbf{g}_i$ as $\mathbf{g}_i^*$. Notably, vector $\mathbf{u}_i$ can be rewritten as

$$\mathbf{u}_i = w(x_i) \cdot \left\{ \sigma(r_{\theta}(x_i, \bar{\mathbf{y}}_i^a) - r_{\theta}(x_i, \bar{\mathbf{y}}_i^b)) - \mathbb{1}\{\bar{\mathbf{y}}_i^a = \bar{\mathbf{y}}_i^w, \bar{\mathbf{y}}_i^b = \bar{\mathbf{y}}_i^{\ell}\} \right\} \cdot \mathbf{g}_i,\tag{54}$$

where $\mathbf{g}_i$ is given by

$$\mathbf{g}_i := \nabla_{\theta} r_{\theta}(x_i, \bar{\mathbf{y}}_i^a) - \nabla_{\theta} r_{\theta}(x_i, \bar{\mathbf{y}}_i^b).$$

From the structure of the BT model, it holds that

$$\mathbb{E}[\mathbb{1}\{\bar{\mathbf{y}}_i^a = \bar{\mathbf{y}}_i^w, \bar{\mathbf{y}}_i^b = \bar{\mathbf{y}}_i^{\ell}\} \mid x_i] = \sigma(r^*(x_i, \bar{\mathbf{y}}_i^a) - r^*(x_i, \bar{\mathbf{y}}_i^b)),$$

which implies $\mathbb{E}[\mathbf{u}_i^*] = \mathbf{0}$.

To analyze the asymptotic distribution of $\nabla_{\theta} \hat{\mathcal{L}}(\theta^*)$, we apply the central limit theorem (CLT) to its empirical form given in equation equation 53. By the CLT, we have

$$\sqrt{n}(\nabla_{\theta} \hat{\mathcal{L}}(\theta^*) - \nabla_{\theta} \mathcal{L}(\theta^*)) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \tilde{\Omega}), \quad n \rightarrow \infty, \quad (55)$$

where the covariance matrix $\tilde{\Omega} \in \mathbb{R}^{d \times d}$ is given by

$$\tilde{\Omega} := \text{Cov}(\mathbf{u}_1^*) = \mathbb{E}[\mathbf{u}_1^* (\mathbf{u}_1^*)^{\top}].$$

Here we have used the property $\mathbb{E}[\mathbf{u}_i^*] = \mathbf{0}$ in the second equality.

We now compute the explicit form of the covariance matrix $\tilde{\Omega}$. Using the definition of $\mathbf{u}_i$ from expression equation 54, we find that

$$\begin{aligned} \tilde{\Omega} &= \mathbb{E}[\mathbf{u}_1^* (\mathbf{u}_1^*)^{\top}] \\ &= \mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot \mid x)} \left[w^2(x) \cdot \left\{ \sigma(r^*(x_1, \bar{\mathbf{y}}_1^a) - r^*(x_1, \bar{\mathbf{y}}_1^b)) - \mathbb{1}\{\bar{\mathbf{y}}_1^a = \bar{\mathbf{y}}_1^w, \bar{\mathbf{y}}_1^b = \bar{\mathbf{y}}_1^{\ell}\} \right\}^2 \cdot \mathbf{g}_1^* (\mathbf{g}_1^*)^{\top} \right]. \end{aligned} \quad (56)$$

Taking the conditional expectation over the outcomes of winners and losers, and using the relation

$$\begin{aligned} &\mathbb{E} \left[\left\{ \sigma(r^*(x_1, \bar{\mathbf{y}}_1^a) - r^*(x_1, \bar{\mathbf{y}}_1^b)) - \mathbb{1}\{\bar{\mathbf{y}}_1^a = \bar{\mathbf{y}}_1^w, \bar{\mathbf{y}}_1^b = \bar{\mathbf{y}}_1^{\ell}\} \right\}^2 \mid x_1, \bar{\mathbf{y}}_1^a, \bar{\mathbf{y}}_1^b \right] \\ &= \text{Var} \left(\mathbb{1}\{\bar{\mathbf{y}}_1^a = \bar{\mathbf{y}}_1^w, \bar{\mathbf{y}}_1^b = \bar{\mathbf{y}}_1^{\ell}\} \mid x_1, \bar{\mathbf{y}}_1^a, \bar{\mathbf{y}}_1^b \right) \\ &= \sigma(r^*(x_i, \bar{\mathbf{y}}_i^a) - r^*(x_i, \bar{\mathbf{y}}_i^b)) \sigma(r^*(x_i, \bar{\mathbf{y}}_i^b) - r^*(x_i, \bar{\mathbf{y}}_i^a)), \end{aligned}$$

we reduce equation equation 56 to

$$\tilde{\Omega} = \mathbb{E}_{x \sim \rho; (\bar{\mathbf{y}}^a, \bar{\mathbf{y}}^b) \sim \bar{\mu}(\cdot \mid x)} \left[w^2(x) \cdot \text{Var}(\mathbb{1}\{\bar{\mathbf{y}}_1^a = \bar{\mathbf{y}}_1^w, \bar{\mathbf{y}}_1^b = \bar{\mathbf{y}}_1^{\ell}\} \mid x_1, \bar{\mathbf{y}}_1^a, \bar{\mathbf{y}}_1^b) \cdot \mathbf{g}_1^* (\mathbf{g}_1^*)^{\top} \right].$$

Bounding the weight function $w(x)$ by its uniform bound $\|w\|_{\infty}$, we simplify further:

$$\tilde{\Omega} \preceq \|w\|_{\infty} \cdot \mathbb{E} \left[w(x) \cdot \text{Var}(\mathbb{1}\{\bar{\mathbf{y}}_1^a = \bar{\mathbf{y}}_1^w, \bar{\mathbf{y}}_1^b = \bar{\mathbf{y}}_1^{\ell}\} \mid x_1, \bar{\mathbf{y}}_1^a, \bar{\mathbf{y}}_1^b) \cdot \mathbf{g}_1^* (\mathbf{g}_1^*)^{\top} \right].$$

This ultimately reduces to

$$\tilde{\Omega} \preceq \|w\|_{\infty} \cdot \Sigma_{*} \quad (57)$$

where Σ_{*} is defined in equation equation 11.

Finally, by combining equations equation 55 and equation 57, we establish the asymptotic normality of $\nabla_{\theta} \hat{\mathcal{L}}(\theta^*)$ and complete the proof of Lemma C.6.

D.3 PROOF OF AUXILIARY RESULTS FOR THEOREM 4.3

This section contains the proofs of the auxiliary results supporting Theorem 4.3. In Appendix D.3.1, we derive the explicit form of the Hessian $\nabla_{\theta}^2 J(\pi^*)$. Appendix D.3.2 rigorously establishes the asymptotic distribution of the value gap (equation equation 32). Finally, Appendix D.3.3 proves the tail bound equation 10 on the chi-square distribution χ_d^2 .

D.3.1 PROOF OF EQUATION EQUATION 9 FROM THEOREM 4.3, EXPLICIT FORM OF HESSIAN $\nabla_{\theta}^2 J(\pi^*)$

We begin by differentiating expression equation 39 for the gradient $\nabla_{\theta} J(\pi_{\theta})$ to obtain the Hessian matrix $\nabla_{\theta}^2 J(\pi_{\theta})$. The resulting expression can be written as

$$\nabla_{\theta}^2 J(\pi_{\theta}) = \Gamma_1 + \Gamma_2 + \Gamma_3,$$

where the terms are defined as follows:

$$\begin{aligned}\Gamma_1 &:= \frac{1}{\beta} \mathbb{E}_{x \sim \rho} \left[\int_{\mathcal{Y}} \{r^*(x, \vec{y}) - r_\theta(x, \vec{y})\} \left\{ \nabla_\theta r_\theta(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_\theta(\cdot|x)} [\nabla_\theta r_\theta(x, \vec{y}')] \right\} \nabla_\theta \pi_\theta(d\vec{y} | x)^\top \right], \\ \Gamma_2 &:= -\frac{1}{\beta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_\theta(\cdot|x)} \left[\left\{ \nabla_\theta r_\theta(x, \vec{y}) - \mathbb{E}_{\vec{y}' \sim \pi_\theta(\cdot|x)} [\nabla_\theta r_\theta(x, \vec{y}')] \right\} \nabla_\theta r_\theta(x, \vec{y})^\top \right], \\ \Gamma_3 &:= \frac{1}{\beta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_\theta(\cdot|x)} \left[\{r^*(x, \vec{y}) - r_\theta(x, \vec{y})\} \left\{ \nabla_\theta^2 r_\theta(x, \vec{y}) - \nabla_\theta \mathbb{E}_{\vec{y}' \sim \pi_\theta(\cdot|x)} [\nabla_\theta r_\theta(x, \vec{y}')] \right\} \right].\end{aligned}$$

At the point $\theta = \theta^*$, we know that $r_\theta = r^*$. This simplifies the expression significantly:

$$\Gamma_1 = \mathbf{0} \quad \text{and} \quad \Gamma_3 = \mathbf{0}.$$

Therefore, only term Γ_2 contributes to the Hessian, and it further reduces to

$$\begin{aligned}\Gamma_2 &= -\frac{1}{\beta} \mathbb{E}_{x \sim \rho, \vec{y} \sim \pi_\theta(\cdot|x)} \left[\nabla_\theta r_\theta(x, \vec{y}) \nabla_\theta r_\theta(x, \vec{y})^\top \right] \\ &\quad + \frac{1}{\beta} \mathbb{E}_{x \sim \rho} \left[\mathbb{E}_{\vec{y}' \sim \pi_\theta(\cdot|x)} [\nabla_\theta r_\theta(x, \vec{y}')] \mathbb{E}_{\vec{y} \sim \pi_\theta(\cdot|x)} [\nabla_\theta r_\theta(x, \vec{y})]^\top \right] \\ &= -\frac{1}{\beta} \mathbb{E}_{x \sim \rho} \left[\text{Cov}_{\vec{y} \sim \pi_\theta(\cdot|x)} [\nabla_\theta r_\theta(x, \vec{y}) | x] \right].\end{aligned}$$

From this simplification, we deduce

$$\nabla_\theta^2 J(\pi^*) = -\frac{1}{\beta} \mathbb{E}_{x \sim \rho} \left[\text{Cov}_{\vec{y} \sim \pi^*(\cdot|x)} [\nabla_\theta r^*(x, \vec{y}) | x] \right],$$

which establishes equation 9 as stated in Theorem 4.3.

D.3.2 PROOF OF THE ASYMPTOTIC DISTRIBUTION IN EQUATION 32

The goal of this part is to establish the asymptotic distribution of $n\{J(\pi^*) - J(\hat{\pi})\}$, as stated in equation 32 from Appendix C.2.3. To achieve this, we first recast the value gap into the product of two terms and then invoke Slutsky's theorem.

We start by writing

$$n \cdot \{J(\pi^*) - J(\hat{\pi})\} = \underbrace{n \cdot (\hat{\theta} - \theta^*)^\top \mathbf{H} (\hat{\theta} - \theta^*)}_{U_n} \cdot \underbrace{\frac{J(\pi^*) - J(\hat{\pi})}{(\hat{\theta} - \theta^*)^\top \mathbf{H} (\hat{\theta} - \theta^*)}}_{V_n}. \quad (58)$$

By isolating U_n and V_n in this way, we can handle their limiting behaviors separately:

$$U_n \xrightarrow{d} \mathbf{z}^\top \Omega^{\frac{1}{2}} \mathbf{H} \Omega^{\frac{1}{2}} \mathbf{z} \quad \text{with } \mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \quad (59a)$$

$$V_n \xrightarrow{p} \frac{1}{2}. \quad (59b)$$

If these two results are established, the desired asymptotic distribution of the value gap, as given in equation 32, follows directly from Slutsky's theorem.

To complete the proof, we proceed to verify equations 59a and equation 59b. It is worth noting that equation 59a is a straightforward corollary of Theorem 4.2, so the main task is to establish the convergence result in equation 59b.

Proof of Equation 59b: Since $\Sigma_\star$ is nonsingular, the matrix $\mathbf{H} = (\bar{\mathcal{Z}}_\theta / \beta) \cdot \Sigma_\star$ is also nonsingular. From equation 31, we know that for any $\varepsilon \in (0, 1)$, there exists a threshold $\eta(\varepsilon) > 0$ such that whenever $\|\theta - \theta^*\|_2 \leq \eta(\varepsilon)$, the following inequality holds:

$$\left(\frac{1}{2} - \varepsilon\right) (\theta - \theta^*)^\top \mathbf{H} (\theta - \theta^*) \leq J(\pi^*) - J(\pi_\theta) \leq \left(\frac{1}{2} + \varepsilon\right) (\theta - \theta^*)^\top \mathbf{H} (\theta - \theta^*).$$

This can be reformulated as

$$\left| V_n - \frac{1}{2} \right| \leq \varepsilon.$$

Next, under the condition that $\hat{\theta} \xrightarrow{P} \theta^*$, for any $\delta > 0$, there exists an integer $N(\varepsilon, \delta) \in \mathbb{Z}_+$ such that for any $n \geq N(\varepsilon, \delta)$,

$$\mathbb{P}\{\|\hat{\theta} - \theta^*\|_2 > \eta(\varepsilon)\} \leq \delta.$$

Therefore, for any $n \geq N(\varepsilon, \delta)$, we can conclude

$$\mathbb{P}\left\{\left|V_n - \frac{1}{2}\right| > \varepsilon\right\} \leq \delta.$$

In simpler terms, $V_n \xrightarrow{P} \frac{1}{2}$, which establishes equation equation 59b.

D.3.3 PROOF OF THE TAIL BOUND IN EQUATION EQUATION 10

We now establish the tail bound

$$\mathbb{P}\{\chi_d^2 > (1 + \varepsilon)d\} \leq \exp\left\{-\frac{d}{2}(\varepsilon - \log(1 + \varepsilon))\right\}, \quad (60)$$

as stated in equation equation 10.

We first note that the moment-generating function (MGF) of distribution χ_d^2 is given by

$$M_{\chi_d^2}(t) = (1 - 2t)^{-\frac{d}{2}}, \quad \text{for any } t < \frac{1}{2}.$$

Using Markov's inequality, for any $t > 0$, we have

$$\mathbb{P}\{\chi_d^2 > (1 + \varepsilon)d\} \leq \exp\{-t(1 + \varepsilon)d\} \cdot M_{\chi_d^2}(t) = \exp\{-t(1 + \varepsilon)d\} \cdot (1 - 2t)^{-\frac{d}{2}}, \quad \text{for any } t < \frac{1}{2}. \quad (61)$$

We optimize the bound by choosing t to minimize the exponent $-t(1 + \varepsilon)d - \frac{d}{2} \log(1 - 2t)$. Solving for the optimal t , we obtain

$$t = \frac{\varepsilon}{2(1 + \varepsilon)}.$$

Substituting t back into inequality equation 61, the bound simplifies to the desired inequality equation 60.

E SUPPORTING THEOREM: MASTER THEOREM FOR Z -ESTIMATORS

In this section, we provide a brief introduction to the master theorem for Z -estimators for the convenience of the readers.

Let the parameter space be Θ , and consider a data-dependent function $\Psi_n : \Theta \rightarrow \mathbb{L}$, where $\mathbb{L}$ is a metric space with norm $\|\cdot\|_{\mathbb{L}}$. Assume that the parameter estimate $\hat{\theta}_n \in \Theta$ satisfies $\|\Psi_n(\hat{\theta}_n)\|_{\mathbb{L}} \xrightarrow{P} 0$, making $\hat{\theta}_n$ a Z -estimator. The function Ψ_n is an estimator of a fixed function $\Psi : \Theta \rightarrow \mathbb{L}$, where $\Psi(\theta_0) = 0$ for some parameter of interest $\theta_0 \in \Theta$.

Theorem E.1 (Theorem 2.11 in Kosorok (2008), master theorem for Z -estimators). *Suppose the following conditions hold:*

1. $\Psi(\theta_0) = 0$, where θ_0 lies in the interior of Θ .
2. $\sqrt{n}\Psi_n(\hat{\theta}_n) \xrightarrow{P} 0$ and $\|\hat{\theta}_n - \theta_0\| \xrightarrow{P} 0$ for the sequence of estimators $\{\hat{\theta}_n\} \subset \Theta$.
3. $\sqrt{n}(\Psi_n - \Psi)(\theta_0) \xrightarrow{d} Z$, where Z is a tight⁴ random variable.
4. The following smoothness condition is satisfied:

$$\frac{\|\sqrt{n}(\Psi_n(\hat{\theta}_n) - \Psi(\hat{\theta}_n)) - \sqrt{n}(\Psi_n(\theta_0) - \Psi(\theta_0))\|_{\mathbb{L}}}{1 + \sqrt{n}\|\hat{\theta}_n - \theta_0\|} \xrightarrow{P} 0. \quad (62)$$

⁴A random variable Z is tight if, for any $\epsilon > 0$, there exists a compact set $K \subset \mathbb{R}$ such that $\mathbb{P}(Z \notin K) < \epsilon$.

Additionally, assume that $\theta \mapsto \Psi(\theta)$ is Fréchet differentiable⁵ at θ_0 with derivative $\dot{\Psi}_{\theta_0}$, and that $\dot{\Psi}_{\theta_0}$ is continuously invertible⁶. Then

$$\|\sqrt{n}\dot{\Psi}_{\theta_0}(\hat{\theta}_n - \theta_0) + \sqrt{n}(\Psi_n - \Psi)(\theta_0)\|_{\mathbb{L}} \xrightarrow{P} 0$$

and therefore

$$\sqrt{n}(\hat{\theta}_n - \theta_0) \xrightarrow{d} -\dot{\Psi}_{\theta_0}^{-1} Z.$$

F EXPERIMENTAL DETAILS

We implement our code based on the open-sourced OpenRLHF framework Hu et al. (2024). We will open-source our code in the camera-ready version.

We use both the helpful and the harmless (HH) sets from HH-RLHF (Bai et al., 2022) without additional data selection. We adopt the chat template from the Skywork-Reward-8B model (Liu et al., 2024a) to align with the reward template. This reward model, fine-tuned from Llama-3.1-8B, is used to simulate human preference labeling and matches our network trained for alignment.

For SFT, we apply full-parameter tuning with Adam for one epoch, using a cosine learning rate schedule, a 3% warmup phase, a learning rate of 5×10^{-7} , and a batch size of 256. These hyperparameters are adopted from Hu et al. (2024).

For all the DPO training in both iterative and online settings, we use full-parameter tuning with Adam but with two epochs. The learning rate, warmup schedules, and batch size are all the same.

During generation, we limit the maximum number of new tokens to 896 and employ top- p decoding with $p = 0.95$ for all experiments. For Online DPO, we use a sampling temperature of 1.0, following Guo et al. (2024), while in Iterative DPO, we set the temperature to 0.7 to account for the off-policy nature of the data, following Dong et al. (2024); Shi et al. (2024).

Prompts are truncated to a maximum length of 512 tokens (truncated from the left if the length exceeds this limit) for SFT, DPO, and generation tasks. For SFT data, the maximum length is further restricted to 1024 tokens. When the combined length of the response and the (truncated) prompt exceeds 1024 tokens, the response is truncated from the right. These truncation practices align with the standard methodology described by Rafailov et al. (2023). In contrast, for DPO, responses are not further truncated, as we are already limiting the maximum tokens generated during the generation process.

When reproducing the *Hybrid Sampling* baseline (Exploration Preference Optimization, XPO) from Xie et al. (2024), we use $\alpha = 5 \times 10^{-6}$ as suggested in the paper.

We do not include a comparison with Shi et al. (2024) and Liu et al. (2024d) in our experiments. While Shi et al. (2024) employs a sampling method similar to ours, their approach requires significantly more hyperparameters to tune, whereas our method involves no hyperparameter tuning. On the other hand, Liu et al. (2024d) relies on training an ensemble of 20 reward models to approximate the posterior. Their sampling method requires solving the argmax of these rewards, which is computationally intractable. As a workaround, they generate 20 samples and select the best one using best-of- N with $N = 20$. This approach demands at least six times the computational resources compared to our method.

F.1 ADDITIONAL RESULTS

We present the full results for Online DPO with the overfitted initial policy, including a scatter plot in Figure 5 and a summary of the objective values in Table 4.

⁵Fréchet differentiability: A map $\phi : \mathbb{D} \rightarrow \mathbb{L}$ is Fréchet differentiable at θ if there exists a continuous, linear map $\phi'_\theta : \mathbb{D} \rightarrow \mathbb{L}$ such that $\|\phi(\theta + h_n) - \phi(\theta) - \phi'_\theta(h_n)\|_{\mathbb{L}}/\|h_n\| \rightarrow 0$ for all sequences $\{h_n\} \subset \mathbb{D}$ with $\|h_n\| \rightarrow 0$ and $\theta + h_n \in \mathbb{D}$ for all $n \geq 1$.

⁶Continuous invertibility: A map $A : \Theta \rightarrow \mathbb{L}$ is continuously invertible if A is invertible, and there exists a constant $c > 0$ such that $\|A(\theta_1) - A(\theta_2)\|_{\mathbb{L}} \geq c\|\theta_1 - \theta_2\|$ for all $\theta_1, \theta_2 \in \Theta$.

We observe that *Vanilla Sampling* rapidly increases its KL divergence from the reference model while its reward improvement diminishes over time. In contrast, PILAF undergoes an early phase of training with fluctuating KL values but ultimately achieves a policy with higher reward and substantially lower KL divergence. We hypothesize that PILAF’s interpolation-based exploration enables it to escape the suboptimal region of the loss landscape where *Vanilla Sampling* remains trapped.

Conversely, *Hybrid Sampling*, despite its explicit exploration design, remains biased by the policy model and continues to exhibit high KL values. While KL divergence decreases over training, the reward improvement remains limited. Meanwhile, *Best-of-N Sampling* introduces an implicit exploration mechanism through internal DPO, which selects the best and worst responses, leading to wider coverage than *Vanilla Sampling*. However, despite achieving a KL divergence similar to PILAF, it results in a lower reward. These findings highlight the superiority of PILAF sampling, demonstrating its effectiveness in robustly optimizing an overfitted policy.

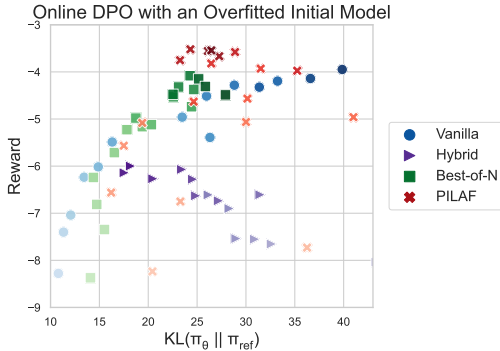


Figure 5: **Online DPO with an overfitted initial policy.** Full results of the Figure 4. Each dot represents an evaluation performed every 50 training steps. Color saturation indicates the training step, with darker colors representing later steps.

Table 4: **Results of Online DPO with an overfitted initial policy.** We report the average reward, KL divergence from the reference model, and objective J on the testset.

METHOD	REWARD ($\uparrow$)	KL ($\downarrow$)	J ($\uparrow$)
<i>Vanilla</i>	<u>-3.95</u>	39.85	-7.93
<i>Best-of-N</i>	-4.49	27.90	<u>-7.28</u>
<i>Hybrid</i>	-6.00	18.20	-7.82
PILAF (OURS)	-3.54	<u>26.45</u>	-6.19

G EXTENSION TO PROXIMAL POLICY OPTIMIZATION (PPO)

In this section, we briefly explore how the core principles of our PILAF sampling approach can be extended to PPO-based RLHF methods.

Integrating Response Sampling in InstructGPT: The PPO-based RLHF pipeline used in InstructGPT (Ouyang et al., 2022) consists of three key steps:

- (i) Supervised Fine-Tuning (SFT) that produces the reference model π_{ref} .
- (ii) Reward Modeling (RM) by solving the optimization problem equation 2, yielding an estimated reward function r_θ .
- (iii) Reinforcement Learning Fine-Tuning, where the policy π_ϕ is optimized against the reward model r_θ using the Proximal Policy Optimization (PPO) algorithm, following the optimization scheme equation 4.

The key distinction between the PPO and DPO approaches lies in how the reward model r_θ is represented—explicitly in PPO and implicitly in DPO. In response sampling for data collection, it is crucial to consider the iterative nature of the InstructGPT pipeline. During each iteration, additional human-labeled data is collected for reward modeling (step (ii)), and steps (ii) and (iii) are repeatedly applied to refine the model. Our proposed PILAF algorithm naturally integrates into this pipeline by improving the data collection process in step (ii), thereby enhancing reward model training and, in turn, policy optimization.

Extensions of T-PILAF and PILAF: Extending our response sampling methods, PILAF and T-PILAF, to the PPO setup with an explicit r_θ is both natural and straightforward.

- Within the theoretical framework of T-PILAF, as introduced in Section 3, the only required modification is replacing π_θ with the language model π_ϕ and redefining the interpolated and extrapolated policies, π_ϕ^+ and π_ϕ^- , following the same formulation as in equations equation 6a and equation 6b. Specifically, we define

$$\pi_\phi^+(\vec{y} | x) := \frac{1}{Z^+(x)} \pi_\phi(\vec{y} | x) \exp \{ r_\theta(x, \vec{y}) \}, \quad (63a)$$

$$\pi_\phi^-(\vec{y} | x) := \frac{1}{Z^-(x)} \pi_\phi(\vec{y} | x) \exp \{ - r_\theta(x, \vec{y}) \}, \quad (63b)$$

where r_θ is now explicitly produced by a reward network, rather than being implicitly derived from π_ϕ , as in equation equation 5.

- To extend our empirical PILAF algorithm, as described in Section 5, we propose applying the same interpolation and extrapolation techniques directly to the logits of the language models π_ϕ and π_{ref} . In particular, we take

$$\pi_\phi^+(\cdot | x, y_{1:t-1}) = \text{softmax} \left(\{ (1 + \beta) \mathbf{h}_\phi - \beta \mathbf{h}_{\text{ref}} \} (x, y_{1:t-1}) \right),$$

$$\pi_\phi^-(\cdot | x, y_{1:t-1}) = \text{softmax} \left(\{ (1 - \beta) \mathbf{h}_\phi + \beta \mathbf{h}_{\text{ref}} \} (x, y_{1:t-1}) \right),$$

where $\mathbf{h}_\phi$ and $\mathbf{h}_{\text{ref}}$ represent the logits of the language models π_ϕ and π_{ref} , respectively.

Adaption of Theoretical Analysis: Our theoretical analyses can be extended to the PPO framework, assuming that the optimization process equation 4 in step (iii) of InstructGPT is solved exactly. In this case, the policy satisfies $\pi_\phi = \pi_{r_\theta}$, where

$$\pi_{r_\theta}(\vec{y} | x) := \frac{1}{Z_\theta(x)} \pi_{\text{ref}}(\vec{y} | x) \exp \left\{ \frac{1}{\beta} r_\theta(x, \vec{y}) \right\}.$$

Under this assumption, the output language model π_ϕ is implicitly a function of the parameter θ . Building on this, we can adapt our optimization and statistical analyses as follows:

- **Optimization Consideration:** Using the same argument as in Theorem 4.1, we can prove that

$$\nabla_\theta \mathcal{L}(\theta) = -C' \cdot \nabla_\theta J(\pi_\phi) + T_2,$$

where $C' > 0$ is a universal constant, and T_2 represents a second-order approximation error. In other words, if the policy optimization step is sufficiently accurate for the reward model r_θ , then performing gradient descent on the MLE loss with respect to θ is equivalent to applying gradient ascent on the oracle objective J , following the steepest direction in the parameter space of θ .

- **Statistical Consideration:** Even with the new parameterization, the asymptotic distribution of $\hat{\theta}$ from Theorem 4.2 remains unchanged. Moreover, the gradient and Hessian of J with respect to θ retain the same form as in Theorem 4.1. As a result, the statistical analysis extends naturally to PPO, allowing us to conclude that PILAF also maintains structure-invariant statistical efficiency for PPO methods.