

Sample-Efficient Alignment for LLMs

Anonymous Author(s)

Affiliation

Address

email

Abstract

We study methods for efficiently aligning large language models (LLMs) with human preferences given budgeted online feedback. We first formulate the LLM alignment problem in the frame of contextual dueling bandits. This formulation, subsuming recent paradigms such as online RLHF and online DPO, inherently quests for sample-efficient algorithms that incorporate *online active exploration*. Leveraging insights from bandit theory, we introduce a unified algorithm based on **Thompson sampling** and highlight its applications in two distinct LLM alignment scenarios. The practical agent that efficiently implements this algorithm, named **SEA** (Sample-Efficient Alignment), is empirically validated through extensive experiments across three model scales (1B, 2.8B, 6.9B) and three preference learning algorithms (DPO, IPO, SLiC). The results demonstrate that **SEA** achieves highly sample-efficient alignment with oracle’s preferences, outperforming recent active exploration methods for LLMs. We will release our codebase to hopefully accelerate future research in this field.

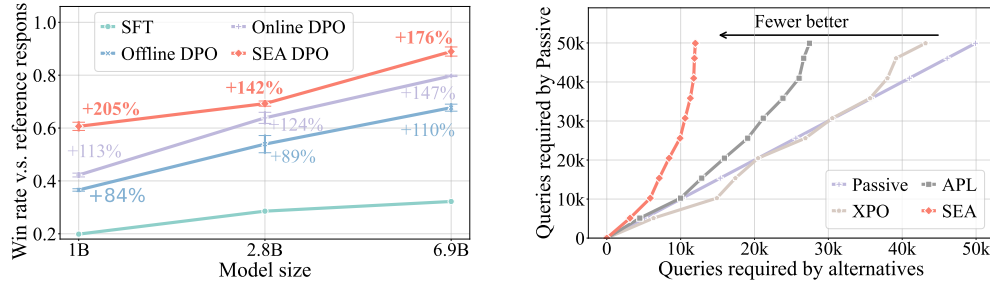


Figure 1: Win rate comparison of model responses against reference responses on the TL;DR task, judged by the preference oracle. All compared methods use the same optimization method (DPO). **(Left)** Performance improvements at convergence over SFT models achieved by offline (Offline DPO), passively online (Online DPO), and our *active exploration* (**SEA** DPO) methods. **(Right)** The number of queries required by the passively online method (Passive) versus that by different active exploration methods to attain various levels of win rates. **SEA** achieves the best sample efficiency for online alignment compared to XPO and APL.

1 Introduction

Aligning LLMs with human preferences is a crucial step to elicit various desirable behaviors, e.g., helpfulness and harmlessness [5]. Moreover, it holds the potential to create superhuman capabilities with only human-level feedback, as verifying is believed to be easier than synthesizing novel behaviors. By iteratively generating new candidates and asking for human feedback, LLMs could learn to reinforce good behaviors and may eventually surpass human capabilities.

Existing methods, either via reinforcement learning from human feedback (RLHF) [65, 50] or direct alignment from preferences (DAP) [55, 4], typically require a large amount of human annotations to achieve effective alignment. As a result, the volume of human feedback becomes a major bottleneck in practical alignment scenarios. This poses a challenging and under-explored research question:

How to align LLMs sample-efficiently?

To seek answers, in Sec. 2, we formalize LLM alignment as a contextual dueling bandit (CDB) [85, 20], where the agent (i.e., the learner and decision maker, in our case the LLM) interacts with the environment (i.e., human) to collect experience for improving its policy. This formulation naturally calls for two key properties for alignment algorithms to be sample-efficient:

Property 1 (Online interaction). Interacting and learning *online* allows the agent to act with the latest learned policy and then use that experience to immediately improve the policy.

Property 2 (Active exploration). An *actively exploring* agent strategically selects actions such that the collected experience leads to maximal policy improvement.

Since the CDB formulation is general and almost subsumes all existing LLM alignment methods, it provides us a lens to scrutinize prior methods on the axes of Properties 1 and 2. In Sec. 3, we thoroughly discuss prior alignment approaches, ranging from offline learning [55, 4] and passive learning with iterative [15, 18] or online interaction [24], to active exploration for learning preference models [21] or aligning LLMs [47, 86, 79]. As will be revealed, most prior methods (partially) fail to satisfy the two properties, resulting in inferior sample efficiency. Moreover, through the CDB formulation, we identify two LLM alignment scenarios, namely aligning from online users’ feedback (e.g., ChatGPT [13]) and aligning from crowdsourcing [15, 50], and shed light on their correspondences to two bandit settings (explore & exploit and best arm identification). Understanding their differences is important for designing efficient alignment algorithms for respective scenarios. We detail these two settings in Sec. 2 and discuss how prior works approach them in Sec. 3.

Leveraging algorithmic insights from bandit theory, our answer to the research question above is a principled alignment algorithm based on Thompson sampling (TS) [71]. Our method fulfills Properties 1 and 2 to enhance sample efficiency, and it solves either of the two settings depending on practical scenarios (Sec. 4.1). We incorporate techniques including *epistemic reward model*, *policy-guided search* and *mixed preference learning* to implement the proposed TS algorithm (Sec. 4.2), yielding a practical agent which we call **SEA** (Sample-Efficient Alignment). In addition, we develop and will open source a highly efficient, distributed learning system for studying online LLM alignment methods (Sec. 5), eliminating barriers to *fair* empirical comparisons of different alignment algorithms. Through extensive experiments (Sec. 6), **SEA** shows strong empirical results (see Fig. 1), consistently achieving higher win rates and improved sample efficiency compared to baseline approaches across three model scales. We will open source the codebase to hopefully accelerate future research in this field.

In summary, the contributions of this work are:

- Through the lens of contextual dueling bandits, we propose a principled *Thompson sampling* algorithm for LLM online exploration, addressing both *explore & exploit* and *best arm identification* settings.
- We develop two novel techniques to approximate Thompson sampling in LLM’s large action space: policy-guided search and mixed preference learning. Thompson sampling requires sampling a reward function from the posterior distribution and generating the sequence that maximizes the sampled reward function. For **policy-guided search**, we use an existing epistemic reward model for approximating the posterior and propose an approximate maximization method based on sampling a finite set of sequences from the LLM, and doing maximization on the finite sample. However, maintaining and updating a separate LLM for each reward function as suggested by Thompson sampling would be prohibitively expensive, thus **mixed preference learning** is introduced to align the LLM with internal reward functions to better approximate the maximization.
- To our knowledge, we are the first to study active exploration for LLM alignment with fully online experimental verification. The online alignment codebase will be open sourced to accelerate future studies.

2 LLM alignment as contextual dueling bandits

We first review the definitions and two typical objectives of *Contextual Dueling Bandits* (Sec. 2.1), then translate them into the language of *LLM alignment* (Sec. 2.2). The tight connection between them, as we will see, allows us to leverage insights from bandit algorithms to design efficient alignment algorithms for LLMs.

2.1 Contextual dueling bandits

Contextual dueling bandits (CDB) [85, 20] is proposed to study online learning problems where the feedback consists of relative pairwise comparisons. A CDB problem can be characterized by a tuple $(\mathcal{C}, \mathcal{A}, \mathbb{P})$, where $\mathcal{C}$ is the context space, $\mathcal{A}$ is the action space, and $\mathbb{P} : \mathcal{A} \times \mathcal{A} \times \mathcal{C} \mapsto [0, 1]$ denotes the unknown *preference oracle*. An agent learns by iteratively interacting with the environment (i.e., the preference oracle $\mathbb{P}$) as follows. At each round t of the learning process, a context $c_t \sim p_{\mathcal{C}}$ is presented to the agent, who needs to take two actions $\mathbf{a}_t, \mathbf{a}'_t \in \mathcal{A}$ for a “dueling” comparison. The agent then receives stochastic feedback in the form of a comparison result $z_t \sim \text{Ber}(\mathbb{P}(\mathbf{a}_t \succ \mathbf{a}'_t | c_t))$ from the environment, where $\text{Ber}(\cdot)$ is the Bernoulli distribution and $\succ$ denotes that the first action is preferred.

Regret. The quality of the dueling actions selected by the agent is measured by the *immediate regret*: $R_t = \mathbb{P}(\mathbf{a}_t^* \succ \mathbf{a}_t | c_t) + \mathbb{P}(\mathbf{a}_t^* \succ \mathbf{a}'_t | c_t) - 1$, where $\mathbf{a}_t^*$ is the best action¹ the agent would take at round t if it had complete knowledge of $\mathbb{P}$. Intuitively, if the agent has learned how to act optimally from round t onwards, it would no longer suffer any regret since its actions would be indistinguishable from the best action ($\mathbb{P}(\mathbf{a}_t^* \succ \mathbf{a}_t | c_t) = \frac{1}{2}$ hence $R_t = 0$ for $t \geq t$).

Optimal policy. A policy $\pi \in \Delta_{\mathcal{A}}^{\mathcal{C}}$ ² associates each context $c \in \mathcal{C}$ with a probability distribution $\pi(\cdot | c) \in \Delta_{\mathcal{A}}$ over the action space. The *total preference* of policy π over policy μ given a context sampling distribution $p_{\mathcal{C}}$ and a preference oracle $\mathbb{P}$ is defined as

$$P_{p_{\mathcal{C}}, \mathbb{P}}(\pi \succ \mu) = \mathbb{E}_{c \sim p_{\mathcal{C}}} [\mathbb{E}_{\mathbf{a} \sim \pi(\cdot | c)} \mathbb{E}_{\mathbf{a}' \sim \mu(\cdot | c)} [\mathbb{P}(\mathbf{a} \succ \mathbf{a}' | c)]] . \quad (1)$$

We adopt the *von Neumann winner* [20] as the solution concept, which requires the optimal policy π^* to satisfy that

$$\forall \pi' \in \Delta_{\mathcal{A}}^{\mathcal{C}}, P_{p_{\mathcal{C}}, \mathbb{P}}(\pi^* \succ \pi') \geq \frac{1}{2}. \quad (2)$$

In words, the von Neumann winner policy should beat or tie with every policy (i.e., is zero-regret) on average.

Learning objectives. The goal of bandit agents is to learn an optimal policy through interactions with the environment. There are two subtypes of objectives that focus on different learning scenarios. The first type considers the conventional *explore and exploit* (E&E) setting [59, 3], where the agent learns fully **online** and tries to minimize the cumulative regret over T rounds: $\sum_{t=1}^T R_t$. The second type of objective concerns the *best arm identification* (BAI) setting [9, 2], where the agent is only evaluated **offline** on its average performance, possibly at any round (a.k.a., anytime regret), and tries to learn the optimal policy with minimum interaction. Both settings call for effective *online exploration* strategies that satisfy Properties 1 and 2. Their differences will be made clearer with real scenarios in Sec. 2.2.

2.2 Online alignment as CDB

Online LLM alignment can be framed as a CDB problem. Specifically, at time t a text prompt (cf. context) $\mathbf{x}_t \in \mathcal{X}$ is sampled from a prompt distribution $p_{\mathcal{X}}$. Then, two distinct responses (cf. actions), $\mathbf{y}_t, \mathbf{y}'_t \in \mathcal{Y}$, are chosen by the agent, and presented to human annotators (cf. the environment) for preference ranking. The winning and losing responses are labeled as $(\mathbf{y}_t^+, \mathbf{y}_t^-)$ based on a binary stochastic feedback $z_t \sim \text{Ber}(\mathbb{P}(\mathbf{y}_t \succ \mathbf{y}'_t | \mathbf{x}_t))$. The agent is expected to produce *good* responses satisfying either E&E or BAI objectives, with knowledge learned from the experience accumulated so far: $\mathcal{D}_t = \{(\mathbf{x}_{\tau}, \mathbf{y}_{\tau}^+, \mathbf{y}_{\tau}^-)\}_{\tau=1}^t$. A standard assumption is that human preferences follow the Bradley-Terry (BT) model [8]:

$$\mathbb{P}(\mathbf{y}_t \succ \mathbf{y}'_t | \mathbf{x}_t) = \frac{\exp(r^*(\mathbf{x}_t, \mathbf{y}_t))}{\exp(r^*(\mathbf{x}_t, \mathbf{y}_t)) + \exp(r^*(\mathbf{x}_t, \mathbf{y}'_t))} = \sigma(r^*(\mathbf{x}_t, \mathbf{y}_t) - r^*(\mathbf{x}_t, \mathbf{y}'_t)), \quad (3)$$

where σ is the sigmoid function and r^* encodes human’s implicit reward. The immediate regret of LLM alignment can be rewritten as $R_t = r^*(\mathbf{x}_t, \mathbf{y}_t^*) - (r^*(\mathbf{x}_t, \mathbf{y}_t) + r^*(\mathbf{x}_t, \mathbf{y}'_t)) / 2$ with the BT assumption [62, 39], where $\mathbf{y}_t^*$ is the best response for prompt $\mathbf{x}_t$ given human’s implicit reward, i.e., $r^*(\mathbf{x}_t, \mathbf{y}_t^*) \geq r^*(\mathbf{x}_t, \mathbf{y}), \forall \mathbf{y} \in \mathcal{Y}$. The von Neumann winner policy is also redefined as

$$\pi^* \in \arg \max_{\pi \in \Delta_{\mathcal{Y}}^{\mathcal{X}}} J(\pi), \text{ where } J(\pi) = \mathbb{E}_{\mathbf{x} \sim p_{\mathcal{X}}} \mathbb{E}_{\mathbf{y} \sim \pi(\cdot | \mathbf{x})} [r^*(\mathbf{x}, \mathbf{y})] \text{ is the objective,} \quad (4)$$

¹We assume that a best action $\mathbf{a}^*$ in the sense that $\mathbb{P}(\mathbf{a}^* \succ \mathbf{a} | c) \geq \frac{1}{2}, \forall \mathbf{a} \in \mathcal{A}$ exists for all context $c \in \mathcal{C}$.

²We denote by $\Delta_{\mathcal{A}}^{\mathcal{C}}$ the set of all mappings $\mathcal{C} \mapsto \Delta_{\mathcal{A}}$, where $\Delta_{\mathcal{A}}$ denotes the set of all probability distributions over $\mathcal{A}$.

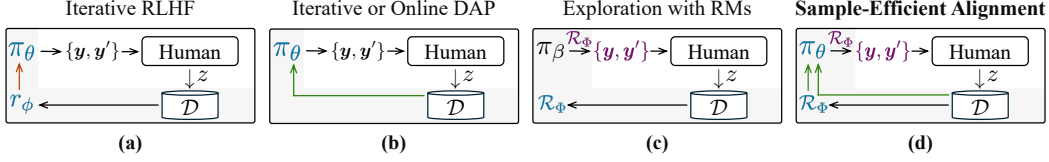


Figure 2: Different paradigms to solve online LLM alignment in the CDB interface. The CDB agent is shaded in gray. We use colors to denote **learnable components**, **RL optimizer**, **direct optimizer**, and **active exploration**. r_ϕ denotes a point estimate of human’s implicit reward, while $\mathcal{R}_\phi$ refers to an uncertainty-aware reward model. Please see Sec. 3 for detailed comparisons with references to prior works. [fdaL: updated fig.3 to highlight the differences between (b) and (d).]

by substituting Eq. (3) into Eq. (1) and maximizing $P_{p_{\mathcal{X}}, \mathbb{P}}(\pi \succ \pi^*)$ towards $1/2$.

The **two settings in bandits** have their respective applications in LLM alignment. (1) The E&E setting applies to the scenario of serving an LLM-based application online and aligning it continually with users’ preferences. In this setting, the agent needs to balance exploration with exploitation, thus the cumulative regret is of interest because the quality of *every* response matters. In fact, commercial systems like ChatGPT would strategically ask users to make a dueling comparison, while upholding the quality of both responses. Please see Fig. 11 in App. I for an example. (2) The BAI setting corresponds to the other scenario where annotators are paid to provide human feedback [15, 50]. The desideratum in this scenario is to align the LLM at the minimum labeling cost, while the quality of the dueling responses is not important as long as the experience helps sample-efficiently learn the von Neumann winner policy.

After formalizing LLM alignment in the framework of CDB and uncovering their tight connections, we next thoroughly discuss existing alignment methods in the CDB framework and reveal the sources of their sample inefficiencies.

3 How prior works (partially) solve LLM alignment as CDB

We first align the notations and terminology used in CDB with commonly referred ones in the LLM community. Previously, we used the term “agent” to denote the learner and decision maker, and referred to its overall behavior as the “policy” π (as in Eq. (4)), following the standard abstraction in RL [67, 68]. However, in the LLM literature, “policy” typically refers to the generative language model alone, excluding components like reward models (RMs) that the agent might additionally build. To avoid confusion, from now on we use π_{θ^t} to denote the generative language model (policy) and r_{ϕ^t} to denote the (optional) RM at time t , both of which are learned from preference data $\mathcal{D}_t$ collected up to time t . We will omit t when the time-indexing is not applicable (i.e., no online interaction) or not important in the context.

RLHF and DAP. Commonly adopted RLHF pipelines [15, 65, 5, 50] first learn a proxy RM with a negative log-likelihood loss:

$$\mathcal{L}_r(\phi|\mathcal{D}) = -\mathbb{E}_{(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-) \sim p_{\mathcal{D}}} [\log \sigma(r_\phi(\mathbf{x}, \mathbf{y}^+) - r_\phi(\mathbf{x}, \mathbf{y}^-))], \quad (5)$$

where $\mathcal{D}$ is collected by querying human annotators using a behavior policy π_{ref} (typically the supervised fine-tuned policy π_{sft}). Afterwards, *offline RL*³ [36, 37] is conducted to learn π_θ with respect to the learned reward r_ϕ internally within the agent (Fig. 2a). However, the learned model π_θ might be inaccurate at regions out of the distribution (o.o.d.) of π_{ref} because little training data can be collected. An effective remedy is to incorporate a pessimistic term to combat the distributional shift, leading to a reformulation of the von Neumann winner policy objective in Eq. (4) as

$$J(\pi_\theta) = \mathbb{E}_{\mathbf{x} \sim p_{\mathcal{X}}} \mathbb{E}_{\mathbf{y} \sim \pi_\theta(\cdot|\mathbf{x})} \left[\underbrace{r_\phi(\mathbf{x}, \mathbf{y})}_{\text{estimated } r^*} - \beta \log \frac{\pi_\theta(\mathbf{y}|\mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}|\mathbf{x})} \right] \quad (6)$$

$$= \mathbb{E}_{\mathbf{x} \sim p_{\mathcal{X}}} \left[\mathbb{E}_{\mathbf{y} \sim \pi_\theta(\cdot|\mathbf{x})} [r_\phi(\mathbf{x}, \mathbf{y})] - \beta D_{\text{KL}}(\pi_\theta(\cdot|\mathbf{x}) || \pi_{\text{ref}}(\cdot|\mathbf{x})) \right], \quad (7)$$

which converts an online objective regarding the human’s implicit reward r^* to an offline objective regarding the proxy reward r_ϕ . The KL penalty in Eq. (15) is widely used for language model

³Offline in the sense that π_θ is not directly learned from online human feedback. See App. C for details.

fine-tuning [29, 80], and PPO [64] has become a default *RL optimizer* to maximize the KL-regularized reward. However, the performance of RLHF is guaranteed only if the preference data $\mathcal{D}$ induced by π_{ref} adequately covers π^* [90], which is often approximated by updating π_{ref} with the latest (improved) π_θ for re-sampling a batch of online experience and repeating Eq. (13) and (15). Prior works typically focus on offline or iterative online (with only a few iterations) settings [80, 18], which may compromise sample efficiency (Property 1).

True online RLHF is difficult due to the complexity and instability of RL optimizers. For example, Huang et al. [27] openly reproduces offline RLHF scaling behaviors but requires many implementation tricks for training, highlighting the difficulties of an online counterpart. Fortunately, the introduction of DAP (or *direct optimizers*) largely simplifies and stabilizes fine-tuning by conducting contrastive supervised learning directly on $\mathcal{D}$ (Fig. 2b). While most DAP works focus on learning from a fixed offline preference dataset, including Zhao et al. [88], Rafailov et al. [55], Azar et al. [4], Meng et al. [46], Zhang et al. [87]), iterative DPO [81] observes improved results when allowing iterative online interaction. Guo et al. [24] further propose OAIF to make DAP faithfully online, satisfying Property 1, and demonstrate that online learning prevents over-fitting and yields continual performance improvement. Nevertheless, it still employs passive exploration strategies (using $\mathbf{y}, \mathbf{y}' \sim \pi_\theta$), hindering sample efficiency (Property 2).

Online exploration in LLMs. A line of recent works [44, 17, 45, 21] adopts the fully online bandit formulation and incorporates *active exploration* with uncertainty-aware RMs for response selection (Fig. 2c). In particular, Mehta et al. [44] consider the E&E setting and develop a UCB-style [3] algorithm; Das et al. [17] instead select the dueling responses with the most uncertain preference estimate, targeting the BAI setting in a pure exploration way; unlike the above, Melo et al. [45] view the problem from the angle of pool-based active learning and propose an acquisition function based on both entropy and epistemic uncertainty; finally, the work by Dwaracherla et al. [21] is the closest to ours in the sense that they apply double Thompson sampling (DTS) [78] for exploration, but DTS is designed for the E&E setting while they evaluate anytime average performance as in the BAI setting. We will show in App. G.1 that pure exploration by Das et al. [17] is not the best choice for BAI, and the objective mismatch in Dwaracherla et al. [21] could lead to suboptimal performance in respective settings. Meanwhile, all these works primarily focus on learning uncertainty-aware RMs online without updating LLM policies. Therefore, all responses are sampled from a fixed proposal policy π_β (or even a fixed dataset), making the data coverage a critical concern.

Another line of research updates LLMs online while incorporating exploration. Zhang et al. [86] and Xie et al. [79] independently propose to learn an optimistic RM to encourage exploration. They leverage the property of DPO [55] to reparameterize RM with policy and conclude with an extra optimistic term in the DPO loss function. Thus, their learning processes are like Fig. 2b but with an optimistic direct optimizer. Muldrew et al. [47] adopt the vanilla DPO loss but utilize the implicit reward margin to actively select dueling responses. Yet, these methods are tightly coupled with DPO and not compatible to other direct optimizers. Their experiments are also limited to a few online iterations, possibly due to the implementation difficulty of a faithfully online learning system. Given their relevance to our approach, we will reproduce them in a fully online manner for fair comparisons in Sec. 6.1. We summarize prior works in Table 2 in App. I.

4 SEA: sample-efficient alignment for LLMs

In this section we present our online exploration agent SEA (Fig. 2d). We first introduce a principled Thompson sampling algorithm inspired by bandit theory (Sec. 4.1), and then derive SEA as its practically efficient implementation (Sec. 4.2). Interestingly, SEA can also be viewed as an instantiation of a classical model-based RL architecture called Dyna [66], for which we defer the discussion to App. C.

4.1 Thompson sampling for LLM alignment

Thompson sampling (TS) [71] is widely adopted for solving bandit problems at scale due to its efficiency and strong empirical performance in general online learning problems [12, 61]. A bandit agent using Thompson sampling typically maintains and incrementally updates a posterior distribution of the oracle reward $p(r|\mathcal{D})$. Meanwhile, the agent takes actions following a greedy policy with respect to a sampled RM: $\mathbf{a}_t = \arg \max_{\mathbf{a}} r(\mathbf{a})$ with $r \sim p_r(\cdot|\mathcal{D})$. This simple yet effective algorithm naturally balances exploration and exploitation: when the agent has limited knowledge about the environment, the posterior estimate exhibits high uncertainty so that the sampled RM could guide the

Algorithm 1 Thompson sampling for LLM alignment (intractable).

Input: Prompt distribution $p_{\mathcal{X}}$, unknown but queryable preference oracle $\mathbb{P}$.

- 1: Initialize experience $\mathcal{D}_0 \leftarrow \emptyset$.
- 2: **for** $t = 1, \dots, T$ **do**
- 3: Receive a prompt $x_t \sim p_{\mathcal{X}}$.
- 4: Sample $r \sim p_r(\cdot | \mathcal{D}_{t-1})$ and set $y_t \leftarrow \arg \max_{b \in \mathcal{Y}} r(x_t, b)$. // Select 1st response y .
- 5: **repeat**
 Sample $r \sim p_r(\cdot | \mathcal{D}_{t-1})$ and set $y'_t \leftarrow \arg \max_{b \in \mathcal{Y}} r(x_t, b)$. // Select 2nd response y' .
 until $y'_t \neq y_t$
 // BAI objective: labeling via crowdsourcing.
- 6: Set $y'_t \leftarrow \arg \max_{b \in \mathcal{Y}} \mathbb{V}[\sigma(r(x_t, y_t) - r(x_t, b))]$, // OR select 2nd response y' .
 where $\mathbb{V}[\cdot]$ computes variance over the posterior $p_r(\cdot | \mathcal{D}_{t-1})$.
- 7: Query $\mathbb{P}$ to label $\{y_t, y'_t\}$, and update experience $\mathcal{D}_t \leftarrow \mathcal{D}_{t-1} \cup \{(x_t, y_t^+, y_t^-)\}$.
- 8: **end for**

// See Algorithm 2 for a practical version.

209 greedy policy to explore; after sufficient experience is gathered, the sampled RM approximates the
210 oracle more closely, allowing the agent to exploit near-optimal policies.

211 In the context of LLM alignment, we leverage the BT assumption (Eq. (3)) to replace the preference
212 oracle $\mathbb{P}$ with human’s implicit reward r^* . This substitution enables us to model the reward posterior
213 $p(r|\mathcal{D})$ in the standard TS framework, preserving the probabilistic structure necessary for effective
214 posterior sampling. Inspired by prior works [78, 23] on non-contextual K -arm bandits and preferential
215 Bayesian optimization problems, we generalize them for LLM alignment and develop a unified algo-
216 rithm as shown in Algorithm 1. Note that we assume for now the LLM agent can be fully described
217 by the posterior $p(r|\mathcal{D})$, and we defer practical reward (r_ϕ) and policy (π_θ) learning to Sec. 4.2.

218 As Algorithm 1 presents, the first response of the duel is always selected via standard TS (Line 4).
219 The selection of the second response varies across different settings. Line 5 will be used for scenarios
220 where preference feedback is collected from online users (the E&E setting). The dueling responses
221 selected in this case will both try to maximize a sampled RM, so that the online user experience is
222 warranted with best effort. However, such algorithm can have poor asymptotic performance for BAI
223 problems [60], because sub-optimal responses with confidently high rewards might be tried for a
224 long time at the expense of not exploring other potentially better choices. In light of this, Line 6
225 provides an alternative for scenarios where we could hire annotators for feedback and low-quality but
226 exploratory responses are safe to try. Specifically, Line 6 selects the second response as the one that
227 maximizes the variance of the preference (Eq. (3)) over the first response y_t . This variance quantifies
228 the *epistemic uncertainty* of the RM, pointing the agent to the maximally informative direction to
229 explore for better sample efficiency.

230 However, Algorithm 1 is yet to be practical for LLM alignment for three main reasons. First,
231 computing and sampling from a reward posterior is intractable for nearly all RMs at LLM scale,
232 which are mostly based on large transformers [35]. Second, even if we managed to approximate the
233 reward posterior, the $\arg \max$ operations for response selection are still intractable since the search
234 space $\mathcal{Y}$ is discrete and massive for token sequences of arbitrary length. Last but not least, an LLM
235 agent [1, 72] typically consists in a generative model π_θ (e.g., a transformer [73]), while the algorithm
236 above is centered around a reward posterior $p(r|\mathcal{D})$ that cannot be easily converted into a generative
237 model. We next detail how [SEA](#) practically addresses the three aforementioned issues.

238 4.2 Practical implementation

239 4.2.1 Epistemic reward model for posterior sampling

240 To implement active exploration with TS, we seek an efficient way to maintain and incrementally
241 update the reward posterior $p(r|\mathcal{D})$. We consider *deep ensemble* for our purpose, due to its capability
242 to model epistemic uncertainty [34] and provable results when applied to TS in linear bandits [54].
243 Specifically, we update a set of plausible RMs independently and online, using the preference data

244 and a regularized negative log-likelihood loss:

$$\mathcal{L}_{\mathcal{R}}(\Phi^t|\mathcal{D}_t) = \sum_{k=1}^K (\mathcal{L}_r(\phi_k^t|\mathcal{D}_t) - \lambda \|\phi_k^t - \phi_k^0\|), \quad (8)$$

245 where $\mathcal{L}_r$ is defined in Eq. (13), $\Phi^t = \{\phi_k^t\}_{k=1}^K$ contains the weights of the ensemble of size K ,
 246 and λ controls the regularization towards individual initial weights ϕ_k^0 . Each ensemble member is
 247 initialized independently with random weights, and then trained with regularization to maintain the
 248 diversity across ensemble members [21]. Randomly picking a ϕ_k^t from Φ^t would approximate the
 249 posterior sampling ($r \sim p_r(\cdot|\mathcal{D}_t)$) for the RM [43, 25]. In practice, we train K MLP heads on top of
 250 a pretrained and frozen transformer. We refer to the ensemble as the Epistemic Reward Model (ERM,
 251 denoted as $\mathcal{R}_{\Phi}$).

252 4.2.2 Policy-guided search to approximate $\arg \max$

253 With the ERM approximating the reward posterior, we need to further approximate the response
 254 selection steps (Lines 4 to 6) which generally take the form of $\arg \max_{b \in \mathcal{Y}} U(b)$, where U absorbs the
 255 sampled prompt, the sampled RM, and optionally the selected first response (for BAI, Line 6). To ob-
 256 tain the maximum, bandit algorithms for large action spaces typically resort to an action optimization
 257 oracle [31, 91], but they assume a linear structure of U with respect to b , which might be impractical
 258 for LLMs. Therefore, we instead replace the optimization over $\mathcal{Y}$ with sampling from a policy-guided
 259 distribution conditioned on U , $\pi_{\text{prior}}(\cdot|\mathbf{x}) \exp(U(\cdot)/\eta)$, which is appropriate since it favors responses
 260 $\mathbf{y}$ that approximately maximize $U(\mathbf{y})$. In practice, for a given prompt $\mathbf{x}_t$, we sample M candidate
 261 responses from the prior policy $\pi_{\text{prior}}(\cdot|\mathbf{x}_t)$ to construct a proposal set $\mathcal{S}_t = \{\mathbf{y}_i^t\}_{i=1}^M$. We then con-
 262 duct a greedy search in $\mathcal{S}_t$ (taking $\eta \rightarrow 0$) to identify the response $\mathbf{y}_t$ (or $\mathbf{y}_t'$) that locally maximizes
 263 the utility function U , which is subsequently used in the duel. We also reuse the same $\mathcal{S}_t$ for different
 264 U functions at time t to save computation. The choice of π_{prior} will be discussed in the next section.

265 4.2.3 Online policy learning from mixed preferences

266 We finally resolve two remaining questions: (Q1) how to choose a sensible π_{prior} at each time t
 267 and (Q2) how to get a good generative policy online. To this end, we propose a simple approach to
 268 approximately address both questions simultaneously. That is, we can utilize any direct optimizer to
 269 learn the policy π_{θ^t} online with the following loss and use the latest online policy as π_{prior} :

$$\mathcal{L}_{\pi}(\theta^t|\mathcal{B}_t, \pi_{\text{ref}}, F) = \mathbb{E}_{(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-) \sim p_{\mathcal{B}_t}} [F_{\theta^t}(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-, \pi_{\text{ref}})], \quad (9)$$

270 where $\mathcal{B}_t$ is a batch of preference data labeled by the oracle wherein the responses are proposed by
 271 π_{prior} and selected by $\mathcal{R}_{\Phi^t}$, F could be any DAP loss (see App. A for some examples), and π_{ref} is
 272 chosen to be π_{sft} . Note that we use π_{θ^t} as π_{prior} at any time t , thus $\mathcal{B}^t$ is a batch of on-policy data.
 273 By *contrastive training* on these *on-policy* data, we leverage their orthogonal benefits to achieve
 274 maximal policy improvement [69, 70].

275 Now that optimizing Eq. (9) yields a good online policy π_{θ^t} (answering Q2), we need to assess
 276 whether π_{θ^t} can serve as a suitable π_{prior} for approximating the $\arg \max$ in TS (Q1). If we optimize
 277 π_{θ^t} with oracle preference data, $\mathcal{S}_t$ will be biased towards responses with high oracle reward r^* .
 278 Bias towards high- r^* region is generally helpful because it aligns with $\arg \max_{b \in \mathcal{Y}} r(\mathbf{x}, b)$ that
 279 seeks high-reward responses. However, optimizing π_{θ^t} *only* with oracle data can average out the
 280 epistemic uncertainty of $\mathcal{R}$, hindering the exploration efficiency. To mitigate this issue, we further
 281 align π_{θ^t} with $\mathcal{R}_{\Phi^t}$ using the same direct optimizer to encourage π_{θ^t} to propose high- $r_{\phi_k^t}$ responses
 282 for individual $r_{\phi_k^t}$, leading to better approximation of $\arg \max_{b \in \mathcal{Y}} r(\mathbf{x}, b)$ for any sampled r . To
 283 implement, we optimize Eq. (9) over a batch of data mixture $p_{\mathcal{B}_t^{\text{mix}}} = \gamma p_{\mathcal{B}_t} + (1 - \gamma) p_{\mathcal{B}_t^{\text{ERM}}}$, where
 284 $\gamma \in [0, 1]$ controls the mixture ratio and $\mathcal{B}_t^{\text{ERM}} = \{(\mathbf{x}_i, \tilde{\mathbf{y}}_i^+, \tilde{\mathbf{y}}_i^-)\}_{i=1}^b$ consists of preference data
 285 labeled by randomly sampled individual ensemble members $r_{\phi_k^t}$. Interestingly, learning from mixed
 286 preferences further boosts sample efficiency because it utilizes the internal ERM to get pseudo labels
 287 instead of querying humans. This relates closely to model-based RL, for which we discuss further in
 288 App. C. We summarize our practical algorithm (Algorithm 2) in App. A.

289 5 Experimental setup

290 **Software.** To facilitate our empirical studies, we develop a distributed learning framework for
 291 online LLM alignment. The framework is based on an Actor-Learner-Oracle architecture, drawing

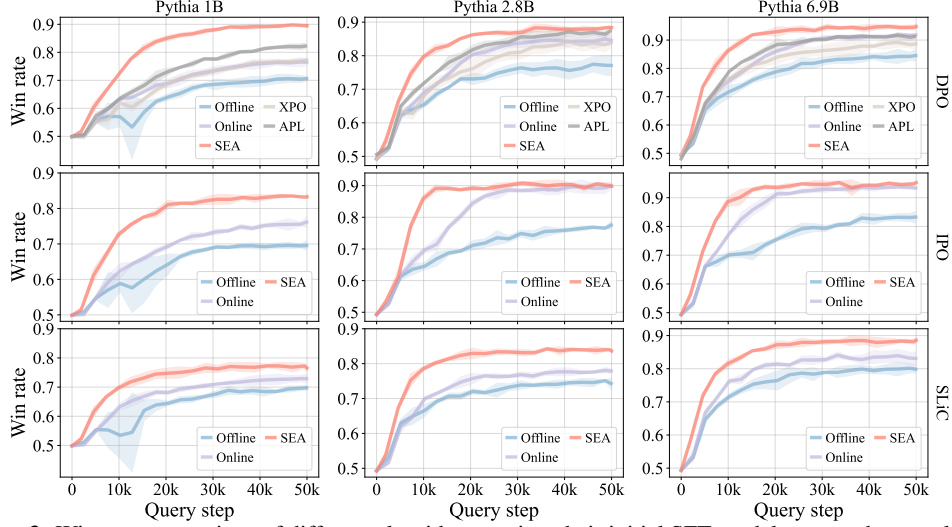


Figure 3: Win rate comparison of different algorithms against their initial SFT models across three scales and three direct optimizers.

inspiration from Espeholt et al. [22]. We incorporate various optimizations for each component: vLLM [33] for actors, DeepSpeed [58] for learners, and Mosec [83] for oracles. Detailed descriptions of the framework and its efficiency benchmarks are provided in App. D & H.

Settings. We adopt SFT models tuned on TL;DR [65] from Huang et al. [27], which cover three scales (1B, 2.8B, 6.9B) of the Pythia family [7], as starting points for our experiments. We use a strong scalar RM [40]⁴ to simulate the preference oracle. To verify the effectiveness of **SEA**, we employ three direct optimizers: DPO [55], IPO [4], and SLiC [88] to serve as F in Eq. (9). Besides, two LLM exploration methods built on DPO, APL [47] and XPO [79], are fairly compared when using DPO as the optimizer. Our experiments primarily focus on the BAI setting (crowdsourcing labeling), where we report the win rate of learned models against initial SFT models. All experiments are repeated three times to ensure statistical significance. Please see App. F for more details.

6 Empirical studies

We next present our empirical studies highlighting five results: (1) Comparisons with baselines across various direct optimizers and model scales demonstrate **SEA**’s superior *sample efficiency* (Sec. 6.1). (2) Ablations confirm that both *online policy learning* and *active exploration* contribute to sample-efficient alignment, and using the learned ERM for Best-of-N sampling further improves the performance (Sec. 6.2). (3) Different exploration strategies (Line 5 or Line 6 in Algorithm 1) are verified to work best in respective settings. (4) **SEA** robustly outperforms baselines when GPT4o-mini is used as a judge to simulate human feedback. (5) Beyond the summarization task, **SEA** can effectively enhance general capabilities of LLMs. Results for (3-5) are deferred to App. G due to space constraints.

6.1 Overall comparison

We first compare **SEA** with all baselines across three model scales and three direct optimizers. APL and XPO are only compared when DPO is used as the direct optimizer, because they are incompatible with IPO or SLiC. Fig. 3 shows the win rate curves versus the number of query steps. Across all settings, Online agents consistently improve sample efficiency over their Offline counterparts, validating the necessity of **Property 1** for alignment algorithms. Focusing on the first row, we observe that among prior active exploration methods, XPO gives a small improvement in final performance over Online (passive) at the 1B scale, but falls short for larger scales. On the other hand, APL shows a significant sample efficiency boost at the 1B scale, but this advantage diminishes when scaling up and it performs almost the same as Online at 6.9B scale. Our method, **SEA**, outperforms both offline and online passive methods across all scales and all direct optimizers, confirming the critical role that **Property 2** plays for sample-efficient alignment. Meanwhile, in the special case of using DPO as the direct

⁴<https://huggingface.co/Skywork/Skywork-Reward-Llama-3.1-8B>.

Table 1: Decomposition of different driving factors of online active alignment algorithms.

Variant	Inference (Test)	Exploration	Learn	Remark
1	π_θ	passive	π_θ	Online DAP [24]
2	π_θ	active	$(\pi_\theta, \mathcal{R}_\Phi)$	SEA without ERM sync (Sec. 4.2.3)
3	π_θ	active	$(\pi_\theta \leftrightarrow \mathcal{R}_\Phi)$	SEA
4	$\text{BoN}(\pi_\theta, \mathcal{R}_\Phi)$	passive	$(\pi_\theta, \mathcal{R}_\Phi)$	-
5	$\text{BoN}(\pi_\theta, \mathcal{R}_\Phi)$	active	$(\pi_\theta, \mathcal{R}_\Phi)$	-
6	$\text{BoN}(\pi_\theta, \mathcal{R}_\Phi)$	active	$(\pi_\theta \leftrightarrow \mathcal{R}_\Phi)$	SEA with Best-of-N sampling
7	$\text{BoN}(\pi_{\text{ref}}, \mathcal{R}_\Phi)$	active	$\mathcal{R}_\Phi$	Not learn policy [21]

optimizer, **SEA** also shows superior performance to prior online active exploration methods including APL and XPO. We invite readers to revisit Fig. 1, where we show that **SEA** not only attains significantly improved final performance (Top) but also achieves $2\text{-}5\times$ better sample efficiency (Bottom).

Additionally, we note that the choice of direct optimizer is crucial for both online learning and active exploration. When comparing different optimizers at the 1B scale (the first column), all Offline agents demonstrate comparable learning efficiency and reach the same level of final performance (around 70% win rate), but SLiC Online agent deliver slightly less improvement than DPO and IPO Online agents. Besides, when incorporating active exploration, the **SEA** agent using DPO shows much larger improvement than the other two. This suggests that selecting the most suitable policy optimizer coupled with active exploration would yield the best agent.

6.2 Ablation analysis

We decompose **SEA** into distinct components to evaluate their individual contributions. Table 1 shows the three axes we dissect **SEA** on, including inference methods, exploration strategies, and learning components. We construct seven agent variants from different combinations, which cover two closely related baselines [21, 24]. We show in Fig. 4 the performance curves of each variant, all trained with DPO on 1B scale.

The top plot compares variants that directly use the policy for inference. Comparing with the vanilla online method (Variant-1), we observe learning ERM for active exploration (Variant-2) is beneficial, and aligning π_{θ^t} with $\mathcal{R}_{\Phi^t}$ (Variant-3) further improves sample efficiency, which validate our algorithm. Additionally, since a reward model is learned within the agent, we can incorporate inference-time alignment via Best-of-N (BoN) sampling [48, 72]. This also facilitates a direct comparison between **SEA** and Dwaracherla et al. [21], which learns a similar ERM for both exploration and BoN but does not align the LLM policy. Results in the bottom plot of Fig. 4 suggest a similar trend that Variant-6 $\succ$ Variant-5 $\succ$ Variant-4. The Variant-7 [21], however, ceases to improve after ERM converges due to the limited capability of its fixed policy.

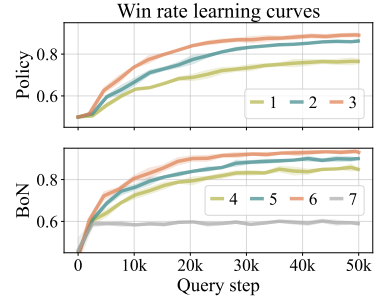


Figure 4: Win rate comparison of different agent variants when using (Left) policy and (Right) Best-of-N sampling for inference.

7 Conclusion

In this paper, we study the problem of LLM alignment through the lens of contextual dueling bandits and propose a Thompson sampling-based algorithm to achieve sample-efficient alignment. We incorporate three techniques, including epistemic reward model, policy-guided search and mixed preference learning to yield a practically efficient online alignment method. Extensive empirical evaluation demonstrates the superior sample efficiency of our method compared to existing baselines. To our knowledge, this is the first work to study active exploration for online LLM alignment with fully online experimental verification. We hope our positive empirical results, along with the open-sourced codebase, will encourage future research in this direction, ultimately enabling LLMs to achieve superhuman intelligence with an affordable amount of human feedback.

References

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [2] Jean-Yves Audibert and Sébastien Bubeck. Best arm identification in multi-armed bandits. In *Conference on learning theory*, pages 41–53, 2010.
- [3] Peter Auer, Nicolò Cesa-Bianchi, and Paul Fischer. Finite-time analysis of the multiarmed bandit problem. *Machine Learning*, 47:235–256, 2002.
- [4] Mohammad Gheshlaghi Azar, Zhaohan Daniel Guo, Bilal Piot, Remi Munos, Mark Rowland, Michal Valko, and Daniele Calandriello. A general theoretical paradigm to understand learning from human preferences. In *International Conference on Artificial Intelligence and Statistics*, pages 4447–4455. PMLR, 2024.
- [5] 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.
- [6] Christopher Berner, Greg Brockman, Brooke Chan, Vicki Cheung, Przemysław Dębiak, Christy Dennison, David Farhi, Quirin Fischer, Shariq Hashme, Chris Hesse, et al. Dota 2 with large scale deep reinforcement learning. *arXiv preprint arXiv:1912.06680*, 2019.
- [7] Stella Biderman, Hailey Schoelkopf, Quentin Gregory Anthony, Herbie Bradley, Kyle O’Brien, Eric Hallahan, Mohammad Aflah Khan, Shivanshu Purohit, USVSN Sai Prashanth, Edward Raff, et al. Pythia: A suite for analyzing large language models across training and scaling. In *International Conference on Machine Learning*, pages 2397–2430. PMLR, 2023.
- [8] Ralph Allan Bradley and Milton E. Terry. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952.
- [9] Sébastien Bubeck, Rémi Munos, and Gilles Stoltz. Pure exploration in multi-armed bandits problems. In *Algorithmic Learning Theory: 20th International Conference, ALT 2009, Porto, Portugal, October 3-5, 2009. Proceedings 20*, pages 23–37. Springer, 2009.
- [10] Sébastien Bubeck, Varun Chandrasekaran, Ronen Eldan, Johannes Gehrke, Eric Horvitz, Ece Kamar, Peter Lee, Yin Tat Lee, Yuanzhi Li, Scott Lundberg, et al. Sparks of artificial general intelligence: Early experiments with gpt-4. *arXiv preprint arXiv:2303.12712*, 2023.
- [11] Róbert Busa-Fekete, Balázs Szörényi, Paul Weng, Weiwei Cheng, and Eyke Hüllermeier. Preference-based reinforcement learning: evolutionary direct policy search using a preference-based racing algorithm. *Machine learning*, 97:327–351, 2014.
- [12] Olivier Chapelle and Lihong Li. An empirical evaluation of thompson sampling. *Advances in neural information processing systems*, 24, 2011.
- [13] OpenAI ChatGPT. ChatGPT. <https://chatgpt.com/>, 2024. Accessed: 2024-09-30.
- [14] Changyu Chen, Zichen Liu, Chao Du, Tianyu Pang, Qian Liu, Arunesh Sinha, Pradeep Varakantham, and Min Lin. Bootstrapping language models with dpo implicit rewards. *arXiv preprint arXiv:2406.09760*, 2024.
- [15] Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. Deep reinforcement learning from human preferences. *Advances in neural information processing systems*, 30, 2017.
- [16] Ganqu Cui, Lifan Yuan, Ning Ding, Guanming Yao, Wei Zhu, Yuan Ni, Guotong Xie, Zhiyuan Liu, and Maosong Sun. Ultrafeedback: Boosting language models with high-quality feedback, 2023.

- [17] Nirjhar Das, Souradip Chakraborty, Aldo Pacchiano, and Sayak Ray Chowdhury. Provably sample efficient rlhf via active preference optimization. *arXiv preprint arXiv:2402.10500*, 2024.
- [18] 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. *arXiv preprint arXiv:2405.07863*, 2024.
- [19] Yann Dubois, Balázs Galambosi, Percy Liang, and Tatsunori B Hashimoto. Length-controlled alpacaeval: A simple way to debias automatic evaluators. *arXiv preprint arXiv:2404.04475*, 2024.
- [20] Miroslav Dudík, Katja Hofmann, Robert E Schapire, Aleksandrs Slivkins, and Masrour Zoghi. Contextual dueling bandits. In *Conference on Learning Theory*, pages 563–587. PMLR, 2015.
- [21] Vikranth Dwaracherla, Seyed Mohammad Asghari, Botao Hao, and Benjamin Van Roy. Efficient exploration for llms. In *International Conference on Machine Learning*, 2024.
- [22] Lasse Espeholt, Hubert Soyer, Remi Munos, Karen Simonyan, Vlad Mnih, Tom Ward, Yotam Doron, Vlad Firoiu, Tim Harley, Iain Dunning, et al. Impala: Scalable distributed deep-rl with importance weighted actor-learner architectures. In *International conference on machine learning*, pages 1407–1416. PMLR, 2018.
- [23] Javier González, Zhenwen Dai, Andreas Damianou, and Neil D Lawrence. Preferential bayesian optimization. In *International Conference on Machine Learning*, pages 1282–1291. PMLR, 2017.
- [24] 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.
- [25] Fredrik K Gustafsson, Martin Danelljan, and Thomas B Schon. Evaluating scalable bayesian deep learning methods for robust computer vision. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition workshops*, pages 318–319, 2020.
- [26] Jian Hu, Xibin Wu, Weixun Wang, Dehao Zhang, Yu Cao, et al. Openrlhf: An easy-to-use, scalable and high-performance rlhf framework. *arXiv preprint arXiv:2405.11143*, 2024.
- [27] Shengyi Huang, Michael Noukhovitch, Arian Hosseini, Kashif Rasul, Weixun Wang, and Lewis Tunstall. The n+ implementation details of rlhf with ppo: A case study on tl; dr summarization. *arXiv preprint arXiv:2403.17031*, 2024.
- [28] Michael Janner, Justin Fu, Marvin Zhang, and Sergey Levine. When to trust your model: Model-based policy optimization. *Advances in neural information processing systems*, 32, 2019.
- [29] Natasha Jaques, Judy Hanwen Shen, Asma Ghandeharioun, Craig Ferguson, Agata Lapedriza, Noah Jones, Shixiang Shane Gu, and Rosalind Picard. Human-centric dialog training via offline reinforcement learning. *arXiv preprint arXiv:2010.05848*, 2020.
- [30] Dongfu Jiang, Xiang Ren, and Bill Yuchen Lin. Llm-blender: Ensembling large language models with pairwise comparison and generative fusion. In *Proceedings of the 61th Annual Meeting of the Association for Computational Linguistics (ACL 2023)*, 2023.
- [31] Julian Katz-Samuels, Lalit Jain, Kevin G Jamieson, et al. An empirical process approach to the union bound: Practical algorithms for combinatorial and linear bandits. *Advances in Neural Information Processing Systems*, 33:10371–10382, 2020.
- [32] Rahul Kidambi, Aravind Rajeswaran, Praneeth Netrapalli, and Thorsten Joachims. Morel: Model-based offline reinforcement learning. *Advances in neural information processing systems*, 33:21810–21823, 2020.
- [33] Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the ACM SIGOPS 29th Symposium on Operating Systems Principles*, 2023.

- [34] Balaji Lakshminarayanan, Alexander Pritzel, and Charles Blundell. Simple and scalable predictive uncertainty estimation using deep ensembles. *Advances in neural information processing systems*, 30, 2017.
- [35] 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, 2024.
- [36] Sascha Lange, Thomas Gabel, and Martin Riedmiller. Batch reinforcement learning. In *Reinforcement learning: State-of-the-art*, pages 45–73. Springer, 2012.
- [37] Sergey Levine, Aviral Kumar, George Tucker, and Justin Fu. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *arXiv preprint arXiv:2005.01643*, 2020.
- [38] Xuechen Li, Tianyi Zhang, Yann Dubois, Rohan Taori, Ishaan Gulrajani, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. AlpacaEval: An automatic evaluator of instruction-following models. https://github.com/tatsu-lab/alpaca_eval, 5 2023.
- [39] Xuheng Li, Heyang Zhao, and Quanquan Gu. Feel-good thompson sampling for contextual dueling bandits. *arXiv preprint arXiv:2404.06013*, 2024.
- [40] Chris Yuhao Liu, Liang Zeng, Liu Jiakai, Rui Yan, Jujie He, Chaojie Wang, Shuicheng Yan, Yang Liu, and Yahui Zhou. Skywork reward model series. *arXiv preprint arXiv:2410.18451*, 2024.
- [41] Zichen Liu, Siyi Li, Wee Sun Lee, Shuicheng Yan, and Zhongwen Xu. Efficient offline policy optimization with a learned model. In *International Conference on Learning Representations*, 2023.
- [42] Zichen Liu, Chao Du, Wee Sun Lee, and Min Lin. Locality sensitive sparse encoding for learning world models online. In *International Conference on Learning Representations*, 2024.
- [43] Xiuyuan Lu and Benjamin Van Roy. Ensemble sampling. *Advances in neural information processing systems*, 30, 2017.
- [44] 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.
- [45] Luckeciano C Melo, Panagiotis Tigas, Alessandro Abate, and Yarin Gal. Deep bayesian active learning for preference modeling in large language models. *arXiv preprint arXiv:2406.10023*, 2024.
- [46] Yu Meng, Mengzhou Xia, and Danqi Chen. Simpo: Simple preference optimization with a reference-free reward. *arXiv preprint arXiv:2405.14734*, 2024.
- [47] William Muldrew, Peter Hayes, Mingtian Zhang, and David Barber. Active preference learning for large language models. In *International Conference on Machine Learning*, 2024.
- [48] Reiichiro Nakano, Jacob Hilton, Suchir Balaji, Jeff Wu, Long Ouyang, Christina Kim, Christopher Hesse, Shantanu Jain, Vineet Kosaraju, William Saunders, et al. Webgpt: Browser-assisted question-answering with human feedback. *arXiv preprint arXiv:2112.09332*, 2021.
- [49] Andrew Y Ng and Stuart Russell. Algorithms for inverse reinforcement learning. In *International Conference on Machine Learning*, volume 1, page 2, 2000.
- [50] 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.
- [51] Xue Bin Peng, Aviral Kumar, Grace Zhang, and Sergey Levine. Advantage-weighted regression: Simple and scalable off-policy reinforcement learning. *arXiv preprint arXiv:1910.00177*, 2019.

- [52] Jan Peters and Stefan Schaal. Reinforcement learning by reward-weighted regression for operational space control. In *International Conference on Machine Learning*, pages 745–750, 2007.
- [53] Moritz Philipp and Nishihara Robert. Plasma: A high-performance shared-memory object store, 2017. URL <https://arrow.apache.org/blog/2017/08/08/plasma-in-memory-object-store/>.
- [54] Chao Qin, Zheng Wen, Xiuyuan Lu, and Benjamin Van Roy. An analysis of ensemble sampling. *Advances in Neural Information Processing Systems*, 35:21602–21614, 2022.
- [55] 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*, 37, 2023.
- [56] Rafael Rafailov, Joey Hejna, Ryan Park, and Chelsea Finn. From r to q*: Your language model is secretly a q-function. In *Conference on Language Modeling*, 2024.
- [57] Samyam Rajbhandari, Jeff Rasley, Olatunji Ruwase, and Yuxiong He. Zero: Memory optimizations toward training trillion parameter models. In *SC20: International Conference for High Performance Computing, Networking, Storage and Analysis*, pages 1–16. IEEE, 2020.
- [58] Jeff Rasley, Samyam Rajbhandari, Olatunji Ruwase, and Yuxiong He. Deepspeed: System optimizations enable training deep learning models with over 100 billion parameters. In *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 3505–3506, 2020.
- [59] Herbert Robbins. Some aspects of the sequential design of experiments. *Bulletin of the American Mathematics Society*, 58:527–535, 1952.
- [60] Daniel Russo. Simple bayesian algorithms for best arm identification. In *Conference on Learning Theory*, pages 1417–1418. PMLR, 2016.
- [61] Daniel J Russo, Benjamin Van Roy, Abbas Kazerouni, Ian Osband, Zheng Wen, et al. A tutorial on thompson sampling. *Foundations and Trends® in Machine Learning*, 11(1):1–96, 2018.
- [62] Aadirupa Saha. Optimal algorithms for stochastic contextual preference bandits. *Advances in Neural Information Processing Systems*, 34:30050–30062, 2021.
- [63] Julian Schrittwieser, Thomas Hubert, Amol Mandhane, Mohammadamin Barekatain, Ioannis Antonoglou, and David Silver. Online and offline reinforcement learning by planning with a learned model. *Advances in Neural Information Processing Systems*, 34:27580–27591, 2021.
- [64] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- [65] Nisan Stiennon, Long Ouyang, Jeffrey Wu, Daniel Ziegler, Ryan Lowe, Chelsea Voss, Alec Radford, Dario Amodei, and Paul F Christiano. Learning to summarize with human feedback. *Advances in Neural Information Processing Systems*, 33:3008–3021, 2020.
- [66] Richard S. Sutton. Integrated architectures for learning, planning, and reacting based on approximating dynamic programming. In *Machine Learning Proceedings*, pages 216–224. Morgan Kaufmann, 1990.
- [67] Richard S. Sutton and Andrew G. Barto. *Reinforcement Learning: An Introduction*. The MIT Press, second edition, 2018.
- [68] Richard S Sutton, Michael Bowling, and Patrick M Pilarski. The alberta plan for ai research. *arXiv preprint arXiv:2208.11173*, 2022.
- [69] Fahim Tajwar, Anikait Singh, Archit Sharma, Rafael Rafailov, Jeff Schneider, Tengyang Xie, Stefano Ermon, Chelsea Finn, and Aviral Kumar. Preference fine-tuning of llms should leverage suboptimal, on-policy data. *arXiv preprint arXiv:2404.14367*, 2024.

- [70] Yunhao Tang, Daniel Zhaohan Guo, Zeyu Zheng, Daniele Calandriello, Yuan Cao, Eugene Tarassov, Rémi Munos, Bernardo Ávila Pires, Michal Valko, Yong Cheng, et al. Understanding the performance gap between online and offline alignment algorithms. *arXiv preprint arXiv:2405.08448*, 2024.
- [71] William R Thompson. On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika*, 25(3-4):285–294, 1933.
- [72] 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.
- [73] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Advances in Neural Information Processing Systems*, volume 30, 2017.
- [74] Haoxiang Wang, Wei Xiong, Tengyang Xie, Han Zhao, and Tong Zhang. Interpretable preferences via multi-objective reward modeling and mixture-of-experts. *arXiv preprint arXiv:2406.12845*, 2024.
- [75] Jiayi Weng, Min Lin, Shengyi Huang, Bo Liu, Denys Makoviichuk, Viktor Makoviyshuk, Zichen Liu, Yufan Song, Ting Luo, Yukun Jiang, Zhongwen Xu, and Shuicheng Yan. EnvPool: A highly parallel reinforcement learning environment execution engine. In *Advances in Neural Information Processing Systems*, volume 35, pages 22409–22421, 2022.
- [76] Ronald J Williams. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine learning*, 8:229–256, 1992.
- [77] Christian Wirth, Riad Akrou, Gerhard Neumann, and Johannes Fürnkranz. A survey of preference-based reinforcement learning methods. *Journal of Machine Learning Research*, 18(136):1–46, 2017.
- [78] Huasen Wu and Xin Liu. Double thompson sampling for dueling bandits. *Advances in neural information processing systems*, 29, 2016.
- [79] 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.
- [80] 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.
- [81] 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.
- [82] Fan Yang, Gabriel Barth-Maron, Piotr Stańczyk, Matthew Hoffman, Siqi Liu, Manuel Kroiss, Aedan Pope, and Alban Rustemi. Launchpad: A programming model for distributed machine learning research. *arXiv preprint arXiv:2106.04516*, 2021.
- [83] Keming Yang, Zichen Liu, and Philip Cheng. MOSEC: Model Serving made Efficient in the Cloud. <https://github.com/mosecorg/mosec>, 2021.
- [84] Tianhe Yu, Aviral Kumar, Rafael Rafailov, Aravind Rajeswaran, Sergey Levine, and Chelsea Finn. Combo: Conservative offline model-based policy optimization. *Advances in neural information processing systems*, 34:28954–28967, 2021.
- [85] Yisong Yue, Josef Broder, Robert Kleinberg, and Thorsten Joachims. The k-armed dueling bandits problem. *Journal of Computer and System Sciences*, 78(5):1538–1556, 2012.
- [86] Shenao Zhang, Donghan Yu, Hiteshi Sharma, 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.

- 599 [87] Xuan Zhang, Chao Du, Tianyu Pang, Qian Liu, Wei Gao, and Min Lin. Chain of preference
600 optimization: Improving chain-of-thought reasoning in llms. *Advances in Neural Information*
601 *Processing Systems*, 38, 2024.
- 602 [88] Yao Zhao, Rishabh Joshi, Tianqi Liu, Misha Khalman, Mohammad Saleh, and Peter J Liu. Slic-
603 hf: Sequence likelihood calibration with human feedback. *arXiv preprint arXiv:2305.10425*,
604 2023.
- 605 [89] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang,
606 Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. Judging llm-as-a-judge with mt-bench and
607 chatbot arena. *Advances in Neural Information Processing Systems*, 36:46595–46623, 2023.
- 608 [90] Banghua Zhu, Michael Jordan, and Jiantao Jiao. Principled reinforcement learning with human
609 feedback from pairwise or k-wise comparisons. In *Proceedings of the 40th International*
610 *Conference on Machine Learning*, pages 43037–43067. PMLR, 2023.
- 611 [91] Yinglun Zhu, Dylan J Foster, John Langford, and Paul Mineiro. Contextual bandits with
612 large action spaces: Made practical. In *International Conference on Machine Learning*, pages
613 27428–27453. PMLR, 2022.

614 A Algorithm details

615 While Algorithm 1 presents our Thompson sampling algorithm for LLM alignment, it is intractable
 616 and centered around the reward posterior modeling. We next present a practical sample-efficient
 617 alignment agent that learns both an LLM policy and an epistemic reward model (ERM) online.

Algorithm 2 Sample-efficient alignment (SEA) for LLMs

Input: Reference policy π_{ref} , DAP loss function F , prompt distribution $p_{\mathcal{X}}$, unknown but queryable preference oracle $\mathbb{P}$, mixture ratio γ .

```

1: Initialize experience  $\mathcal{D}_0 \leftarrow \emptyset$ , policy  $\pi_{\theta^0} \leftarrow \pi_{\text{ref}}$ , and ERM weights  $\Phi^0 = \{\phi_k^0\}_{k=1}^K$  randomly.
2: for  $t = 1, \dots, T$  do
3:   Receive a prompt  $\mathbf{x}_t \sim p_{\mathcal{X}}$ .
4:   Sample  $M$  responses  $\mathbf{y}_t^i \sim \pi_{\theta^{t-1}}(\cdot|\mathbf{x}_t)$  to construct  $\mathcal{S}_t = \{\mathbf{y}_t^i\}_{i=1}^M$ .
5:   Sample  $\phi \sim \text{Uniform}(\Phi^{t-1})$  and set  $\mathbf{y}_t \leftarrow \arg \max_{\mathbf{b} \in \mathcal{S}_t} r_{\phi}(\mathbf{x}_t, \mathbf{b})$ . // Select 1st response  $\mathbf{y}$ .
// E&E objective: aligning an online system.
6:   repeat
     Sample  $\phi \sim \text{Uniform}(\Phi^{t-1})$  and set  $\mathbf{y}'_t \leftarrow \arg \max_{\mathbf{b} \in \mathcal{S}_t} r_{\phi}(\mathbf{x}_t, \mathbf{b})$ . // Select 2nd response  $\mathbf{y}'$ .
     until  $\mathbf{y}'_t \neq \mathbf{y}_t$ 
// BAI objective: labeling via crowdsourcing.
7:   Set  $\mathbf{y}'_t \leftarrow \arg \max_{\mathbf{b} \in \mathcal{S}_t} \mathbb{V}_{\phi} [\sigma(r_{\phi}(\mathbf{x}_t, \mathbf{y}_t) - r_{\phi}(\mathbf{x}_t, \mathbf{b}))]$ . // OR select 2nd response  $\mathbf{y}'$ .
     where  $\mathbb{V}_{\phi}[\cdot]$  computes variance across ensemble members of  $\Phi^{t-1}$ .
8:   if  $g < \gamma$  for  $g \sim \text{Uniform}(0, 1)$  then
     Label  $\{\mathbf{y}_t, \mathbf{y}'_t\}$  with  $\mathbb{P}$  to obtain  $\mathcal{B}_t = \{(\mathbf{x}_t, \mathbf{y}_t^+, \mathbf{y}_t^-)\}$  and update experience  $\mathcal{D}_t \leftarrow \mathcal{D}_{t-1} \cup \mathcal{B}_t$ .
     else
     Use  $\mathcal{R}_{\Phi^{t-1}}$  to get synthetic labels and obtain  $\mathcal{B}_t = \{(\mathbf{x}_t, \tilde{\mathbf{y}}_t^+, \tilde{\mathbf{y}}_t^-)\}$ .
     end if
9:   Update ERM with the regularized NLL loss (Eq. (8)):

$$\Phi^t \leftarrow \Phi^{t-1} - \alpha_{\mathcal{R}} \nabla_{\Phi} \mathcal{L}_{\mathcal{R}}(\Phi^{t-1} | \mathcal{D}_t).$$

// Reward learning.
10:  Update policy with the direct optimizer (Eq. (9)):

$$\theta^t \leftarrow \theta^{t-1} - \alpha_{\pi} \nabla_{\theta} \mathcal{L}_{\pi}(\theta^{t-1} | \mathcal{B}_t, \pi_{\text{ref}}, F).$$

// Policy learning.
11: end for

```

618 In Algorithm 2, we describe an online setting where a single example is processed at each time t
 619 (batch size $b = 1$). This is mainly for notational convenience, while in implementation we set b to
 620 be the training batch size (e.g., 128). We instantiate the reward posterior with an epistemic reward
 621 model, which allows for efficient incremental update and sampling. We also replace the global
 622 optimization ($\arg \max_{\mathbf{b} \in \mathcal{Y}}$) with a policy-guided local search among proposals sampled from the
 623 latest online policy $\pi_{\theta^{t-1}}$. At each time t , we update ERM weights Φ with m gradient steps with
 624 randomly sampled batches from the experience $\mathcal{D}_t$. We find setting $m = 5$ suffices to achieve a
 625 reasonable accuracy. The policy parameters θ are updated using mixed preference data, with a
 626 γ proportion being the real environment experience and the remaining $(1 - \gamma)$ from the ERM's
 627 synthetic experience. Note that the synthetic experience is not added into $\mathcal{D}_t$ to ensure reward
 628 learning always uses ground truth environment data.

629 We consider the following three direct optimizers in our experiments:

630 • DPO [55]:

$$F_{\theta}(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-, \pi_{\text{ref}}) = -\log \sigma \left(\beta \log \frac{\pi_{\theta}(\mathbf{y}^+ | \mathbf{x}) \pi_{\text{ref}}(\mathbf{y}^- | \mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}^+ | \mathbf{x}) \pi_{\theta}(\mathbf{y}^- | \mathbf{x})} \right) \quad (10)$$

631 • IPO [4]:

$$F_{\theta}(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-, \pi_{\text{ref}}) = \left(\log \left(\frac{\pi_{\theta}(\mathbf{y}^+ | \mathbf{x}) \pi_{\text{ref}}(\mathbf{y}^- | \mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}^+ | \mathbf{x}) \pi_{\theta}(\mathbf{y}^- | \mathbf{x})} \right) - \frac{1}{2\beta} \right)^2 \quad (11)$$

632 • SLiC [88]:

$$F_{\theta}(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-, \pi_{\text{ref}}) = \max \left(0, 1 - \beta \log \frac{\pi_{\theta}(\mathbf{y}^+ | \mathbf{x}) \pi_{\text{ref}}(\mathbf{y}^- | \mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}^+ | \mathbf{x}) \pi_{\theta}(\mathbf{y}^- | \mathbf{x})} \right) \quad (12)$$

where β controls the rate of deviation of π_θ from π_{ref} .

B Full related works

RLHF and DAP. Commonly adopted RLHF pipelines [15, 65, 5, 50] first learn a proxy RM with a negative log-likelihood loss:

$$\mathcal{L}_r(\phi|\mathcal{D}) = -\mathbb{E}_{(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-) \sim \mathcal{D}} [\log \sigma(r_\phi(\mathbf{x}, \mathbf{y}^+) - r_\phi(\mathbf{x}, \mathbf{y}^-))], \quad (13)$$

where $\mathcal{D}$ is collected by querying human annotators using a behavior policy π_{ref} (typically the supervised fine-tuned policy π_{sft}). Afterwards, *offline RL*⁵ [36, 37] is conducted to learn π_θ with respect to the learned reward r_ϕ internally within the agent (Fig. 2a). However, the learned model π_θ might be inaccurate at regions out of the distribution (o.o.d.) of π_{ref} because little training data can be collected. An effective remedy is to incorporate a pessimistic term to combat the distributional shift, leading to a reformulation of the von Neumann winner policy objective in Eq. (4) as

$$J(\pi_\theta) = \mathbb{E}_{\mathbf{x} \sim p_{\mathcal{X}}} \mathbb{E}_{\mathbf{y} \sim \pi_\theta(\cdot|\mathbf{x})} \left[\underbrace{r_\phi(\mathbf{x}, \mathbf{y})}_{\text{estimated } r^*} - \beta \underbrace{\log \frac{\pi_\theta(\mathbf{y}|\mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}|\mathbf{x})}}_{\text{o.o.d. reward penalty}} \right] \quad (14)$$

$$= \mathbb{E}_{\mathbf{x} \sim p_{\mathcal{X}}} \left[\mathbb{E}_{\mathbf{y} \sim \pi_\theta(\cdot|\mathbf{x})} [r_\phi(\mathbf{x}, \mathbf{y})] - \beta D_{\text{KL}}(\pi_\theta(\cdot|\mathbf{x}) || \pi_{\text{ref}}(\cdot|\mathbf{x})) \right] \quad (15)$$

which converts an online objective regarding the human’s implicit reward r^* to an offline objective regarding the proxy reward r_ϕ . The KL penalty in Eq. (15) is widely used for language model fine-tuning [29, 80], and PPO [64] has become a default *RL optimizer* to maximize the KL-regularized reward. However, the performance of RLHF is guaranteed only if the preference data $\mathcal{D}$ induced by π_{ref} adequately covers π^* [90], which is often approximated by updating π_{ref} with the latest (improved) π_θ for re-sampling a batch of online experience and repeating Eq. (13) and (15). Prior works typically focus on offline or iterative online (with only a few iterations) settings [80, 18], which may compromise sample efficiency (Property 1).

True online RLHF is difficult due to the complexity and instability of RL optimizers. For example, Huang et al. [27] openly reproduces offline RLHF scaling behaviors but requires many implementation tricks for training, highlighting the difficulties of an online counterpart. Fortunately, the introduction of DAP (or *direct optimizers*) largely simplifies and stabilizes fine-tuning by conducting contrastive supervised learning directly on $\mathcal{D}$ (Fig. 2b). While most DAP works focus on learning from a fixed offline preference dataset (, including Zhao et al. [88], Rafailov et al. [55], Azar et al. [4], Meng et al. [46], Zhang et al. [87]), iterative DPO [81] observes improved results when allowing iterative online interaction. Guo et al. [24] further propose OAIF to make DAP faithfully online, satisfying Property 1, and demonstrate that online learning prevents over-fitting and yields continual performance improvement. Nevertheless, it still employs passive exploration strategies (using $\mathbf{y}, \mathbf{y}' \sim \pi_\theta$), hindering sample efficiency (Property 2).

Online exploration in LLMs. A line of recent works [44, 17, 45, 21] adopts the fully online bandit formulation and incorporates *active exploration* with uncertainty-aware RMs for response selection (Fig. 2c). In particular, Mehta et al. [44] consider the E&E setting and develop a UCB-style [3] algorithm; Das et al. [17] instead select the dueling responses with the most uncertain preference estimate, targeting the BAI setting in a pure exploration way; unlike the above, Melo et al. [45] view the problem from the angle of pool-based active learning and propose an acquisition function based on both entropy and epistemic uncertainty; finally, the work by Dwaracherla et al. [21] is the closest to ours in the sense that they apply double Thompson sampling (DTS) [78] for exploration, but DTS is designed for the E&E setting while they evaluate anytime average performance as in the BAI setting. We will show in App. G.1 that pure exploration by Das et al. [17] is not the best choice for BAI, and the objective mismatch in Dwaracherla et al. [21] could lead to suboptimal performance in respective settings. Meanwhile, all these works primarily focus on learning uncertainty-aware RMs online without updating LLM policies. Therefore, all responses are sampled from a fixed proposal policy π_β (or even a fixed dataset), making the data coverage a critical concern.

Another line of research updates LLMs online while incorporating exploration. Zhang et al. [86] and Xie et al. [79] independently propose to learn an optimistic RM to encourage exploration. They

⁵Offline in the sense that π_θ is not directly learned from online human feedback. See App. C for details.

leverage the property of DPO [55] to reparameterize RM with policy and conclude with an extra optimistic term in the DPO loss function. Thus, their learning processes are like Fig. 2b but with an optimistic direct optimizer. Muldrew et al. [47] adopt the vanilla DPO loss but utilize the implicit reward margin to actively select dueling responses. Yet, these methods are tightly coupled with DPO and not compatible to other direct optimizers. Their experiments are also limited to a few online iterations, possibly due to the implementation difficulty of a faithfully online learning system. Given their relevance to our approach, we reproduce them in a fully online manner for fair comparisons in Sec. 6.1. We summarize prior works in Table 2.

Method	Exploration		Interaction			Proposal Policy	
	Active	Passive	Online	Iterative	Offline	π_θ	π_β
RL Optimizer	[15]	✓		✓	✓	✓	
	[65]	✓		✓	✓	✓	
	[5]	✓		✓	✓	✓	
	[50]	✓		✓	✓	✓	
Direct Optimizer	[88]	✓			✓	✓	
	[55]	✓			✓	✓	
	[4]	✓			✓	✓	
	[46]	✓			✓	✓	
	[81]	✓		✓		✓	
	[24]	✓	✓			✓	
	[44]	✓	✓				✓
	[17]	✓	✓				✓
	[45]	✓	✓				✓
	[21]	✓	✓				✓
	[86]	✓	✓			✓	
	[79]	✓	✓			✓	
	[47]	✓	✓			✓	

Table 2: A summary of prior work. π_θ denotes the proposal policy that is continuously updated based on newly collected preference data, while π_β denotes a fixed proposal policy. Algorithms that encompass online interaction (Property 1), active exploration (Property 2), and learnable π_θ offer the best sample efficiency. Notably, only three methods (listed at the bottom of the table) satisfy these characteristics, and we include them for comparisons in our experiments.

C On connections with single-step RL

By viewing contextual dueling bandits as *single-step* preference-based RL (PbRL) [11, 77] problems, we can interpret paradigms shown in Fig. 2 from the RL perspective.

RLHF approaches (Fig. 2a) are instances of **offline model-based RL** [32, 84, 63, 41, 69], where they learn a reward model (no need for a transition model since the prompt-response interaction is single-step) of the environment from a batch of offline collected data, and train a policy (i.e., LLM) to maximize the return (i.e., expected one-step reward) with respect to the *learned* reward.

In contrast, DAP methods (Fig. 2b) are similar to **policy-based model-free RL** algorithms, e.g., REINFORCE [76] which conducts policy gradient update:

$$\mathbb{E}_{\mathbf{x} \sim \mathcal{X}} \mathbb{E}_{\mathbf{y} \sim \pi_\theta(\cdot | \mathbf{x})} [R(\mathbf{x}, \mathbf{y}) \nabla_\theta \log \pi_\theta(\mathbf{y} | \mathbf{x})], \quad (16)$$

where $R(\mathbf{x}, \mathbf{y})$ is the return (i.e., cumulative reward) of the trajectory. To connect with DAP, we could set R as arbitrary scalar values based on the binary preference outcomes, e.g., $R(\mathbf{x}, \mathbf{y}^+) = \zeta$ and $R(\mathbf{x}, \mathbf{y}^-) = -\zeta$ for preference triplet $\{\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-\}$. In this way we could rewrite Eq. (16) as

$$\mathbb{E}_{\mathbf{x} \sim \mathcal{X}} \mathbb{E}_{\mathbf{y}, \mathbf{y}' \sim \pi_\theta(\cdot | \mathbf{x})} \mathbb{E}_{(\mathbf{y}^+ \succ \mathbf{y}^-) \sim \mathbb{P}} [\zeta (\nabla_\theta \log \pi_\theta(\mathbf{y}^+ | \mathbf{x}) - \nabla_\theta \log \pi_\theta(\mathbf{y}^- | \mathbf{x}))], \quad (17)$$

by repeating action sampling twice and querying the oracle for preference labeling. This matches the gradient direction of contrastive DAP losses (e.g., see Section 4 of DPO [55]) if we optimize them online [24].

Additionally, active reward learning from behavior policy’s data distribution (Fig. 2c) can be regarded as **inverse RL** [49], which tries to recover environment’s reward function given expert trajectories. In the context of LLM alignment, the preference data $\{\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-\}_{i=1}^N$ directly encodes human’s implicit reward r^* , which can be inversely learned with assumptions such as the BT model [8]. However, existing methods belonging to this paradigm mostly rely on a fixed (and suboptimal) behavior policy for response sampling, whose coverage inherently limits the quality of the recovered reward function.

Last but not least, **SEA** depicted in Fig. 2d resembles a class of **online model-based RL** algorithms, known as Dyna [66, 28], that learns a *world model* from environment experience and trains a base agent (consisting of reactive policies and value functions) from both environment experience and model experience. Compared to model-free methods, Dyna naturally enables more sample-efficient learning by planning with the learned world model to update the base agent. In **SEA**, we learn the reward model online and update the LLM (i.e., the reactive policy) with model-planing experience by mixed preference learning (Sec. 4.2.3). Online model-based RL algorithms could suffer from catastrophic forgetting in the face of nonstationary data [42], and we leave it for future work. Overall, this model-based RL formulation is powerful and explains popular LLM techniques, e.g., Best-of-N sampling [72] can be viewed as planning for acting, which trades compute for performance. We believe it is a promising path leading us to unlock superhuman capabilities of LLMs.

D Distributed learning framework

The interactive nature of LLM alignment necessitates an integrated online learning system that simulates the interface. The absence of a performant open-source online alignment system has restricted many existing works to only a few iterations of batch learning [47, 18, 14, 86, 79], which creates a mismatch with their theories that typically require a large number of online interaction rounds. Even worse, such absence also makes the comparison between different LLM exploration methods difficult, often restricting evaluations to the simplest iterative DAP baselines [86, 79].

To fill this gap, we build a highly efficient learning system for experimenting with online LLM alignment algorithms. We notice that the computational bottleneck lies in online response sampling (i.e., autoregressive generation) and preference labeling (e.g., human, large RMs, or large LLMs), which mirrors the slow actor-environment interaction seen in RL systems. Inspired by distributed deep RL systems which spawn many actors or environments in parallel [22, 75], we design an Actor-Learner-Oracle architecture for online LLM alignment, which is depicted in Fig. 5. The three types of workloads (i.e., actor, learner and oracle) are heterogeneous and require different optimization. In particular, we adopt vLLM [33] for the actor to accelerate the autoregressive response generation. We also use DeepSpeed’s ZeRO [58, 57] strategies to enhance the memory efficiency of the learner. The updated model weights are broadcasted from the learner master to all actors after every optimizer step efficiently via NCCL, similar to Hu et al. [26]. Furthermore, to improve the scalability, we wrap the oracle RM as a service using Mosec [83], which supports dynamic batching and parallel processing, to minimize preference query latency. Finally, we leverage DeepMind Launchpad [82] to compose all workloads into a distributed program and adopt Plasma [53] to efficiently transfer data across process boundaries.

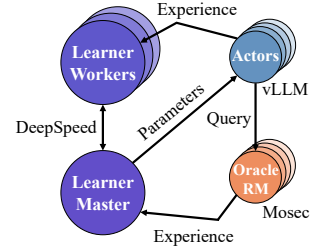


Figure 5: The learning system for experimenting online LLM alignment algorithms.

We benchmark our system’s efficiency against a concurrent implementation of online DPO by HuggingFace⁶, which utilizes only DeepSpeed for memory optimization. Our system achieves up to **2.5×** latency reduction compared to this counterpart, demonstrating its computational efficiency. Due to space constraints, detailed benchmarking methods and results are presented in App. H.

E Baseline methods

We review four baseline methods that are relevant to this work and used for comparisons in our experiments.

Offline DAP. We review DPO [55], which is a representative work in the direction of Direct Alignment from Preferences (DAP). It simplifies the two-stage pipeline of offline RLHF as a single step of supervised learning by leveraging the closed-form solution [52, 51] of the RL objective in Eq. (15):

$$\pi_r(\mathbf{y}|\mathbf{x}) = \frac{1}{Z(\mathbf{x})} \pi_{\text{ref}}(\mathbf{y}|\mathbf{x}) \exp\left(\frac{1}{\beta} r(\mathbf{x}, \mathbf{y})\right), \quad (18)$$

⁶https://huggingface.co/docs/trl/main/en/online_dpo_trainer.

where $Z(\mathbf{x})$ normalizes such that $\sum_{\mathbf{y}} \pi_r(\mathbf{y}|\mathbf{x}) = 1$, to reparametrize r as a function of π :

$$r(\mathbf{x}, \mathbf{y}) = \beta \log \frac{\pi_r(\mathbf{y}|\mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}|\mathbf{x})} + \beta \log Z(\mathbf{x}). \quad (19)$$

Consequently, plugging Eq. (19) into the reward model loss (Eq. (13)) yields a contrastive loss that directly optimizes the policy:

$$\min_{\pi_{\theta}} \mathbb{E}_{(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-) \sim p_{\mathcal{D}}} \left[-\log \sigma \left(\beta \log \frac{\pi_{\theta}(\mathbf{y}^+|\mathbf{x}) \pi_{\text{ref}}(\mathbf{y}^-|\mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}^+|\mathbf{x}) \pi_{\theta}(\mathbf{y}^-|\mathbf{x})} \right) \right], \quad (20)$$

where $\mathcal{D}$ is a pre-collected offline preference dataset.

We also experiment different DAP methods⁷ besides DPO, such as IPO [4] and SLiC [88], whose loss functions are shown in Eq. (11) and (12).

Online DAP [24]. In contrast to the conventional DAP methods that learn a policy from a fixed dataset $\mathcal{D}$, online DAP proposes to collect on-policy preference data to update the policy online. It first samples responses from the current policy $(\mathbf{y}, \mathbf{y}') \sim \pi_{\theta_t}$, then acquires preference labels to form a batch $\mathcal{B}_t = \{(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-)\}_{i=1}^b$. One gradient step minimizing the DAP loss over this data batch to get $\pi_{\theta_{t+1}}$, which is used for the next iteration. Such approach not only mitigates the over-fitting issue faced by offline DAP methods [24], but also facilitates online interaction (Property 1) with the environment, falling into the second paradigm of CDB solution algorithms (Fig. 2b).

Active Preference Learning (APL) [47]. APL follows the online DAP paradigm, but is restricted to DPO due to its reliance on DPO implicit rewards. Two techniques are proposed by APL to actively select both prompts and dueling responses for querying the preference oracle:

1. *Predictive entropy (PE)* for selecting prompts. In this step APL computes a Monte-Carlo estimate of PE for each prompt as $\mathcal{H}_{\pi_{\theta}}(\mathbf{y}|\mathbf{x}) \approx -\sum_{n=1}^N \log \pi_{\theta}(\mathbf{y}_n|\mathbf{x})/N$, where $\mathbf{y}_n \sim \pi_{\theta}(\cdot|\mathbf{x})$ and $\log \pi_{\theta}(\mathbf{y}_n|\mathbf{x})$ is the summation of log probabilities of each token. Then, APL filters a subset of prompts with high PE to form $\mathcal{X}_S$.
2. *Preference model certainty* for selecting dueling responses. For prompts in $\mathcal{X}_S$, APL generates many responses for each prompt, then selects the pair with largest reward margin measured as $|\hat{r}(\mathbf{x}_i, \mathbf{y}_i) - \hat{r}(\mathbf{x}_i, \mathbf{y}'_i)|$, where $\hat{r}$ is the DPO implicit reward $\hat{r}(\mathbf{x}, \mathbf{y}) = \beta(\log \pi_{\theta}(\mathbf{y}|\mathbf{x}) - \log \pi_{\text{ref}}(\mathbf{y}|\mathbf{x}))$.

By above two steps, APL actively explores more uncertain prompts and responses in an online DPO paradigm, satisfying both Properties 1 and 2.

Exploratory Preference Optimization (XPO) [79]. XPO studies LLM alignment in the framework of token-level MDP, and leverages the property that DPO conducts *implicit Q^* -approximation* [56], so that

$$\beta \log \frac{\pi^*(\mathbf{y}|\mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}|\mathbf{x})} = r^*(\mathbf{x}, \mathbf{y}) - V^*(\mathbf{x}) \quad \forall \mathbf{y}, \quad (21)$$

where V^* is the optimal value function depending only on the prompt $\mathbf{x}$. XPO incorporates the *implicit (global) optimism* for exploration by overestimating the value $V_{\pi_{\theta}}(\mathbf{x}) = r^*(\mathbf{x}, \mathbf{y}) - \beta \log \frac{\pi_{\theta}(\mathbf{y}|\mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}|\mathbf{x})}$. This is achieved by optimizing the policy with a modified DPO loss:

$$\min_{\pi_{\theta}} \mathbb{E}_{(\mathbf{x}, \mathbf{y}^+, \mathbf{y}^-, \mathbf{y}^{\text{ref}}) \sim p_{\mathcal{B}^t}} \left[\alpha \log \pi_{\theta}(\mathbf{y}^{\text{ref}}|\mathbf{x}) - \log \sigma \left(\beta \log \frac{\pi_{\theta}(\mathbf{y}^+|\mathbf{x}) \pi_{\text{ref}}(\mathbf{y}^-|\mathbf{x})}{\pi_{\text{ref}}(\mathbf{y}^+|\mathbf{x}) \pi_{\theta}(\mathbf{y}^-|\mathbf{x})} \right) \right], \quad (22)$$

where $\mathbf{y}^{\text{ref}} \sim \pi_{\text{ref}}(\cdot|\mathbf{x})$ and $\mathcal{B}^t$ is an on-policy data batch in the same vein as online DPO. Intuitively, the first term in Eq. (22) biases the policy toward a large value estimation such that $V_{\pi_{\theta}} \gtrsim V^*$, implementing the optimism in the face of uncertainty (OFU) for exploration. Theoretically, Xie et al. [79] also prove the sample complexity bound of XPO, making it a promising algorithm for online LLM alignment.

Self-exploring language model (SELM) [86] is a concurrent work of Xie et al. [79] that proposes nearly the same theoretic algorithm to achieve OFU. However, the practical implementation of SELM involves offline preference dataset for training, making it hard to benchmark in an online alignment setting like ours. Therefore, we will keep XPO as our baseline for comparison.

⁷We use ‘‘DAP method’’ and ‘‘direct optimizer’’ interchangeably.

F Full experimental details

In the main text we focus on the task of summarization using the TL;DR dataset. This provides a lightweight and clean setting to extensively study different algorithmic designs with affordable computational resources. App. F.1 provides the full details of this setting.

To further validate the sample efficiency of SEA in aligning LLMs to perform general tasks, we adopt the UltraFeedback dataset [16] and evaluate trained LLMs on AlpacaEval 2.0 [38]. App. F.2 provides more details of this setting.

F.1 Details of TL;DR task

Models. We experiment three model scales (1B, 2.8B, 6.9B) from the Pythia family [7]. We take pretrained SFT models from [27] as π_{ref} for the starting model in all experiments. Except in Sec. 6.1, we use 1B model for other experiments to save computation.

Preference oracle. We simulate the process of human feedback with a strong scalar RM and refer it as preference oracle. We choose Skywork-Reward-Llama-3.1-8B⁸ [40], which is top-ranked in RewardBench leaderboard [35], as the preference oracle.

Epistemic reward model. We build ERM on top of a pretrained 0.4B transformer [30], by removing its head and adding an ensemble of MLPs. The size of ensemble is set to $K = 20$, and all MLPs contain 2 hidden layers of 128 nodes. Note that the ERM is chosen to be much smaller than the preference oracle following Dwaracherla et al. [21], which reflects the fact that human preferences can be more complex than what the agent can model. The regularization coefficient λ is fixed to be 0.5 after a coarse hyperparameter search.

Data. We employ the widely adopted TL;DR dataset [65] for our experiments. It consists of Reddit posts as prompts, and the agent is required to give summaries that align with human preferences. We fix 50k prompts for training and limit the query budget to 50k as well.

DAP methods. We adopt three DAP methods (direct optimizers) to thoroughly validate our algorithm, including DPO [55], IPO [4] and SLiC [88]. Except in Sec. 6.1, all experiments are done with DPO as the direct optimizer.

Baselines. Similar to Guo et al. [24], we include the offline and online variants of different DAP methods as baselines. Additionally, we compare with two active exploration baselines built on online DPO: APL [47] and XPO [79]. A detailed review of all baselines can be found in App. E.

Metrics. We use the win rate of agent’s responses against reference responses judged by the preference oracle as the performance metric. This metric can reflect both the agent’s cumulative regret and anytime regret (i.e., average performance). In the E&E setting, we measure the “online” win rate of the agent’s dueling responses that are executed during experience collection and take the average. In the BAI setting, we measure the “offline” win rate by evaluating the latest agent’s responses given a fixed set of 1000 holdout prompts periodically. We mainly focus on the BAI setting because crowdsourcing seems a major scenario for most practitioners, and present one set of experiments for comparing different exploration strategies in both settings. When the comparison is only made within a model scale, we report the relative win rate against the initial STF models. When the comparison is across scales (Fig. 1 Left), we report the absolute win rate against the ground truth responses in the dataset.

Hyperparameters. We set $\beta = 0.1$ for DPO and $\beta = 0.2$ for SLiC and find they are robust for all scales. We tune β from $\{0.2, 0.3, 0.5, 1.0\}$ for IPO across scales and report the best performing results. We sample $M = 20$ on-policy responses with a temperature $\eta = 0.7$ during training, and use greedy decoding for offline evaluation (BAI’s metric). We use the Adam optimizer with learning rate of 5×10^{-7} and cosine scheduling, and set the batch size to be 128. We initialize the mixture ratio γ of SEA to be 1 and adjust it to 0.7 after a burn-in period of 1k samples.

All hyperparameters are kept the same for offline and online baselines, except that online methods update the sampling policy after every gradient step as the latest π_{θ_t} . For APL and XPO, we keep the learning rate and DPO’s β the same for apple-to-apple comparisons. Specifically for APL, we initially sample 1024 prompts per batch and use the predictive entropy to filter a subset of 128 prompts. Then, we sample 8 responses per prompt and use the preference model certainty to finalize two responses

⁸<https://huggingface.co/Skywork/Skywork-Reward-Llama-3.1-8B>.

847 for the duel. Specifically for XPO, we follow the their recommended optimism coefficient to set
848 $\alpha = 5 \times 10^{-6}$.

849 **Statistical significance.** There are various factors to introduce randomness during online learning.
850 We thus launch 3 independent runs for every experiment with different random seeds. All the results
851 are reported with mean and standard error to indicate their statistical significance.

852 **Computational resources.** Experiments at all scales are conducted on a single machine with 8
853 A100 GPUs to run the learner and actors. We additionally host a separate remote server with workers
854 spawned on 16 A100 GPUs for the oracle RM⁹, so that it can be queried by all concurrently running
855 experiments. All experiments conducted for this research consume about 2 A100 GPU years.

856 F.2 Details of general tasks

857 **Model.** Following Meng et al. [46], Zhang et al. [86], we employ Llama3-8B-Instruct¹⁰ as our
858 initial model π_{ref} .

859 **Preference oracle.** We follow Meng et al. [46] to adopt ArmoRM-Llama3-8B-v0.1¹¹ [74] as the
860 preference oracle to provide online preference feedback.

861 **Data.** We take the UltraFeedback dataset [16], which is widely used for LLM alignment in the
862 literature. We filter out samples whose prompt is longer than 1800 tokens and result in 61k samples.
863 We extract prompts from the filtered dataset while excluding the responses. The prompt set are
864 collected from multiple sources and cover diverse domains, making it suitable to improve LLM’s
865 capability on general tasks.

866 **DAP method and baselines.** We employ the state-of-the-art DAP method, SimPO [46], as our direct
867 optimizer. Since SimPO is originally an offline algorithm, we extend it to Online SimPO and take
868 both offline and online variants as baselines.

869 **Evaluation.** We evaluate SEA and baselines using AlpacaEval 2.0 [38]. It consists of 805 test
870 prompts, and uses GPT4-Turbo to judge the quality of model responses against reference responses
871 generated by GPT-4-Turbo. We follow the standard protocol to report both the win rate (WR) and
872 the Length-Controlled win rate (LC) [19].

873 **Hyperparameters.** We follow SimPO’s recommended hyperparameters to set $\beta = 10$ and $\gamma/\beta = 0.3$.
874 We use a learning rate of 8×10^{-7} and batch size of 128. The decoding temperature is set to be 0.9
875 for generating evaluation outputs. The same hyperparameters apply to baselines and our method.
876 Configurations of SEA are kept the same as those in the TL ; DR task (App. F.1).

877 G Extended empirical studies

878 We present additional empirical studies in this section, including investigation on different exploration
879 strategies (App. G.1) and preference oracles (App. G.2) on the TL ; DR task, as well as the performance
880 comparison on AlpacaEval 2.0 for general tasks (App. G.3).

881 G.1 Choice of exploration strategies

882 Recalling that different LLM alignment scenarios (online system or crowdsourcing) require different
883 exploration strategies to meet their respective learning objectives (Sec. 2.2). We investigate three strate-
884 gies based on posterior sampling and compare them on both online and offline performance. The first
885 strategy (Uncertainty) focuses on pure exploration with information maximization. It seeks the pair of
886 dueling responses that exhibits the largest epistemic uncertainty, which is implemented by selecting
887 the pair whose logits difference has the largest variance across ensemble members. The second (E&E-
888 TS) and the third (BAI-TS) strategies follow the principles in Algorithm 1, and their differences are
889 between Line 5 and Line 6. The comparison results are shown in Fig. 6 (Left and Middle). Focusing on
890 the left plot, we observe that E&E-TS strategy achieves the best online performance, which is within
891 our expectation. In contrast, Uncertainty shows the worst online performance because it tries to maxi-
892 mize the information gain but does not prioritize reward maximization. On the other hand, conclusions
893 are interestingly different when taking the offline performance as the metric. In this case, BAI-TS and

⁹We utilize the Kubernetes service for routing requests to multiple Mosec [83] instances.

¹⁰<https://huggingface.co/meta-llama/Meta-Llama-3-8B-Instruct>.

¹¹<https://huggingface.co/RLHFlow/ArmoRM-Llama3-8B-v0.1>.

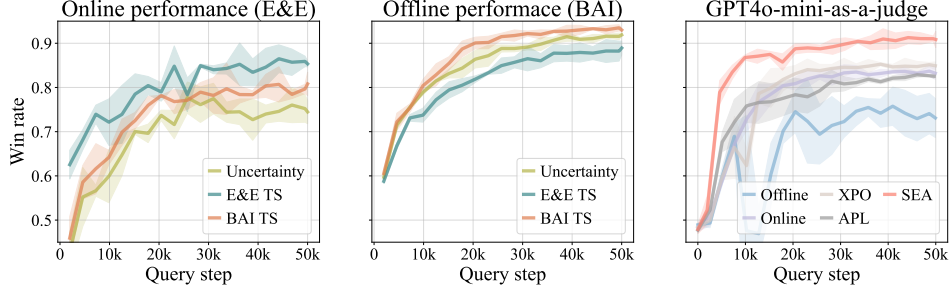


Figure 6: (Left and Middle) Win rate comparison of different exploration strategies measured in E&E and BAI settings. **(Right)** Win rate comparison of different agents when using GPT4o-mini to simulate human feedback via LLM-as-a-judge.

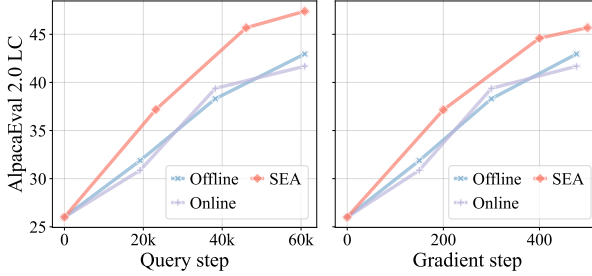


Figure 7: LC win rates on AlpacaEval 2.0 with respect to query budget and gradient update budget.

Table 3: AlpacaEval 2.0 results. LLM exploration methods are highlighted in blue.

Model	LC	WR
GPT-4 Omni (05/13)	57.5	51.3
GPT-4 Turbo (04/09)	55.0	46.1
Yi-Large Preview	51.9	57.5
SEA+SimPO	47.4	41.1
Claude 3 Opus (02/29)	40.5	26.1
SELM	34.7	34.8
XPO	29.4	-
Llama 3 8B Instruct	22.9	22.6

Uncertainty both exhibit more efficient offline performance improvement than E&E-TS. This can be attributed to that exploration for uncertainty minimizing helps to identify more informative responses to train the LLM policy. Moreover, BAI-TS $\succ$ Uncertainty indicates exploration with both reward and information maximization is better than exploration with only information maximization. E&E-TS, however, always chooses two responses with similarly high quality to exploit. This can not only lead to less efficient exploration, but also result in less efficient policy learning due to smaller DAP loss gradients.

G.2 Aligning LLMs with a human simulator

Results presented so far are based on experimenting LLM alignment with the preference oracle being a scalar reward model, which is deterministic and does not capture the potential randomness of the choice by real humans. To test different agents in a more realistic setting, we use generative models as human simulator in an LLM-as-a-judge [10, 89] manner. In particular, we directly query the OpenAI API and use gpt-4o-mini-2024-07-18 as the judge to provide preference feedback. We use a similar prompt template to Li et al. [38]’s, which is shown in Fig. 10. We also randomly swap the order of two responses to mitigate the known position bias of LLM judges. The results are shown in Fig. 6 (Right). We can observe the performance curves generally exhibit higher variance, possibly due to the randomness introduced in the feedback process, which puts more stringent requirements for learning algorithms. The two active exploration methods demonstrate opposite results to those in Sec. 6.1—APL learns fast initially but is eventually outperformed by Online, while XPO improves over Online after stabilizing its training and delivers a better final performance. Our agent, **SEA**, is shown to offer the best sample efficiency as well as asymptotic performance, further validating the importance of online learning and well-designed active exploration mechanism.

G.3 Performance on general tasks

We investigate the generalizability of **SEA** by training with the prompt set from UltraFeedback [16] and evaluating the model performance on AlpacaEval 2.0 [38]. Fig. 7 shows the Length-Controlled (LC) win rate of different models against GPT-4-Turbo. The left plot compares the sample efficiency (in terms of the number of queries) of offline, online and **SEA** SimPO. The results suggest that enabling online interaction does not improve the sample efficiency over the offline counterpart. Such observation is in stark contrast to what we have seen in the TL ; DR task, where the online agent always improves over the offline ones. We hypothesize that this is due to the different coverage of π_{ref} in

these two tasks. For TL;DR, which is a much easier task, the initial SFT models already have good coverage, permitting online DAP with only *passive exploration* to work reasonably well; however, for more challenging tasks, the insufficient coverage of π_{ref} would lead to sample complexity exponential in $\frac{1}{\beta}$ [79], which necessitates *deliberate exploration*, such as Thompson sampling proposed in this work. The above claim is justified by observing that SEA largely improves the sample efficiency over the online and offline variants.

Attentive readers may have noticed that comparing query budget could be advantageous to SEA because pseudo labels are used in mixed preference learning (Sec. 4.2.3), which results in more gradient steps given the same query budget. In the right plot of Fig. 7, we show the performance versus gradient step. We can observe SEA has the steepest learning curve, verifying that it explores more informative samples to yield faster improvement.

Last but not least, in Table 3, we show the AlpacaEval 2.0 LC win rates of XPO and SELM (as reported in their papers), along with ours and several cutting-edge LLMs. SEA is agnostic to direct optimizers, thus it can leverage the state-of-the-art SimPO to achieve a high LC of 47.4%. On the other hand, XPO and SELM can only be applied to DPO, restricting their potential to incorporate future advances in direct optimization algorithms.

H System benchmarking

We conduct a rigorous benchmarking comparison on the efficiency of online DPO training using our learning system, alongside the trl’s implementation¹².

Settings. In alignment with the examples provided by trl, we use the TL;DR [65] dataset and evaluate training efficiency at three model scales: 1B, 2.8B and 6.9B parameters for both SFT-ed LLMs¹³ and exclusively trained RMs¹⁴. This is similar to the settings in our experiments (see App. F) except that we fix the preference oracle to be a strong general-purpose RM.

Hardware & Software. All benchmarking experiments are conducted on a single machine with eight A100-40G GPUs and 96 AMD EPYC 7352 CPUs. To ensure fair comparison, we align all key hyperparameters for both our codebase and trl. The DeepSpeed ZeRO-2 strategy is employed by default when GPU memory suffices; otherwise, ZeRO-3 or ZeRO-2-offload is utilized as applicable. Notably, the distributed architecture of our implementation provides flexibility in system configuration, enabling adjustments to accommodate memory and computational time constraints. Fig. 8 illustrates two example configurations employed in our benchmarking experiments. We will provide all benchmarking scripts in our codebase for reproducibility.

- **Config 1** collocates all three workloads on each of the GPUs. Specifically, eight vLLM instances (for actors) and eight Mosec workers (for oracle RMs) are spawned to run independently on each GPU. After a batch of responses is generated (by actors) and labeled (by oracle RMs), it is sent to the learner, which runs on all eight GPUs coordinated through ZeRO strategies for policy learning. The updated policy weights are then broadcasted to all actors for *on-policy* response sampling on subsequent prompt batch. While this configuration maximizes GPU utilization, it requires substantial GPU memory to accommodate all workloads and is thus employed only for 1B scale experiments.
- **Config 2** only collocates actor and oracle workloads on half of the GPUs, reserving the remaining four GPUs exclusively for the learner. This is suited for larger-scale experiments (e.g., 2.8B or 6.9B), where additional GPU memory is allocated to the learner. However, this setup incurs idle time on half of the GPUs due to data dependency, as the learner must await new preference data, and the actor must await updated policies. An alternative is to implement *asynchronous* data collection, where minor data staleness is allowed by using θ_{t-1} to generate data for updating θ_t . Although this data would not be strictly on-policy,

¹²https://github.com/huggingface/trl/blob/main/trl/trainer/online_dpo_trainer.py.

¹³<https://huggingface.co/trl-lib/pythia-1b-deduped-tldr-sft>; <https://huggingface.co/trl-lib/pythia-2.8b-deduped-tldr-sft>; <https://huggingface.co/trl-lib/pythia-6.9b-deduped-tldr-sft>

¹⁴<https://huggingface.co/trl-lib/pythia-1b-deduped-tldr-rm>; <https://huggingface.co/trl-lib/pythia-2.8b-deduped-tldr-rm>; <https://huggingface.co/trl-lib/pythia-6.9b-deduped-tldr-rm>

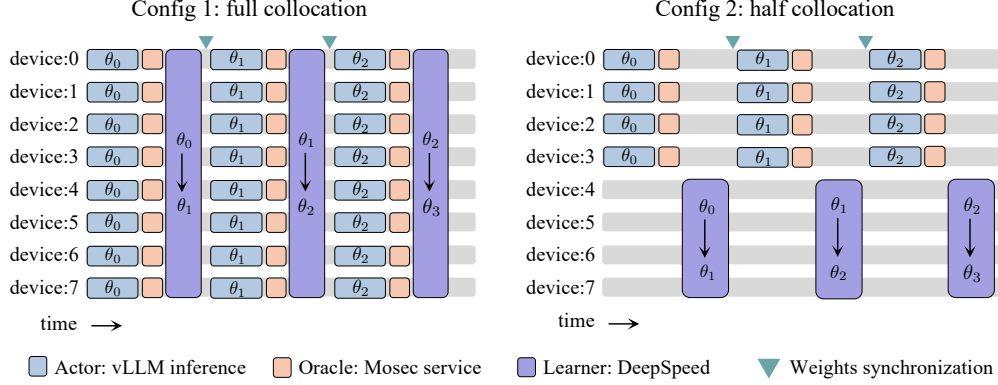


Figure 8: Two example configurations of our learning system used in benchmarking experiments.

asynchronous training could reduce idle time and enhance GPU utilization. This approach has proven effective in large-scale RL systems [6], and we leave this optimization to future work.

Results. Benchmarking results for the latency of training a batch of 128 samples are presented in Fig. 9. Overall, training with the config 2 demonstrates consistently greater efficiency than trl, achieving up to a $2.5\times$ reduction in latency at the 2.8B scale.

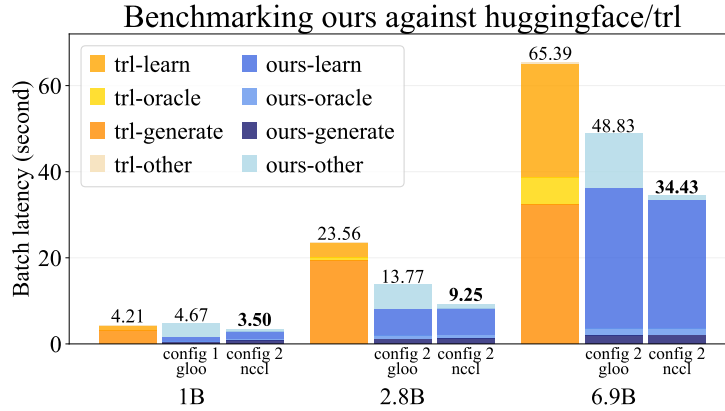


Figure 9: Averaged training latency (over 10 batches, equivalent to 1280 samples) comparing ours against huggingface/trl.

We next analyze the time costs for individual stages: generate, oracle and learn. Across all scales and configurations, ours demonstrates significantly lower *generate* time than trl, due to distributed actors utilizing vLLM. Additionally, at the 6.9B scale, ours requires substantially less *oracle* time than trl, as trl employs ZeRO-3 to prevent GPU memory overflow, thereby slowing inference. In contrast, ours config 2 allows for flexible collocation, enabling oracle RMs hosted via Mosec to operate in parallel without sharding. However, ours config 2 incurs longer *learn* time compared to trl due to the use of only half the available GPUs. This limitation also explains why, at the 1B scale, config 2 has higher latency than config 1 across all stages.

The *other* category accounts for time costs associated with data loading, tokenization, and communication. Here, inter-process communication is the primary cost, with trl showing minimal overhead as all three stages operate within the same process on identical micro-batches, avoiding weight synchronization. By contrast, ours requires considerable time to transfer updated policy weights from the learner to all actors. While NCCL is recommended for synchronization over GLOO, it requires older vLLM packages (prior to version 0.4.3), which may lack support for newer LLM architectures. Moreover, NCCL is incompatible with config 1 due to its restriction on the learner master process establishing two separate process groups (one for DeepSpeed, the other for weight synchronization). In summary, we recommend future researchers prioritize the config 2 and employ NCCL when feasible.

993 I Additional materials

994 In this section, we include a comparison of prior works (Table 2), the prompt template we use for
995 LLM-as-a-judge (Fig. 10), and an example showing ChatGPT actively explores (Fig. 6).

LLM-as-a-Judge Prompt Template

I require a leaderboard for various large language models. I'll provide you with prompts given to these models and their corresponding outputs. Your task is to assess these responses, and select the model that produces the best output from a human perspective.

Instruction

{instruction}

Model Outputs

Here are the unordered outputs from the models. Each output is associated with a specific model, identified by a unique model identifier.

```
{  
  "model_identifier": "0",  
  "output": {response0}  
}
```

```
{  
  "model_identifier": "1",  
  "output": {response1}  
}
```

Task

Evaluate the models on the basis of the quality and relevance of their results, and select the model that generated the best result. Reply with the identifier of the best model. Our evaluation will only take into account the first character of your answer, so make sure it contains only one of the identifiers and nothing else (no quotation marks, no spaces, no new lines, ...).

Figure 10: We provide the prompt with two responses and ask GPT to choose the preferred one. We parse the first token's logits of "0" and "1" and apply the BT model to obtain the relative preference.

I want to use this for k8s readiness check, which requires a return value "0" if success, how to write the script?

You're giving feedback on a new version of ChatGPT.
Which response do you prefer? Responses may take a moment to load.

Response 1

To create a script for a Kubernetes readiness check using `curl`, where you return a `0` status on success (typically when the HTTP status code is `200`), you can combine `curl` with a conditional check on the HTTP status code. If the status is `200`, the script will exit with `0` (indicating success); otherwise, it will return a non-zero value (indicating failure).

Here's an example of how you can write this script:

```
bash
#!/bin/bash

# Send a POST request and get the HTTP status code
STATUS=$(curl -X POST -H "Content-Type: application/octet-stream" -s -o /dev/null -w "%{http_code}")

# Check if the status code is 200
if [ "$STATUS" -eq 200 ]; then
    exit 0 # Readiness check successful
else
    exit 1 # Readiness check failed
fi
```

Explanation:

- `curl` sends the request and captures the HTTP status code in the `STATUS` variable.
- The `if` statement checks if the status code is `200`.
 - If the status code is `200`, the script exits with `0`, indicating success (readiness).
 - If the status is anything other than `200`, the script exits with `1`, indicating failure.

You can modify the binary string or URL as necessary for your specific readiness check.

I prefer this response

Response 2

Memory updated

For a Kubernetes readiness check, you need to return `0` if the `curl` command succeeds (e.g., status code `200`), and a non-zero value if it fails.

Here's a simple bash script that achieves this:

```
bash
#!/bin/bash

URL="<your-endpoint>"
STATUS_CODE=$(curl -X POST -H "Content-Type: application/octet-stream" -s -o /dev/null -w "%{http_code}")

if [ "$STATUS_CODE" -eq 200 ]; then
    exit 0 # success
else
    exit 1 # failure
fi
```

Breakdown:

- `curl ... -w "%{http_code}"`: Captures the HTTP status code from the `curl` request.
- The `if` condition checks if the status code is `200` (indicating success).
- `exit 0`: Indicates a successful readiness check.
- `exit 1`: Indicates failure, so Kubernetes will consider the pod not ready.

You can adjust the script depending on the specific status code(s) that you consider successful.

I prefer this response

Figure 11: ChatGPT system asks for users' preference feedback to strategically explore better answers. In this case, algorithms should be designed around the objective of *minimizing cumulative regret* (i.e., the E&E setting), because the quality of both responses generated by the system affects user experience.

NeurIPS Paper Checklist

1. Claims

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

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

Justification: We include the paper's contributions and scope in the abstract and introduction.

Guidelines:

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

2. Limitations

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

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

Justification: There is no obvious limitation.

Guidelines:

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

3. Theory assumptions and proofs

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

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

Justification: We do not have theorems.

Guidelines:

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

4. Experimental result reproducibility

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

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

Justification: We have explained our implementation details in Sec. 5.

Guidelines:

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

5. Open access to data and code

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

1101 Answer: [No]

1102 Justification: We haven't include the code in the submission but we will fully open source

1103 the codes after the reviewing process.

1104 Guidelines:

- 1105 • The answer NA means that paper does not include experiments requiring code.
- 1106 • Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- 1107 • While we encourage the release of code and data, we understand that this might not be
- 1108 possible, so "No" is an acceptable answer. Papers cannot be rejected simply for not
- 1109 including code, unless this is central to the contribution (e.g., for a new open-source
- 1110 benchmark).
- 1111 • The instructions should contain the exact command and environment needed to run to
- 1112 reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- 1113 • The authors should provide instructions on data access and preparation, including how
- 1114 to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- 1115 • The authors should provide scripts to reproduce all experimental results for the new
- 1116 proposed method and baselines. If only a subset of experiments are reproducible, they
- 1117 should state which ones are omitted from the script and why.
- 1118 • At submission time, to preserve anonymity, the authors should release anonymized
- 1119 versions (if applicable).
- 1120 • Providing as much information as possible in supplemental material (appended to the
- 1121 paper) is recommended, but including URLs to data and code is permitted.

1122

1123

1124 **6. Experimental setting/details**

1125 Question: Does the paper specify all the training and test details (e.g., data splits, hyper-

1126 parameters, how they were chosen, type of optimizer, etc.) necessary to understand the

1127 results?

1128 Answer: [Yes]

1129 Justification: We have presented all the details in App. F.

1130 Guidelines:

- 1131 • The answer NA means that the paper does not include experiments.
- 1132 • The experimental setting should be presented in the core of the paper to a level of detail
- 1133 that is necessary to appreciate the results and make sense of them.
- 1134 • The full details can be provided either with the code, in appendix, or as supplemental
- 1135 material.

1136

1137 **7. Experiment statistical significance**

1138 Question: Does the paper report error bars suitably and correctly defined or other appropriate

1139 information about the statistical significance of the experiments?

1140 Answer: [Yes]

1141 Justification: All experiments are repeated for 3 times with different seeds for statistical

1142 significance.

1143 Guidelines:

- 1144 • The answer NA means that the paper does not include experiments.
- 1145 • The authors should answer "Yes" if the results are accompanied by error bars, confi-
- 1146 dence intervals, or statistical significance tests, at least for the experiments that support
- 1147 the main claims of the paper.
- 1148 • The factors of variability that the error bars are capturing should be clearly stated (for
- 1149 example, train/test split, initialization, random drawing of some parameter, or overall
- 1150 run with given experimental conditions).
- 1151 • The method for calculating the error bars should be explained (closed form formula,
- 1152 call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).

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

8. Experiments compute resources

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

Answer: [Yes]

Justification: We have included the description about computational resources at line 857.

Guidelines:

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

9. Code of ethics

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

Answer: [Yes]

Justification: We conform with the NeurIPS Code of Ethics.

Guidelines:

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

10. Broader impacts

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

Answer: [NA]

Justification: Our work studies the sample efficiency of LLM fine-tuning from a research point of view, which should not directly lead to harmful societal impact.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

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

11. Safeguards

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

Answer: [NA]

Justification: No such risks.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [Yes]

Justification: We properly cited all the existing resources that we used in this research.

Guidelines:

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

- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. New assets

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

Answer: [NA]

Justification: We do not release new assets as of the paper itself. We will open source our code for reproducibility with documents.

Guidelines:

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

14. Crowdsourcing and research with human subjects

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

Answer: [NA]

Justification: We do not involve crowdsourcing nor research with human subjects.

Guidelines:

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

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

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

Answer: [NA]

Justification: We do not involve crowdsourcing nor research with human subjects.

Guidelines:

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

1306 **16. Declaration of LLM usage**
1307 Question: Does the paper describe the usage of LLMs if it is an important, original, or
1308 non-standard component of the core methods in this research? Note that if the LLM is used
1309 only for writing, editing, or formatting purposes and does not impact the core methodology,
1310 scientific rigorousness, or originality of the research, declaration is not required.
1311 Answer: [Yes]
1312 Justification: We used LLMs for editing (e.g., checking grammar errors).
1313 Guidelines:
1314 • The answer NA means that the core method development in this research does not
1315 involve LLMs as any important, original, or non-standard components.
1316 • Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for
1317 what should or should not be described.