
Trajectory Bellman Residual Minimization: A Simple Value-Based Method for LLM Reasoning

Yurun Yuan¹ Fan Chen² Zeyu Jia² Alexander Rakhlin^{2*} Tengyang Xie^{1*}

¹University of Wisconsin–Madison ²Massachusetts Institute of Technology
{yurun_yuan, tx}@cs.wisc.edu
{fanchen, zyjia, rakhlin}@mit.edu

Abstract

Policy-based methods currently dominate reinforcement learning (RL) pipelines for large language model (LLM) reasoning, leaving value-based approaches largely unexplored. We revisit the classical paradigm of Bellman Residual Minimization and introduce Trajectory Bellman Residual Minimization (TBRM), an algorithm that naturally adapts this idea to LLMs, yielding a simple yet effective off-policy algorithm that optimizes a single trajectory-level Bellman objective using the model’s own logits as Q -values. TBRM removes the need for critics, importance-sampling ratios, or clipping, and can operate with only one rollout per prompt. We prove convergence to the near-optimal KL-regularized policy from arbitrary off-policy data via an improved change-of-trajectory-measure analysis. Experiments on standard mathematical-reasoning benchmarks show that TBRM matches or surpasses policy-based baselines, like PPO and GRPO, with comparable or lower computational and memory overhead. Our results indicate that value-based RL might be a principled and efficient alternative for enhancing reasoning capabilities in LLMs. The codebase for TBRM is publicly available at <https://github.com/r1x-lab/TBRM>.

1 Introduction

Large language models (LLMs) have become the de-facto backbone for modern natural-language understanding and generation (Brown et al., 2020; Ouyang et al., 2022). While ever-larger pre-training corpora push the frontier of *knowledge*, high-value downstream usage increasingly hinges on *reasoning*: the capacity to carry out multi-step thinking, apply abstract rules to practical situations, and generalize from observed patterns to solve complex, structured problems. Reinforcement learning (RL) with verifiable, outcome-based rewards has emerged as a powerful paradigm for enhancing this reasoning capability in LLMs (Guo et al., 2025), especially for mathematical problem-solving where correctness can be objectively determined.

Recent advances in LLM post-training for mathematical reasoning have primarily employed policy-based variants—Proximal Policy Optimization (PPO; Schulman et al., 2017b) and Group Relative Policy Optimization (GRPO; Shao et al., 2024; Guo et al., 2025)—which optimize policies to maximize objective rewards that indicate successful task completion. These approaches leverage the clear evaluation criteria of mathematical tasks, where responses can be automatically verified as correct or incorrect without requiring human judgment.

Despite their empirical success, policy-based methods face several practical challenges. They typically require fresh on-policy rollouts from the current model, increasing computational demands. They often rely on additional components like critic models, advantage normalization, and clipping

*Corresponding authors.

mechanisms, adding complexity to implementation and tuning. Moreover, for token-level decisions, these methods (1) simplistically attribute outcome-based rewards (e.g., correctness of the entire response) to individual tokens, often assigning credit primarily to the final token, and (2) bootstrap advantages with truncated rollout horizons when additional critic models are used, potentially compromising effective credit assignment during training.

Classical RL offers a compelling alternative: *value-based* methods learn an action-value function Q and act, for instance, greedily with respect to it. In the LLM setting, each token is an action, and the model’s logits (that is, raw network outputs) naturally provide a parametric family for Q , under the KL-regularized RL framework (e.g., Schulman et al., 2017a). Despite this natural alignment and the inherent strengths of value-based methods for LLMs, their application to LLMs has been limited, with policy-based techniques being more prevalent. One potential reason, we conjecture, is the perceived difficulty in reconciling traditional iterative bootstrapping in value-based RL (e.g., the Q-learning family, Watkins and Dayan, 1992; Mnih et al., 2015) with the scale of LLM training, as iterative-style algorithms are typically less stable than their optimization-style counterparts (e.g., policy-based methods).

Our starting point is *Bellman Residual Minimization* (BRM; Schweitzer and Seidmann, 1985; Baird et al., 1995), a decades-old idea that fits the Q -function by directly minimizing its Bellman residual, designed for deterministic environments (which LLMs naturally are). By leveraging the recent theoretical advances in trajectory-level change of measure (Jia et al., 2025), we recognize that BRM can be lifted from the *step (token)* to the *trajectory* level: we square a single residual spanning the whole rollout and regress the model’s logits onto it. This approach eliminates the aforementioned per-step-signal barrier, removes the need for critics, importance weights, and clipping, and provably maintains fully off-policy optimization due to the value-based nature of the algorithm. The resulting algorithm is *Trajectory BRM* (TBRM), which builds on classical BRM. Below we state our contributions, focused on theoretical analysis and extensive experiments on math reasoning tasks.

1.1 Our Results

1. **Algorithm.** Building explicitly on the classical idea of Bellman Residual Minimization, we present *TBRM*, a single-objective, off-policy algorithm that fits the trajectory-level Bellman residual using the LLM logits as Q -values. TBRM dispenses with critics, advantage estimates, importance ratios, or clipping, and can operate with only *one rollout per prompt*, while scaling effectively with multiple rollouts in practice.
2. **Theory.** We prove that, under standard realizability assumptions, the algorithm converges to the *optimal KL-regularized policy* even when training data are generated by *arbitrary* behavior policies in deterministic environments (such as LLMs). Our results build upon the recent change-of-trajectory-measure result of Jia et al. (2025). We significantly simplify that proof and improve the rate of convergence in terms of horizon factors. Overall, our results offer a theoretically grounded alternative to popular (yet ad-hoc) methods like GRPO.
3. **Experiments.** On six mathematical-reasoning benchmarks—namely AIME24/25, AMC23, MATH500, Minerva-Math, and OlympiadBench—TBRM performs on par with or better than PPO and GRPO baselines. Notably, TBRM achieves up to 30.5% accuracy on AIME24 with Qwen2.5-Math-7B. Compared to GRPO, it improves the average benchmark score by 1.3% absolute, while under comparable conditions to PPO, it achieves better performance with 22.5% less training time and 33% lower GPU memory. We further demonstrate that TBRM benefits from additional rollouts and the model learns emergent reasoning patterns, such as verification, backtracking, and decomposition, that align with human mathematical practice.

Collectively, our findings suggest that value-based approaches offer a compelling alternative to policy gradient methods especially for enhancing mathematical reasoning capabilities in LLMs. By recasting value learning at the trajectory level, TBRM provides a *principled, efficient, and theoretically grounded* approach for improving performance on mathematical reasoning tasks while dramatically reducing computational requirements.

2 Preliminaries

This section provides the necessary background for our work. We first review the fundamentals of KL-regularized RL (Section 2.1), then discuss prominent reinforcement learning algorithms applied to

large language models (Section 2.2), and finally introduce the autoregressive function approximation framework (Section 2.3) that serves as the foundation for our proposed approach.

2.1 KL-Regularized Reinforcement Learning

Reinforcement learning (RL) provides a framework for sequential decision-making problems where an agent interacts with an environment to maximize cumulative rewards. In the context of Markov Decision Processes (MDPs), which provide the theoretical foundation for RL, we consider an episodic finite-horizon framework. Formally, a horizon- H episodic MDP $M = (H, \mathcal{S}, \mathcal{A}, \mathcal{P}, r, \rho)$ consists of a (potentially very large) state space $\mathcal{S}$, an action space $\mathcal{A}$, a probability transition function $\mathcal{P} : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$, a reward function $r : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$, and an initial state distribution $\rho \in \Delta(\mathcal{S})$. The state space is typically layered such that $\mathcal{S} = \mathcal{S}_1 \cup \mathcal{S}_2 \cup \dots \cup \mathcal{S}_H$, where $\mathcal{S}_h$ is the set of states reachable at step h . A policy $\pi : \mathcal{S} \rightarrow \Delta(\mathcal{A})$ maps states to distributions over actions and induces a distribution over trajectories $\tau = (s_1, a_1, \dots, s_H, a_H)$ and rewards $(r_1, \dots, r_H)$, where the initial state is sampled as $s_1 \sim \rho$, and for $h = 1, \dots, H$: $a_h \sim \pi(s_h)$, $r_h = r(s_h, a_h)$, and $s_{h+1} \sim \mathcal{P}(s_h, a_h)$. We let $\mathbb{E}_{\tau \sim \pi}[\cdot]$ and $\mathbb{P}_{\tau \sim \pi}[\cdot]$ denote expectation and probability under this process, and $\mathbb{E}_\pi[\cdot]$ and $\mathbb{P}_\pi[\cdot]$ for brevity when τ is not explicitly mentioned.

For any policy π , we define the occupancy measures that characterize the probabilities of visiting states and selecting actions when following π . Specifically, the state occupancy measure $d^\pi(s_h) := \mathbb{P}_{\tau \sim \pi}[s_h \in \tau]$ represents the probability of visiting state $s_h \in \mathcal{S}_h$ under policy π . Similarly, the state-action occupancy measure $d^\pi(s_h, a_h) := \mathbb{P}_{\tau \sim \pi}[(s_h, a_h) \in \tau]$ gives the probability of the state-action pair (s_h, a_h) occurring in a trajectory. We also define the trajectory occupancy measure $d^\pi(\tau) := \mathbb{P}_{\tau' \sim \pi}[\tau' = \tau]$, which is the probability of generating the exact trajectory τ when following policy π . It is important to distinguish between $d^\pi(\tau)$ and $\pi(\tau) := \prod_{(s_h, a_h) \in \tau} \pi(a_h | s_h)$, as they differ when the transition dynamics are stochastic.

In standard RL, the objective is to find a policy π that maximizes the expected cumulative reward $J(\pi) = \mathbb{E}_{\tau \sim \pi}[r(\tau)]$, where $r(\tau) = \sum_{h=1}^H r(s_h, a_h)$. In many practical applications, particularly in the context of large language models, it is beneficial to incorporate a regularization term that encourages the learned policy to stay close to a reference policy π_{ref} . This leads to the KL-regularized RL objective (Ziebart et al., 2008; Ziebart, 2010; Neu et al., 2017; Ouyang et al., 2022)

$$J_\beta(\pi) = \mathbb{E}_{\tau \sim \pi}[r(\tau)] - \beta \cdot \mathbb{E}_{\tau \sim \pi} \left[\log \frac{\pi(\tau)}{\pi_{\text{ref}}(\tau)} \right] = \mathbb{E}_{\tau \sim \pi} \left[\sum_{h=1}^H \left(r(s_h, a_h) - \beta \log \frac{\pi(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right) \right],$$

where $\beta > 0$ is a regularization parameter that controls the strength of the penalty $D_{\text{KL}}(\pi \| \pi_{\text{ref}}) = \mathbb{E}_{\tau \sim \pi} \left[\log \frac{\pi(\tau)}{\pi_{\text{ref}}(\tau)} \right]$, known as the Kullback-Leibler divergence.

2.2 Reinforcement Learning Methods for Large Language Models

In this section, we briefly review popular reinforcement learning methods for large language models. For the sake of generality, we continue to use the terminology of MDPs (i.e., we use s to represent state and a to represent action). This terminology naturally encompasses the case of large language models in both single-turn and multi-turn interaction settings.

In the single-turn setting where $x \sim \rho$ denotes the input prompt and $y_1, y_2, \dots, y_H$ denote the output tokens, we can define $s_1 := x$ and $s_h := (x, y_1, \dots, y_{h-1})$ for $h > 1$, with $a_h := y_h$ for $h = 1, \dots, H$. In the multi-turn setting, which consists of multiple interaction turns $(x^{(1)}, y_{1:H}^{(1)})$, $(x^{(2)}, y_{1:H}^{(2)})$, and so forth, we can adapt the transition function accordingly. Here, $y_{1:H}^{(i)}$ is a shorthand notation for the sequence of tokens $y_1^{(i)}, y_2^{(i)}, \dots, y_H^{(i)}$ in the i -th turn. For instance, if a state-action pair (s, a) contains the complete response for one turn (e.g., in a conversation with three or more turns), where $s = (x^{(1)}, y_{1:H}^{(1)}, x^{(2)}, y_{1:H}^{(2)})$ and $a = y_H^{(2)}$, the next state would transition to $s' = (x^{(1)}, y_{1:H}^{(1)}, x^{(2)}, y_{1:H}^{(2)}, x^{(3)})$, rather than simply concatenating the previous state and action as in the single-turn case.

Proximal Policy Optimization (PPO). PPO (Schulman et al., 2017b) introduces a clipped surrogate objective to constrain policy updates:

$$\mathcal{J}^{\text{PPO}}(\theta) = \mathbb{E}_{(s_h, a_h) \sim \pi_{\theta_{\text{old}}}} \left[\min \left(\frac{\pi_\theta(a_h | s_h)}{\pi_{\theta_{\text{old}}}(a_h | s_h)} \hat{A}_h(s_h, a_h), \text{clip} \left(\frac{\pi_\theta(a_h | s_h)}{\pi_{\theta_{\text{old}}}(a_h | s_h)}, 1 - \varepsilon, 1 + \varepsilon \right) \hat{A}_h(s_h, a_h) \right) \right],$$

where $\hat{A}_h$ is the advantage estimate, and ε is a hyperparameter. The advantage $\hat{A}_h$ is typically computed using Generalized Advantage Estimation (GAE; Schulman et al., 2015): $\hat{A}_h(s_h, a_h) =$

$\sum_l (\lambda)^l \delta_{t+l}$, where $\delta_h = r_h + V_\phi(s_{h+1}) - V_\phi(s_h)$ is the temporal difference error and V_ϕ is an estimate of the value function of the KL regularized reward $r(s_h, a_h) - \beta \log \frac{\pi(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)}$.

For LLMs, PPO has been widely used for enhancing mathematical reasoning capabilities, where objective rewards signal the correctness of the model’s solutions.

Group Relative Policy Optimization (GRPO). GRPO (Shao et al., 2024) is a policy-based method that, in practical implementations for LLMs like DeepSeek-R1, samples G responses $o^{(1)}, \dots, o^{(G)}$ for each prompt x and computes advantages by normalizing rewards within each prompt group. In the MDP terminology above, this corresponds to: $s_1^{(i)} = x \sim \rho$, $o^{(i)} = (a_1^{(i)}, a_2^{(i)}, \dots, a_{|o^{(i)}|}^{(i)})$, $s_h^{(i)} = (x, a_1^{(i)}, \dots, a_{h-1}^{(i)})$, and $r(x, o^{(i)}) = \sum_h r(s_h^{(i)}, a_h^{(i)})$. The advantage for the i -th response $o^{(i)}$ (and implicitly for each token within that response) is computed as: $\hat{A}^{(i)} = \frac{r(x, o^{(i)}) - \text{mean}(\{r(x, o^{(1)}), \dots, r(x, o^{(G)})\})}{\text{std}(\{r(x, o^{(1)}), \dots, r(x, o^{(G)})\})}$, where $r(x, o^{(i)})$ is the outcome for response $o^{(i)}$ to prompt x as we defined above. This response-level advantage $\hat{A}^{(i)}$ is then used to replace the step-wise advantage function $\hat{A}_h(s_h, a_h)$ in the PPO objective $\mathcal{J}^{\text{PPO}}$, but then GRPO objective accommodates the KL-regularization at the end:

$$\mathcal{J}^{\text{GRPO}}(\theta) = \mathbb{E}_{x \sim \rho, \{o^{(i)}\}_{i=1}^G \sim \pi_{\theta_{\text{old}}}(\cdot | x)} \left[\frac{1}{G} \sum_{i=1}^G \frac{1}{|o_i|} \sum_{h=1}^{|o_i|} \left\{ \min \left[\frac{\pi_\theta(a_h^{(i)} | s_h^{(i)})}{\pi_{\theta_{\text{old}}}(a_h^{(i)} | s_h^{(i)})} \hat{A}^{(i)}, \text{clip} \left(\frac{\pi_\theta(a_h^{(i)} | s_h^{(i)})}{\pi_{\theta_{\text{old}}}(a_h^{(i)} | s_h^{(i)})}, 1 - \varepsilon, 1 + \varepsilon \right) \hat{A}^{(i)} \right] - \beta D_{\text{KL}}(\pi \| \pi_{\text{ref}}) \right\} \right].$$

The normalization by mean and standard deviation is intended to stabilize training by reducing variance. GRPO is often considered a simpler alternative to PPO for post-training LLMs. This is partly because PPO typically involves training a separate critic network and incorporates more complex mechanisms for policy updates. In the context of LLMs, the full complexity of PPO might not always be necessary, leading to the adoption of more streamlined policy gradient methods like GRPO.

2.3 Autoregressive Function Approximation

Having established the principles of KL-regularized RL and reviewed current RL methods for LLMs, we now introduce the key formulation that bridges these concepts: autoregressive function approximation. This framework allows us to naturally parameterize value functions and policies using autoregressive models like LLMs, which is central to our proposed method.

Note that KL-regularized RL has been widely studied in classical RL literature (e.g., Schulman et al., 2017a; Nachum et al., 2017; Haarnoja et al., 2018) and the similar idea of autoregressive function approximation has also been introduced by Guo et al. (2022). This subsection should be viewed as a discussion of preliminary background and unified notations that will enable our proposed approach.

Given a reference model π_{ref} , we first define the following modified reward function

$$R_\beta(s_h, a_h) = \frac{r(s_h, a_h)}{\beta} + \log \pi_{\text{ref}}(a_h | s_h), \quad (1)$$

and, therefore, the original KL-regularized RL objective can be rewritten as: $J_\beta(\pi) = \beta \cdot \mathbb{E}_{\tau \sim \pi} [R_\beta(\tau) - \log \pi(\tau)] = \beta \cdot \mathbb{E}_{\tau \sim \pi} \left[\sum_{h=1}^H (R_\beta(s_h, a_h) - \log \pi(a_h | s_h)) \right]$. Here, KL-regularization is equivalent to entropy regularization. However, we will continue to use the KL-regularization terminology throughout the remainder of the paper for consistency.

The optimal policy for the above objective, denoted $\pi_\beta^* = \arg \max_\pi J_\beta(\pi)$, has a closed-form solution that takes the form of a softmax distribution,

$$\pi_\beta^*(a_h | s_h) \propto \pi_{\text{ref}}(a_h | s_h) \exp(Q_r^*(s_h, a_h)/\beta), \quad \text{or} \quad \pi_\beta^*(a_h | s_h) \propto \exp(Q_{R_\beta}^*(s_h, a_h)),$$

where Q_r^* and $Q_{R_\beta}^*$ are the optimal action-value functions for the original reward r and transformed reward functions R_β , respectively.

We now formalize the Bellman operators under the shifted reward function R_β , and discuss the key properties induced by the KL-regularization. For a given policy π and any Q-function Q , we define the Bellman operator as

$$(\mathcal{T}_\beta^\pi Q)(s_h, a_h) := R_\beta(s_h, a_h) + \mathbb{E}_{s_{h+1} \sim \mathcal{P}(\cdot | s_h, a_h), a_{h+1} \sim \pi(\cdot | s_{h+1})} [Q(s_{h+1}, a_{h+1}) - \log \pi(a_{h+1} | s_{h+1})].$$

Throughout this paper, we will use $V_Q(s) := \log \sum_{a \in \mathcal{A}} e^{Q(s,a)}$ to denote the softmax of the given Q . With this, the Bellman optimality operator becomes

$$\begin{aligned} (\mathcal{T}_\beta Q)(s_h, a_h) &:= R_\beta(s_h, a_h) + \mathbb{E}_{s_{h+1} \sim \mathcal{P}(\cdot | s_h, a_h)} \left[\underbrace{\max_{\pi \in \Delta(\mathcal{A})} \mathbb{E}_{a_{h+1} \sim \pi(\cdot | s_{h+1})} [Q(s_{h+1}, a_{h+1}) - \log \pi(a_{h+1} | s_{h+1})]}_{=\log \sum_{a \in \mathcal{A}} \exp(Q(s_{h+1}, a)) =: V_Q(s_{h+1})} \right] \\ &= R_\beta(s_h, a_h) + \mathbb{E}_{s_{h+1} \sim \mathcal{P}(\cdot | s_h, a_h)} [V_Q(s_{h+1})]. \end{aligned}$$

With the above definitions, the optimal policy takes the form of the following Boltzmann distribution:

$$\pi_Q(\cdot | s_h) := \operatorname{argmax}_{\pi \in \Delta(\mathcal{A})} \mathbb{E}_{a_h \sim \pi(\cdot | s_h)} [Q(s_h, a_h) - \log \pi(a_h | s_h)] = e^{Q(s_h, \cdot) - V_Q(s_h)}.$$

We can further define the KL-regularized value functions for a given policy π :

$$\begin{aligned} Q_\beta^\pi(\tilde{s}_h, \tilde{a}_h) &:= \mathbb{E}_\pi \left[R_\beta(s_h, a_h) + \sum_{h'=h+1}^H (R_\beta(s_{h'}, a_{h'}) - \log \pi(a_{h'} | s_{h'})) \mid (s_h, a_h) = (\tilde{s}_h, \tilde{a}_h) \right], \\ V_\beta^\pi(\tilde{s}_h) &:= \mathbb{E}_\pi \left[\sum_{h'=h}^H (R_\beta(s_{h'}, a_{h'}) - \log \pi(a_{h'} | s_{h'})) \mid s_h = \tilde{s}_h \right] = \mathbb{E}_{a_h \sim \pi(\cdot | \tilde{s}_h)} [Q_\beta^\pi(\tilde{s}_h, a_h) - \log \pi(a_h | \tilde{s}_h)]. \end{aligned}$$

The corresponding optimal value functions are $Q_\beta^* := Q_\beta^{\pi_\beta^*}$ and $V_\beta^* := V_\beta^{\pi_\beta^*}$. Here, Q_β^π and Q_β^* are also the fixed points of $\mathcal{T}_\beta^\pi$ and $\mathcal{T}_\beta$, respectively.

For autoregressive function approximation architectures, such as large language models, we can directly leverage logits to parameterize Q , V_Q , and π_Q as follows. Let θ be the model weights. Define

$$Q_\theta(s, a) := \operatorname{logit}_\theta(s, a), \quad V_\theta(s) := \operatorname{softmax} \circ \operatorname{logit}_\theta(s, \cdot), \quad \log \pi_\theta(a | s) := Q_\theta(s, a) - V_\theta(s), \quad (2)$$

where $\operatorname{softmax} \circ \operatorname{logit}_\theta(s, \cdot) := \log \sum_{a \in \mathcal{A}} \exp(\operatorname{logit}_\theta(s, a))$.² While we assumed here temperature to be 1 for simplicity, any temperature can be incorporated by appropriately scaling the reward Eq. (1).

3 Trajectory Bellman Residual Minimization

In this section, we introduce our main algorithm, **Trajectory Bellman Residual Minimization (TBRM)**, designed specifically for large language models problems. As we discussed in Section 2.1, the transition dynamics of large language models can be viewed as deterministic. *For the remainder of this section, we will apply the autoregressive function approximation defined in Eq. (2) and assume deterministic transition dynamics.*

Recall that Bellman error $Q(s_h, a_h) - (\mathcal{T}_\beta Q)(s_h, a_h)$ over state-action pairs (s_h, a_h) is employed as the proxy for controlling the performance $J_\beta(\pi_Q)$ of π_Q . Minimizing the square of Bellman error on (s_h, a_h) in deterministic MDPs is equivalent for minimizing the square of Bellman residual $Q(s_h, a_h) - R_\beta(s_h, a_h) - V_Q(s_{h+1})$ given the (s_h, a_h, s_{h+1}) tuple. This leads to the classical Bellman residual minimization objective (BRM; Schweitzer and Seidmann, 1985; Baird et al., 1995), which we expand using the definition of R_β and autoregressive function approximation in Eq. (2):

$$\begin{aligned} \mathcal{L}_\mathcal{D}^{\text{BRM}}(\theta) &= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \sum_{h=1}^H \left(\operatorname{logit}_\theta(s_h, a_h) - (\mathcal{T}_\beta \operatorname{logit}_\theta)(s_h, a_h) \right)^2 \\ &= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \sum_{h=1}^H \left(\operatorname{logit}_\theta(s_h, a_h) - \frac{r(s_h, a_h)}{\beta} - \log \pi_{\text{ref}}(s_h, a_h) - \log \sum_{a \in \mathcal{A}} \exp(\operatorname{logit}_\theta(s_{h+1}, a)) \right)^2. \end{aligned} \quad (3)$$

Here $\mathcal{D}$ denotes data which can be either purely offline or updated online as a replay buffer. In the context of LLMs, directly minimizing BRM may not be possible because the token-level reward signal is either unavailable (e.g., if we assign the outcome reward to the final token) or very sparse. Crucially, minimizing the square of per-step Bellman error as in $\mathcal{L}_\mathcal{D}^{\text{BRM}}(\theta)$ is sufficient but not necessary for maximization of $J_\beta(\pi_\theta)$. Indeed, a weaker control of Bellman errors over certain distributions is sufficient for optimizing $J_\beta(\pi_\theta)$ (see, e.g., Xie and Jiang (2020) or Corollary 4). As we prove below (Section 3.1), it is sufficient to instead consider a trajectory-based variant of BRM,

$$\mathcal{L}_\mathcal{D}^{\text{TBRM}}(\theta) = \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(\sum_{h=1}^H \operatorname{logit}_\theta(s_h, a_h) - (\mathcal{T}_\beta \operatorname{logit}_\theta)(s_h, a_h) \right)^2$$

²Here, softmax denotes the log-sum-exp operator for notational convenience, rather than the vocabulary softmax layer used to produce token probabilities in language models.

$$= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(\text{logit}_{\theta}(s_1, a_1) - \frac{r(\tau)}{\beta} - \log \pi_{\text{ref}}(\tau) + \sum_{h=2}^H \log \pi_{\theta}(a_h | s_h) \right)^2, \quad (4)$$

where the second equality used the autoregressive function approximation, i.e., $\log \pi_{\theta}(a_h | s_h) = \text{logit}_{\theta}(a_h | s_h) - \log \sum_{a \in \mathcal{A}} \exp(\text{logit}_{\theta}(s_h, a))$, to simplify the expression. An immediate benefit of TBRM is that we only require the trajectory outcome $r(\tau)$, rather than the process reward $r(s_h, a_h)$ as in BRM, and credit assignment will be (provably) carried out in the learning procedure itself. In our experiments, we update $\mathcal{D}$ online, and [Algorithm 1](#) represents the exact implementation we used in [Section 4](#).

Algorithm 1 Trajectory Bellman Residual Minimization (TBRM)

input: Task prompt dataset $\mathcal{D}_{\text{task}}$, reward function r , reward scaling coefficient β , reference policy π_{ref} with parameter θ_{ref} , and number of iterations T .

- 1: Initialize $\theta \leftarrow \theta_{\text{ref}}$.
- 2: **for** $t = 1$ to T **do**
- 3: Sample a batch $\mathcal{D}_{\text{task}}^t \subset \mathcal{D}_{\text{task}}$.
- 4: For each question $q \in \mathcal{D}_{\text{task}}^t$, sample a trajectory τ from policy π_{θ} with initial state $s_1 = q$, and collect these trajectories into dataset $\mathcal{D}_t$.
- 5: Update θ via gradient descent to minimize $\mathcal{L}_{\mathcal{D}_t}^{\text{TBRM}}(\theta)$ as defined in [Eq. \(4\)](#):

$$\nabla_{\theta} \frac{1}{|\mathcal{D}_t|} \sum_{\tau \in \mathcal{D}_t} \left(\text{logit}_{\theta}(s_1, a_1) - \frac{r(\tau)}{\beta} - \log \pi_{\text{ref}}(\tau) + \sum_{h=2}^H \log \pi_{\theta}(a_h | s_h) \right)^2,$$

where $\text{logit}_{\theta}(s_1, a_1)$ is the raw logit of the first output token.

- 6: **end for**
 - 7: **return** π_{θ} .
-

As we formally prove in [Section 3.1](#), TBRM is a fully *off-policy* algorithm with a directly optimizable objective, and it provably converges to a near-optimal policy with any off-policy data (though the degree of off-policyness may affect sample efficiency, [Eq. \(6\)](#)). In contrast, policy-based counterparts are usually on-policy in nature: policy-gradient-based algorithms (like REINFORCE) require to sample new trajectories in an on-policy manner. PPO-like algorithms require on-policy actions for their actor components but optimize a surrogate loss instead. PPO’s critic update also requires on-policy rollouts. In contrast, the off-policy nature of TBRM removes the need for additional techniques such as importance sampling ratios, clipping, critic models, or (multiple) on-policy rollouts.

Readers may question: *why hasn’t this simple variant of BRM received attention in the literature?* We conjecture that TBRM has been suspected to suffer from the curse of horizon, at least from the theoretical perspective. In the theory of offline RL (e.g., [Chen and Jiang, 2019](#); [Xie and Jiang, 2020](#)), a key technique is to control the expected Bellman error $|\mathbb{E}_{\pi}[\sum_h ((\mathcal{T}_{\beta}Q)(s_h, a_h) - Q(s_h, a_h))]|$ on a certain (unavailable) distribution d^{π} by instead minimizing the per-step squared Bellman error $\sum_h \mathbb{E}_{\mu}[(\mathcal{T}_{\beta}Q)(s_h, a_h) - Q(s_h, a_h)]^2$ on the data distribution d^{μ} generated by μ . This step only incurs the cost of the state-wise distribution-shift $\frac{d^{\pi}(s_h, a_h)}{d^{\mu}(s_h, a_h)}$. When it comes to trajectory-level data, minimization of the square of expected Bellman error $\mathbb{E}_{\mu}[(\sum_h [(\mathcal{T}_{\beta}Q)(s_h, a_h) - Q(s_h, a_h)])^2]$, as in TBRM, would appear to incur the trajectory-level distribution-shift cost $\Pi_h \frac{\pi(a_h | s_h)}{\mu(a_h | s_h)}$, and thus possibly cause an exponential blow-up with horizon H compared to the state-wise case. However, the recent theoretical results ([Jia et al., 2025](#)) challenge this conventional wisdom and indicate that the Markov property can be the key to avoiding trajectory-level distribution-shift when conducting trajectory-level change of measure. In [Section 3.1](#), we formally prove that TBRM indeed only incurs state-wise distribution-shift regardless of its trajectory-level objective, and show that TBRM can efficiently converge to a near-optimal policy with finite-sample analysis.

Comparison with other related algorithms. We note that algorithms with similar structure to TBRM have been derived previously from diverse perspectives in both deep RL ([Haarnoja et al., 2017](#); [Schulman et al., 2017a](#); [Nachum et al., 2017](#); [Haarnoja et al., 2018](#)) and LLM applications ([Guo et al., 2022](#); [Ethayarajh et al., 2024](#); [Team et al., 2025](#); [Ji et al., 2024](#); [Wang et al., 2024](#)). This convergence is unsurprising, as TBRM and related algorithms fundamentally aim to minimize Bellman error, albeit through different formulations and optimization approaches. However, to the

best of our knowledge, TBRM is the only optimization algorithm (rather than iterative ones like Q-learning; [Appendix F](#) demonstrates the benefit of optimization over iteration) that requires only one rollout per prompt among all of these approaches. [Appendix E](#) provides a detailed comparison of TBRM with other related algorithms. The present paper formally establishes finite-sample guarantees for TBRM.

3.1 Theoretical Analysis of TBRM

We use Θ to denote the parameter space, equipped with norm $\|\cdot\|$. We assume the following standard realizability condition, which can be relaxed to hold approximately (see, e.g., [Cheng et al., 2022](#)).

Assumption 1 (Realizability). *There exists $\theta^* \in \Theta$ such that $Q_{\theta^*} = Q^*$.*

Motivation. We first show that θ^* is the population minimizer of the TBRM loss (4) through Bellman equation. Under the parametrization (2), $Q_{\theta^*}(s, a) = \text{logit}_{\theta^*}(s, a)$ is the optimal soft Q-function for the transformed reward function $R_\beta(s, a) = \frac{r(s, a)}{\beta} + \log \pi_{\text{ref}}(a | s)$, and the optimal value function is given by $V_{\theta^*}(s) = Q_{\theta^*}(s, a) - \log \pi_{\theta^*}(a | s)$. Therefore, the Bellman equation becomes (deterministically for a trajectory τ drawn from the MDP)

$$Q_{\theta^*}(s_h, a_h) = R_\beta(s_h, a_h) + V_{\theta^*}(s_{h+1}) = R_\beta(s_h, a_h) + Q_{\theta^*}(s_{h+1}, a_{h+1}) - \log \pi_{\theta^*}(a_{h+1} | s_{h+1}).$$

Then, summing over $h = 1, 2, \dots, H-1$ for any admissible trajectory τ , we have

$$\begin{aligned} 0 &\equiv \sum_{h=1}^{H-1} [Q_{\theta^*}(s_h, a_h) - R_\beta(s_h, a_h) - Q_{\theta^*}(s_{h+1}, a_{h+1}) + \log \pi_{\theta^*}(a_{h+1} | s_{h+1})] \\ &= \text{logit}_{\theta^*}(s_1, a_1) - R(\tau) + \sum_{h=2}^H \log \pi_{\theta^*}(a_h | s_h), \end{aligned}$$

where $R(\tau) = \frac{r(\tau)}{\beta} + \log \pi_{\text{ref}}(\tau)$ is the trajectory transformed reward. Hence, it holds that $\mathcal{L}_D^{\text{TBRM}}(\theta^*) \equiv 0$ deterministically, and any approximate minimizer of the loss $\mathcal{L}_D^{\text{TBRM}}$ must also attain low trajectory Bellman residual.

The analysis above establishes a necessary condition for the optimal soft Q-function $Q_{\theta^*}(s, a)$ or $\text{logit}_{\theta^*}(s, a)$. Beyond this, the sub-optimality of a policy π_θ can also be related to the trajectory Bellman residual: through our analysis in [Appendix B](#), we can show that for any $\theta \in \Theta$,

$$J_\beta(\pi^*) - J_\beta(\pi_\theta) \leq 2\beta \max_{\pi \in \{\pi^*, \pi_\theta\}} \left| \mathbb{E}_\pi \left[\text{logit}_\theta(s_1, a_1) - R(\tau) + \sum_{h=2}^H \log \pi_\theta(a_h | s_h) \right] \right|.$$

Therefore, it remains to relate the expected trajectory Bellman residual under the off-policy distribution induced by μ and any policy π through change-of-trajectory-measure.

Change-of-trajectory-measure. A key to our analysis is the following improved version of the *change-of-trajectory-measure lemma* ([Jia et al., 2025](#)). Let $\chi^2(P \parallel Q) = \mathbb{E}_P[dP/dQ] - 1$ be the χ^2 -divergence. Let $d_h^\pi(\cdot), d_h^\mu(\cdot) \in \Delta(\mathcal{S}_h \times \mathcal{A})$ denote the occupancy measures of Markovian policies π and μ .

Lemma 1 (Change-of-Trajectory-Measure Lemma). *Given an MDP $M = (H, \mathcal{S}, \mathcal{A}, \mathcal{P}, r, \rho)$ and a policy π ,*

$$\sup_f \frac{(\mathbb{E}_\pi [\sum_{h=1}^H f(s_h, a_h)])^2}{\mathbb{E}_\mu [\left(\sum_{h=1}^H f(s_h, a_h)\right)^2]} \leq 1 + \sum_{h=1}^H \chi^2(d_h^\pi \parallel d_h^\mu), \quad (5)$$

where the supremum is over all measurable functions $f : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$.

The proof of [Lemma 1](#) significantly simplifies the one in [Jia et al. \(2025\)](#), as shown in [Appendix C](#). As a direct corollary, the RHS of [Eq. \(5\)](#) can be further upper bounded by $H \cdot \max_{h,s,a} \frac{d_h^\pi(s_h, a_h)}{d_h^\mu(s_h, a_h)}$, improving upon [Jia et al. \(2025\)](#) by a factor of H^2 .

Our main result will be stated in terms of the following *concentrability coefficient* of the data collection policy μ :

$$C_{\text{conc}}(\mu) := 1 + \max_{\theta \in \Theta} \max_{h \in [H]} \chi^2(d_h^{\pi_\theta} \parallel d_h^\mu), \quad (6)$$

a notion weaker than the commonly-studied L_∞ -concentrability $C_{\text{conc},\infty}(\mu) := \max_{\pi,h,s,a} \frac{d_h^\pi(s,a)}{d_h^\mu(s,a)}$.

The proof of the following theorem is deferred to [Appendix D](#).

Theorem 2. Suppose $\hat{\theta}$ is a parameter that satisfies $\mathcal{L}_D^{\text{TBRM}}(\hat{\theta}) - \inf_{\theta \in \Theta} \mathcal{L}_D^{\text{TBRM}}(\theta) \leq \varepsilon_{\text{opt}}$, data $\mathcal{D}$ are i.i.d. according to μ , and [Assumptions 1, 2 and 3](#) hold. Then, with high probability, it holds that

$$J_\beta(\pi^*) - J_\beta(\pi_{\hat{\theta}}) \leq \tilde{O}\left(\sqrt{HC_{\text{conc}}(\mu)\left(\beta^2\varepsilon_{\text{opt}} + \frac{H^2 \dim(\Theta)}{|\mathcal{D}|}\right)}\right),$$

where $\dim(\Theta)$ is the measure of the dimension of Θ defined in [Assumption 3](#).

4 Experiments

In this section, we present experiments to evaluate the performance of TBRM on reasoning tasks. We compare TBRM against two policy-based methods: GRPO and PPO. The codebase for the experiments is publicly available at <https://github.com/r1x-lab/TBRM>.

4.1 Experiment Setup

Datasets and models. We train our models using the prompt set from DAPO (Yu et al., 2025, Apache license 2.0), which comprises approximately 17.4k math problems sourced from the AoPS³ website and official competition homepages. All problems are standardized to have integer answers. To demonstrate the generality of our method across model scales, we conduct experiments using Qwen2.5-Math-1.5B and Qwen2.5-Math-7B.

Evaluation. We assess the models’ reasoning abilities on several standard math benchmarks: AIME24, AIME25, AMC23, MATH500 (Hendrycks et al., 2021), Minerva Math (Lewkowycz et al., 2022), and OlympiadBench (He et al., 2024). For MATH500, Minerva Math, and OlympiadBench, we generate a single response per problem and report the overall accuracy, denoted as **Avg@1**. For the smaller benchmarks AIME24, AIME25, and AMC23, where performance can fluctuate due to limited data, we generate 32 responses per problem and average the accuracies to mitigate the intrinsic randomness of LLM outputs; this metric is denoted as **Avg@32**. Responses are sampled with temperature 0 for **Avg@1** and temperature 1.0 for **Avg@32**. We employ Math-Verify (Kydlíček, 2025, Apache-2.0 license) as the verifier.

Implementation details. We implement our methods and baselines using the VERL framework (Sheng et al., 2024, Apache-2.0 license), following most of the recommended hyperparameter settings for GRPO and PPO. To balance performance and efficiency, we use a prompt batch size of 128 and a response length of 2048 tokens per training step. For PPO, we generate one response per prompt ($n = 1$), while for GRPO, which requires multiple rollouts, we generate four responses ($n = 4$). On TBRM, we experiment with both settings ($n = 1$ and $n = 4$). All responses are sampled with a temperature of 1.0. For TBRM, we set $\beta = 0.002$ across all experiments. All models are trained for the same number of steps. More details of our implementation can be found in [Appendix G](#).

4.2 Main Results

The effectiveness of our algorithm is demonstrated in [Table 1](#). Across six challenging math benchmarks, TBRM consistently matches or surpasses its comparable baselines. Specifically, with a single rollout per prompt, $\text{TBRM}_{n=1}$ achieves higher accuracies than $\text{PPO}_{n=1}$ on most benchmarks and matches the performance of $\text{GRPO}_{n=4}$, despite the latter using four times as many samples during training. Notably, our Qwen2.5-Math-1.5B-based model attains 13.2% accuracy on AIME24, outperforming both the Qwen2.5-Math-7B base model and the 1.5B $\text{GRPO}_{n=4}$ model. When increasing the number of sampled responses to four, $\text{TBRM}_{n=4}$ surpasses baselines by a larger margin. On AIME24, our 1.5B model reaches 14.3% accuracy, while the 7B model further advances to 30.5%, exceeding $\text{GRPO}_{n=4}$ by 1.6%. Additional results with more rollouts are presented in [Appendix H.5](#).

³<https://artofproblemsolving.com/>

Method	AIME24	AIME25	AMC23	MATH500	Minerva Math	OlympiadBench
	Avg@32	Avg@32	Avg@32	Avg@1	Avg@1	Avg@1
Qwen2.5-Math-1.5B	5.0	1.9	24.9	63.4	16.5	30.8
PPO $n=1$	11.4	4.5	46.6	72.2	26.8	<u>36.0</u>
TBRM $n=1$	<u>13.2</u>	<u>5.6</u>	<u>48.6</u>	72.2	<u>27.2</u>	35.7
GRPO $n=4$	13.0	7.1	49.9	71.2	28.7	37.5
TBRM $n=4$	14.3	6.9	52.0	72.2	30.5	36.1
Qwen2.5-Math-7B	10.6	2.8	31.6	67.4	13.2	29.3
PPO $n=1$	<u>25.4</u>	13.2	<u>63.4</u>	76.4	33.8	39.3
TBRM $n=1$	24.1	<u>13.2</u>	<u>63.4</u>	<u>78.6</u>	36.4	<u>41.5</u>
GRPO $n=4$	28.9	10.7	66.8	79.8	36.0	42.5
TBRM $n=4$	30.5	13.1	68.4	79.8	36.4	44.1

Table 1: Performance of various methods on math benchmarks, where n denotes the number of responses sampled per prompt during training. For each benchmark, the highest accuracy across all methods is bolded, and the highest accuracy among methods with $n=1$ is underscored.

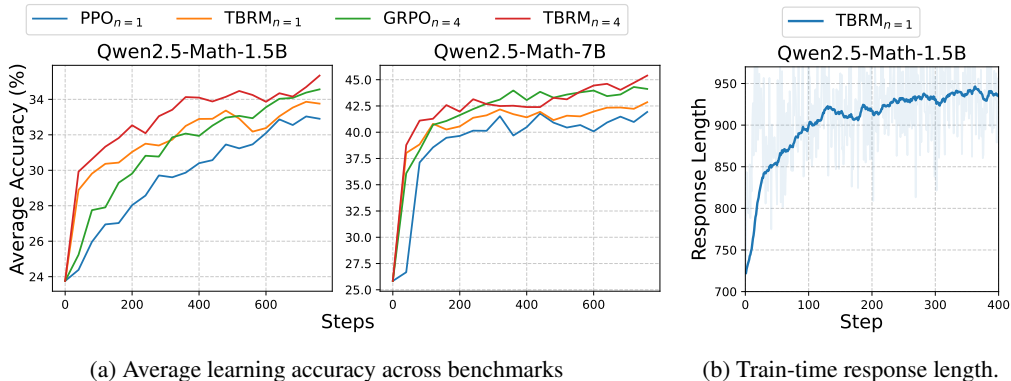


Figure 1: [Figure 1\(a\)](#) shows average learning accuracy across benchmarks for PPO, GRPO, and TBRM. Per-benchmark results can be found in [Appendix H.1](#). [Figure 1\(b\)](#) demonstrates model’s ability to engage in extended reasoning at test time with R1 template.

[Figure 1\(a\)](#) illustrates the step-wise average performance across all benchmarks for TBRM and baseline methods. While all approaches demonstrate improved reasoning with increased training data, TBRM consistently exhibits a superior convergence rate and achieves higher absolute performance than its counterparts. Notably, TBRM_{n=4} attains the highest performance throughout nearly the entire training duration. Furthermore, TBRM_{n=1} outperforms PPO_{n=1} and performs similarly to GRPO_{n=4}, especially on 1.5B model, with only a mild gap in their results.

4.3 Training Dynamics and Performance Analysis

Reward. We present the training reward curves in [Figure 2\(a\)](#), which shows that TBRM achieves comparable reward levels to its baselines. Furthermore, TBRM demonstrates a faster convergence rate during early training. This is particularly evident with the 1.5B model, where TBRM attains significantly higher rewards than PPO and GRPO.

Response length. Prior work has shown that reinforcement learning can enhance a model’s ability to solve increasingly complex reasoning tasks by leveraging extended test-time computation, as reflected in progressively longer responses during training ([Guo et al., 2025](#); [Zeng et al., 2025](#); [Liu et al., 2025](#)). We find that TBRM exhibits a similar capability. Specifically, we adopt the prompt template from DeepSeek-R1 ([Guo et al., 2025](#)) and apply TBRM training to Qwen2.5-Math-1.5B. Following previous studies ([Zeng et al., 2025](#)), we include only responses that terminate under normal conditions—excluding those truncated due to length limits—as truncated outputs often suffer from repetition and incompleteness. As illustrated in [Figure 1\(b\)](#), TBRM encourages the model to explore

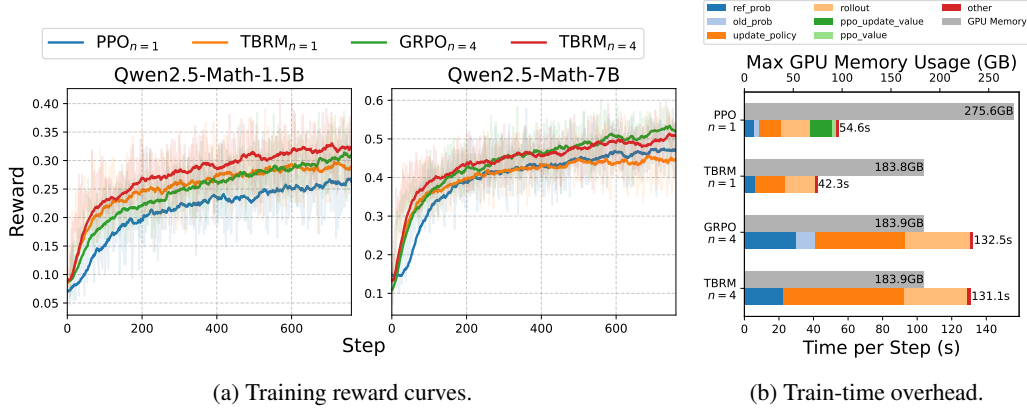


Figure 2: Figure 2(a) demonstrates the increment of rewards during training. Figure 2(b) is a comparison of maximal GPU memory consumption and per-step time cost across different methods with base model Qwen2.5-Math-7B. Time cost is segmented into key partitions, with time partition labels defined in Appendix G.2 (Appendix G.2).

and refine its reasoning more deeply over time, enabling models to take advantage of extended computation at test time to improve reasoning performance.

Training efficiency. TBRM offers substantial implementation simplicity relative to existing approaches. Specifically, it eliminates the need for critic models V_ϕ , as required by PPO, and avoids the necessity of sampling multiple responses per prompt ($n > 1$), as in GRPO. Moreover, due to its fully off-policy nature, TBRM does not require multiple updates per training step. We conduct experiments using Qwen2.5-Math-7B under consistent training conditions (see Appendix G.2 for details) and report the peak GPU memory usage and wall-clock time per training step for each method in Figure 2(b). Overall, TBRM exhibits matched or lower resource consumption compared to its counterparts. When $n = 1$, TBRM _{$n=1$} uses 33.3% less GPU memory than PPO _{$n=1$} and achieves a $1.3\times$ speedup. For $n = 4$, TBRM _{$n=4$} demonstrates comparable resource usage to GRPO _{$n=4$} . Notably, TBRM _{$n=1$} achieves a $3.1\times$ training speedup relative to GRPO _{$n=4$} , while yielding similar performance despite sampling only a single response—GRPO _{$n=4$} outperforms TBRM _{$n=1$} by only 0.80% with the 1.5B model and 1.27% with the 7B model on average across math benchmarks.

Extended Analysis. We further analyze the responses generated by the TBRM models and identify several notable reasoning patterns, including verification, backtracking, decomposition, and enumeration. Illustrative examples of these patterns are provided in Appendix H.2. In Appendix H.3, we compare TBRM with the classical, token-level BRM formulation and show that directly applying BRM to LLMs leads to unstable training and reward collapse, underscoring the importance of the trajectory-level design. In Appendix H.4, we evaluate TBRM on a suite of non-mathematical reasoning tasks, demonstrating its ability to generalize beyond the mathematical domain.

5 Conclusion

In this paper, we present TBRM, a simple yet effective value-based RL algorithm for enhancing LLM reasoning capabilities. TBRM operates efficiently with just a single rollout per prompt and employs a lightweight optimization objective, eliminating the need for critics, importance-sampling ratios, or clipping mechanisms that are commonly required in policy-based approaches. Our theoretical analysis demonstrates that TBRM is guaranteed to converge to a near-optimal policy using off-policy data, while our empirical results show its effectiveness across standard mathematical reasoning benchmarks. We hope this work may contribute to expanding interest in value-based approaches for LLM reasoning, potentially complementing the policy-based algorithms that currently dominate LLM post-training methods.

Acknowledgements

We acknowledge support of the Simons Foundation and the NSF through awards DMS-2031883 and PHY-2019786, ARO through award W911NF-21-1-0328, and the DARPA AIQ award.

References

- Arash Ahmadian, Chris Cremer, Matthias Gallé, Marzieh Fadaee, Julia Kreutzer, Olivier Pietquin, Ahmet Üstün, and Sara Hooker. Back to basics: Revisiting reinforce style optimization for learning from human feedback in llms. *arXiv preprint arXiv:2402.14740*, 2024.
- Mohammad Gheshlaghi Azar, Zhaohan Daniel Guo, Bilal Piot, Remi Munos, Mark Rowland, Michal Valko, and Daniele Calandriello. A general theoretical paradigm to understand learning from human preferences. In *International Conference on Artificial Intelligence and Statistics*, pages 4447–4455. PMLR, 2024.
- 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.
- Leemon Baird et al. Residual algorithms: Reinforcement learning with function approximation. In *Proceedings of the twelfth international conference on machine learning*, pages 30–37, 1995.
- Marc G Bellemare, Will Dabney, and Rémi Munos. A distributional perspective on reinforcement learning. In *International conference on machine learning*, pages 449–458. PMLR, 2017.
- Richard Bellman. A markovian decision process. *Journal of mathematics and mechanics*, pages 679–684, 1957.
- Dimitri Bertsekas and John N Tsitsiklis. *Neuro-dynamic programming*. Athena Scientific, 1996.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901, 2020.
- Huayu Chen, Kaiwen Zheng, Qingsheng Zhang, Ganqu Cui, Yin Cui, Haotian Ye, Tsung-Yi Lin, Ming-Yu Liu, Jun Zhu, and Haoxiang Wang. Bridging supervised learning and reinforcement learning in math reasoning. *arXiv preprint arXiv:2505.18116*, 2025.
- Jinglin Chen and Nan Jiang. Information-theoretic considerations in batch reinforcement learning. In *International conference on machine learning*, pages 1042–1051. PMLR, 2019.
- Ching-An Cheng, Tengyang Xie, Nan Jiang, and Alekh Agarwal. Adversarially trained actor critic for offline reinforcement learning. In *International Conference on Machine Learning*, pages 3852–3878. PMLR, 2022.
- Kawin Ethayarajh, Winnie Xu, Niklas Muennighoff, Dan Jurafsky, and Douwe Kiela. Kto: Model alignment as prospect theoretic optimization. *arXiv preprint arXiv:2402.01306*, 2024.
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- Han Guo, Bowen Tan, Zhengzhong Liu, Eric Xing, and Zhiting Hu. Efficient (soft) q-learning for text generation with limited good data. In *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 6969–6991, 2022.
- Tuomas Haarnoja, Haoran Tang, Pieter Abbeel, and Sergey Levine. Reinforcement learning with deep energy-based policies. In *International conference on machine learning*, pages 1352–1361. PMLR, 2017.
- Tuomas Haarnoja, Aurick Zhou, Pieter Abbeel, and Sergey Levine. Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. In *International conference on machine learning*, pages 1861–1870. PMLR, 2018.
- Hado Hasselt. Double q-learning. *Advances in neural information processing systems*, 23, 2010.
- Chaoqun He, Renjie Luo, Yuzhuo Bai, Shengding Hu, Zhen Leng Thai, Junhao Shen, Jinyi Hu, Xu Han, Yujie Huang, Yuxiang Zhang, et al. Olympiadbench: A challenging benchmark for promoting agi with olympiad-level bilingual multimodal scientific problems. *arXiv preprint arXiv:2402.14008*, 2024.

- Jujie He, Jiakai Liu, Chris Yuhao Liu, Rui Yan, Chaojie Wang, Peng Cheng, Xiaoyu Zhang, Fuxiang Zhang, Jiacheng Xu, Wei Shen, Siyuan Li, Liang Zeng, Tianwen Wei, Cheng Cheng, Bo An, Yang Liu, and Yahui Zhou. Skywork open reasoner series. <https://capricious-hydrogen-41c.notion.site/Skywork-Open-Reasoner-Series-1d0bc9ae823a80459b46c149e4f51680>, 2025. Notion Blog.
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. Measuring mathematical problem solving with the math dataset. *arXiv preprint arXiv:2103.03874*, 2021.
- Matteo Hessel, Joseph Modayil, Hado Van Hasselt, Tom Schaul, Georg Ostrovski, Will Dabney, Dan Horgan, Bilal Piot, Mohammad Azar, and David Silver. Rainbow: Combining improvements in deep reinforcement learning. In *Proceedings of the AAAI conference on artificial intelligence*, volume 32, 2018.
- Jian Hu. Reinforce++: A simple and efficient approach for aligning large language models. *arXiv preprint arXiv:2501.03262*, 2025.
- Jingcheng Hu, Yinmin Zhang, Qi Han, Daxin Jiang, Xiangyu Zhang, and Heung-Yeung Shum. Open-reasoner-zero: An open source approach to scaling up reinforcement learning on the base model, 2025. URL <https://arxiv.org/abs/2503.24290>.
- Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, et al. Openai o1 system card. *arXiv preprint arXiv:2412.16720*, 2024.
- Kaixuan Ji, Guanlin Liu, Ning Dai, Qingping Yang, Renjie Zheng, Zheng Wu, Chen Dun, Quanquan Gu, and Lin Yan. Enhancing multi-step reasoning abilities of language models through direct q-function optimization. *arXiv preprint arXiv:2410.09302*, 2024.
- Zeyu Jia, Alexander Rakhlin, and Tengyang Xie. Do we need to verify step by step? rethinking process supervision from a theoretical perspective. *arXiv preprint arXiv:2502.10581*, 2025.
- Wouter Kool, Herke van Hoof, and Max Welling. Buy 4 REINFORCE samples, get a baseline for free!, 2019. URL <https://openreview.net/forum?id=r1lgTGL5DE>.
- Hynek Kydlíček. Math-Verify: Math Verification Library. 2025. URL <https://github.com/huggingface/math-verify>.
- Aitor Lewkowycz, Anders Andreassen, David Dohan, Ethan Dyer, Henryk Michalewski, Vinay Ramasesh, Ambrose Slone, Cem Anil, Imanol Schlag, Theo Gutman-Solo, et al. Solving quantitative reasoning problems with language models. *Advances in Neural Information Processing Systems*, 35:3843–3857, 2022.
- Ziniu Li, Tian Xu, Yushun Zhang, Zhihang Lin, Yang Yu, Ruoyu Sun, and Zhi-Quan Luo. Remax: A simple, effective, and efficient reinforcement learning method for aligning large language models. *arXiv preprint arXiv:2310.10505*, 2023.
- Zhihang Lin, Mingbao Lin, Yuan Xie, and Rongrong Ji. Cppo: Accelerating the training of group relative policy optimization-based reasoning models. *arXiv preprint arXiv:2503.22342*, 2025.
- Zichen Liu, Changyu Chen, Wenjun Li, Penghui Qi, Tianyu Pang, Chao Du, Wee Sun Lee, and Min Lin. Understanding rl-zero-like training: A critical perspective. *arXiv preprint arXiv:2503.20783*, 2025.
- Volodymyr Mnih, Koray Kavukcuoglu, David Silver, Andrei A Rusu, Joel Veness, Marc G Bellemare, Alex Graves, Martin Riedmiller, Andreas K Fidjeland, Georg Ostrovski, et al. Human-level control through deep reinforcement learning. *nature*, 518(7540):529–533, 2015.
- Ofir Nachum, Mohammad Norouzi, Kelvin Xu, and Dale Schuurmans. Bridging the gap between value and policy based reinforcement learning. *Advances in neural information processing systems*, 30, 2017.

- Gergely Neu, Anders Jonsson, and Vicenç Gómez. A unified view of entropy-regularized markov decision processes. *arXiv preprint arXiv:1705.07798*, 2017.
- 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.
- Yury Polyanskiy and Yihong Wu. *Information theory: From coding to learning*. Cambridge university press, 2025.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36:53728–53741, 2023.
- Rafael Rafailov, Joey Hejna, Ryan Park, and Chelsea Finn. From r to q^* : Your language model is secretly a q -function. *arXiv preprint arXiv:2404.12358*, 2024.
- Pierre Harvey Richemond, Yunhao Tang, Daniel Guo, Daniele Calandriello, Mohammad Gheshlaghi Azar, Rafael Rafailov, Bernardo Avila Pires, Eugene Tarasov, Lucas Spangher, Will Ellsworth, et al. Offline regularised reinforcement learning for large language models alignment. *arXiv preprint arXiv:2405.19107*, 2024.
- Arthur L Samuel. Some studies in machine learning using the game of checkers. *IBM Journal of research and development*, 3(3):210–229, 1959.
- John Schulman, Philipp Moritz, Sergey Levine, Michael Jordan, and Pieter Abbeel. High-dimensional continuous control using generalized advantage estimation. *arXiv preprint arXiv:1506.02438*, 2015.
- John Schulman, Xi Chen, and Pieter Abbeel. Equivalence between policy gradients and soft q-learning. *arXiv preprint arXiv:1704.06440*, 2017a.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017b.
- Paul J Schweitzer and Abraham Seidmann. Generalized polynomial approximations in markovian decision processes. *Journal of mathematical analysis and applications*, 110(2):568–582, 1985.
- Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, YK Li, Y Wu, et al. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*, 2024.
- Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. Hybridflow: A flexible and efficient rlhf framework. *arXiv preprint arXiv: 2409.19256*, 2024.
- Satinder P Singh and Richard S Sutton. Reinforcement learning with replacing eligibility traces. *Machine learning*, 22(1):123–158, 1996.
- Zafir Stojanovski, Oliver Stanley, Joe Sharratt, Richard Jones, Abdulhakeem Adefioye, Jean Kaddour, and Andreas Köpf. Reasoning gym: Reasoning environments for reinforcement learning with verifiable rewards, 2025. URL <https://arxiv.org/abs/2505.24760>.
- Richard S Sutton. Learning to predict by the methods of temporal differences. *Machine learning*, 3: 9–44, 1988.
- Richard S Sutton, Andrew G Barto, et al. *Reinforcement learning: An introduction*. MIT press Cambridge, 1998.
- Yunhao Tang, Zhaohan Daniel Guo, Zeyu Zheng, Daniele Calandriello, Rémi Munos, Mark Rowland, Pierre Harvey Richemond, Michal Valko, Bernardo Ávila Pires, and Bilal Piot. Generalized preference optimization: A unified approach to offline alignment. *arXiv preprint arXiv:2402.05749*, 2024.

- Kimi Team, Angang Du, Bofei Gao, Bowei Xing, Changjiu Jiang, Cheng Chen, Cheng Li, Chenjun Xiao, Chenzhuang Du, Chonghua Liao, et al. Kimi k1. 5: Scaling reinforcement learning with llms. *arXiv preprint arXiv:2501.12599*, 2025.
- Hado Van Hasselt, Arthur Guez, and David Silver. Deep reinforcement learning with double q-learning. In *Proceedings of the AAAI conference on artificial intelligence*, volume 30, 2016.
- Martin J Wainwright. *High-dimensional statistics: A non-asymptotic viewpoint*, volume 48. Cambridge University Press, 2019.
- Huaijie Wang, Shibo Hao, Hanze Dong, Shenao Zhang, Yilin Bao, Ziran Yang, and Yi Wu. Offline reinforcement learning for llm multi-step reasoning. *arXiv preprint arXiv:2412.16145*, 2024.
- Christopher JCH Watkins and Peter Dayan. Q-learning. *Machine learning*, 8:279–292, 1992.
- Liang Wen, Yunke Cai, Fenrui Xiao, Xin He, Qi An, Zhenyu Duan, Yimin Du, Junchen Liu, Lifu Tang, Xiaowei Lv, et al. Light-rl: Curriculum sft, dpo and rl for long cot from scratch and beyond. *arXiv preprint arXiv:2503.10460*, 2025.
- Ronald J Williams. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine learning*, 8:229–256, 1992.
- Ronald J Williams and Jing Peng. Function optimization using connectionist reinforcement learning algorithms. *Connection Science*, 3(3):241–268, 1991.
- Tengyang Xie and Nan Jiang. Q* approximation schemes for batch reinforcement learning: A theoretical comparison. In *Conference on Uncertainty in Artificial Intelligence*, pages 550–559. PMLR, 2020.
- Tian Xie, Zitian Gao, Qingnan Ren, Haoming Luo, Yuqian Hong, Bryan Dai, Joey Zhou, Kai Qiu, Zhirong Wu, and Chong Luo. Logic-rl: Unleashing llm reasoning with rule-based reinforcement learning. *arXiv preprint arXiv:2502.14768*, 2025.
- Wei Xiong, Jiarui Yao, Yuhui Xu, Bo Pang, Lei Wang, Doyen Sahoo, Junnan Li, Nan Jiang, Tong Zhang, Caiming Xiong, et al. A minimalist approach to llm reasoning: from rejection sampling to reinforce. *arXiv preprint arXiv:2504.11343*, 2025.
- Qiyang Yu, Zheng Zhang, Ruofei Zhu, Yufeng Yuan, Xiaochen Zuo, Yu Yue, Tiantian Fan, Gaohong Liu, Lingjun Liu, Xin Liu, et al. Dapo: An open-source llm reinforcement learning system at scale. *arXiv preprint arXiv:2503.14476*, 2025.
- Yufeng Yuan, Qiyang Yu, Xiaochen Zuo, Ruofei Zhu, Wenyuan Xu, Jiaze Chen, Chengyi Wang, TianTian Fan, Zhengyin Du, Xiangpeng Wei, et al. Vapo: Efficient and reliable reinforcement learning for advanced reasoning tasks. *arXiv preprint arXiv:2504.05118*, 2025a.
- Yufeng Yuan, Yu Yue, Ruofei Zhu, Tiantian Fan, and Lin Yan. What’s behind ppo’s collapse in long-cot? value optimization holds the secret. *arXiv preprint arXiv:2503.01491*, 2025b.
- Weihao Zeng, Yuzhen Huang, Qian Liu, Wei Liu, Keqing He, Zejun Ma, and Junxian He. Simplerl-zoo: Investigating and taming zero reinforcement learning for open base models in the wild. *arXiv preprint arXiv:2503.18892*, 2025.
- Chong Zhang, Yue Deng, Xiang Lin, Bin Wang, Dianwen Ng, Hai Ye, Xingxuan Li, Yao Xiao, Zhanfeng Mo, Qi Zhang, et al. 100 days after deepseek-rl: A survey on replication studies and more directions for reasoning language models. *arXiv preprint arXiv:2505.00551*, 2025.
- Yao Zhao, Rishabh Joshi, Tianqi Liu, Misha Khalman, Mohammad Saleh, and Peter J Liu. Slic-hf: Sequence likelihood calibration with human feedback. *arXiv preprint arXiv:2305.10425*, 2023.
- Brian D Ziebart. *Modeling purposeful adaptive behavior with the principle of maximum causal entropy*. Carnegie Mellon University, 2010.
- Brian D Ziebart, Andrew L Maas, J Andrew Bagnell, and Anind K Dey. Maximum entropy inverse reinforcement learning. *AAAI Conference on Artificial Intelligence*, 8:1433–1438, 2008.

Appendix

A Related Works	15
B Technical Tools	16
C Change of Trajectory Measure with Concentrability	18
D Proof of Theorem 2	20
D.1 Proof of Lemma 7	21
E Comparison with Related Algorithms	22
F Hard Instances for Iterative Algorithms	24
G Implementation Details	25
G.1 Training Details	25
G.2 Training Efficiency of TBRM	25
G.3 Prompt Templates	25
H Additional Experimental Results	26
H.1 Training-Time Performance	26
H.2 Qualitative Analysis	26
H.3 Ablation Study: Classical BRM on LLMs	30
H.4 Tasks Beyond Mathematical Problems	30
H.5 TBRM with More Rollouts	31

A Related Works

Value-based RL. Value-based methods are arguably the oldest and most widely studied concepts of reinforcement learning (RL) algorithms (Bellman, 1957; Samuel, 1959). They seek to learn an approximation of the optimal state-action-value function Q^* and act greedily with respect to it, in contrast to policy-gradient methods that directly optimize a parameterized policy. Early works such as Q-learning (Watkins and Dayan, 1992), SARSA (Sutton et al., 1998), Approximate Dynamic Programming (Bertsekas and Tsitsiklis, 1996) established the foundations, while the successive studies introduced function approximation (Sutton, 1988; Bertsekas and Tsitsiklis, 1996), eligibility traces (Singh and Sutton, 1996), and residual updates (Baird et al., 1995). The combination of value-based ideas with deep neural networks culminated in the Deep Q-Network (Mnih et al., 2015), which sparked a wave of extensions including Double DQN (Hasselt, 2010; Van Hasselt et al., 2016), distributional learning (Bellemare et al., 2017), and the integrative DQN-based agent (Hessel et al., 2018).

KL-regularized RL. KL-regularized (or entropy-regularized) reinforcement learning (RL) originated from the maximum-entropy formulation of Ziebart et al. (2008); Ziebart (2010); Neu et al. (2017), where a Kullback–Leibler (KL) penalty encourages policies to stay close to a reference distribution while optimizing reward. Different styles of algorithms have emerged from this line of work, including Soft Q-Learning (SQL) style algorithms such as SQL itself (Haarnoja et al., 2017; Schulman et al., 2017a; Guo et al., 2022), Soft Actor-Critic (SAC) style algorithms like PCL (Nachum et al., 2017), SAC (Haarnoja et al., 2018), DQO (Ji et al., 2024), and OREO (Wang et al., 2024), Point-Wise Direct Alignment Algorithms (DAA-pt) such as KTO (Ethayarajh et al., 2024), DRO (Richemond et al., 2024), and an online policy mirror descent variant (Team et al., 2025), and Pair-Wise Direct Alignment Algorithms (DAA-pair) like DPO (Rafailov et al., 2023, 2024) and IPO (Azar et al., 2024). In Appendix E, we provide a comprehensive discussion of the differences between TBRM and these related algorithms.

RL training for LLM reasoning. RL has played a pivotal role in the post-training of LLMs. The most prominent early example is reinforcement learning from human feedback (RLHF; Ouyang et al., 2022; Bai et al., 2022), which uses PPO to align LLMs with human preferences. A series of subsequent works introduced contrastive learning objectives based on pairwise datasets (Rafailov

et al., 2024; Zhao et al., 2023; Azar et al., 2024; Tang et al., 2024), or verification-driven objectives using a binary verifier (Ethayarajh et al., 2024; Chen et al., 2025). The release of OpenAI’s O1 (Jaech et al., 2024) and DeepSeek’s R1 (Guo et al., 2025) marked a new era of RL algorithms for LLMs—particularly for reasoning tasks—by framing the response generation process as a Markov Decision Process (MDP) and using rule-based verifiers to provide reward signals. Numerous studies have demonstrated and analyzed the effectiveness of RL algorithms in enhancing LLM reasoning capabilities, with PPO and GRPO emerging as the most widely adopted approaches. Prior studies, such as SimpleRL-Zoo (Zeng et al., 2025), Open Reasoner Zero (Hu et al., 2025), Light-r1 (Wen et al., 2025), Logic-rl (Xie et al., 2025), and Skywork-OR1 (He et al., 2025), fall in this category. Variants of these algorithms have been proposed to further improve performance (Zhang et al., 2025). For instance, DAPO (Yu et al., 2025) enhances GRPO with techniques like clip-higher, dynamic sampling, and token-level policy gradient loss, achieving strong results on AIME24. Dr. GRPO (Liu et al., 2025) addresses optimization bias in GRPO to improve token efficiency, while CPPO (Lin et al., 2025) reduces its computational cost by skipping rollouts with low advantages. VC-PPO (Yuan et al., 2025b) resolves PPO’s challenges with value initialization bias and delayed reward signals through value pretraining and decoupled-GAE. Building on this, VAPO (Yuan et al., 2025a) improve DAPO further by incorporating selected techniques from VC-PPO. Additionally, several works explore REINFORCE (Williams and Peng, 1991; Williams, 1992) style RL algorithms, including ReMax (Li et al., 2023), REINFORCE++ (Hu, 2025), RAFT++ (Xiong et al., 2025), and RLOO (Kool et al., 2019; Ahmadian et al., 2024). However, all these approaches rely on policy-based methods or their variants. In contrast, our method adopts a value-based, off-policy RL approach that is principled, efficient, and theoretically grounded.

B Technical Tools

We now present Lemma 3 as a soft performance difference lemma with arbitrary reference function. Lemma 3 holds generally for KL-regularized RL, which uses a slightly different form of Bellman operator defined as below,

$$(\mathcal{T}_\beta^\pi f)(s_h, a_h) := r(s_h, a_h) + \mathbb{E}_{s_{h+1} \sim \mathcal{P}(\cdot | s_h, a_h), a_{h+1} \sim \pi(\cdot | s_{h+1})} \left[f(s_{h+1}, a_{h+1}) - \beta \log \frac{\pi(a_{h+1} | s_{h+1})}{\pi_{\text{ref}}(a_{h+1} | s_{h+1})} \right].$$

We present this slightly different version as this lemma is more generally applicable than the results in this paper.

Lemma 3 (Soft Performance Difference Lemma via Reference Function). *For any function f as well as any policies $\pi^\dagger$ and π , we have*

$$\begin{aligned} J_\beta(\pi^\dagger) - J_\beta(\pi) &= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H ((\mathcal{T}_\beta^\pi f)(s_h, a_h) - f(s_h, a_h)) \right] + \mathbb{E}_\pi \left[\sum_{h=1}^H (f(s_h, a_h) - (\mathcal{T}_\beta^\pi f)(s_h, a_h)) \right] \\ &\quad + \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(f(s_h, a_h) - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - \mathbb{E}_{a_h \sim \pi(\cdot | s_h)} \left[f(s_h, a_h) - \beta \log \frac{\pi(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right] \right) \right]. \end{aligned}$$

Proof of Lemma 3. We first prove the soft performance difference via Q_β^π as the reference,

$$\begin{aligned} &J_\beta(\pi^\dagger) - J_\beta(\pi) \\ &= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(r(s_h, a_h) - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right) \right] - J_\beta(\pi) \\ &= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(r(s_h, a_h) - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - V_\beta^\pi(s_h) + V_\beta^\pi(s_h) \right) \right] - J_\beta(\pi) \\ &= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(r(s_h, a_h) - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - V_\beta^\pi(s_h) + V_\beta^\pi(s_h) \right) \right] - \mathbb{E}_{s_1 \sim \rho} [V_\beta^\pi(s_1)] \\ &= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(r(s_h, a_h) - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - V_\beta^\pi(s_h) + V_\beta^\pi(s_h) \right) - V_\beta^\pi(s_1) \right] \end{aligned}$$

$$\begin{aligned}
&= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(r(s_h, a_h) + \mathbb{E}_{s_{h+1} \sim \mathcal{P}(\cdot | s_h, a_h)} [V_\beta^\pi(s_{h+1})] - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - V_\beta^\pi(s_h) \right) \right] \\
&= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(Q_\beta^\pi(s_h, a_h) - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - \mathbb{E}_{a_h \sim \pi(\cdot | s_h)} \left[Q_\beta^\pi(s_h, \pi) + \beta \log \frac{\pi(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right] \right) \right].
\end{aligned}$$

Next, consider an arbitrary reference function f , we define the augment reward r_f as

$$\begin{aligned}
r_f(s_h, a_h) &:= f(s_h, a_h) - \mathbb{E}_{s_{h+1} \sim \mathcal{P}(\cdot | s_h, a_h), a_{h+1} \sim \pi(\cdot | s_{h+1})} \left[f(s_{h+1}, a_{h+1}) - \beta \log \frac{\pi(a_{h+1} | s_{h+1})}{\pi_{\text{ref}}(a_{h+1} | s_{h+1})} \right] \\
&= f(s_h, a_h) - (\mathcal{T}_\beta^\pi f)(s_h, a_h) + r(s_h, a_h),
\end{aligned}$$

for any $(s_h, a_h) \in \mathcal{S}_h \times \mathcal{A}_h$. This means f is the fixed point of $\mathcal{T}_\beta^\pi$ with replacing reward function r by r_f . We use Q_{β, r_f}^π to denote the soft Q-function with replacing reward function r by r_f , and we immediately have $f \equiv Q_{\beta, r_f}^\pi$. We also use $J_{\beta, r_f}(\pi)$ to denote the soft return of policy π with replacing reward function r by r_f . Then

$$\begin{aligned}
&J_\beta(\pi^\dagger) - J_\beta(\pi) \\
&= J_\beta(\pi^\dagger) - J_{\beta, r_f}(\pi^\dagger) + J_{\beta, r_f}(\pi^\dagger) - J_{\beta, r_f}(\pi) + J_{\beta, r_f}(\pi) - J_\beta(\pi) \\
&= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H (r(s_h, a_h) - r_f(s_h, a_h)) \right] + \mathbb{E}_\pi \left[\sum_{h=1}^H (r_f(s_h, a_h) - r(s_h, a_h)) \right] \\
&\quad + \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(f(s_h, a_h) - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - \mathbb{E}_{a_h \sim \pi(\cdot | s_h)} \left[f(s_h, a_h) - \beta \log \frac{\pi(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right] \right) \right] \\
&\quad \text{(by the soft performance difference via reference as } f \equiv Q_{\beta, r_f}^\pi \text{)} \\
&= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H ((\mathcal{T}_\beta^\pi f)(s_h, a_h) - f(s_h, a_h)) \right] + \sum_{h=1}^H \mathbb{E}_{d_h^\pi} [f(s_h, a_h) - (\mathcal{T}_\beta^\pi f)(s_h, a_h)] \\
&\quad + \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H \left(f(s_h, a_h) - \beta \log \frac{\pi^\dagger(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - \mathbb{E}_{a_h \sim \pi(\cdot | s_h)} \left[f(s_h, a_h) - \beta \log \frac{\pi(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right] \right) \right].
\end{aligned}$$

This completes the proof. $\square$

When specializing to the Bellman operator in this paper as in [Section 2.3](#), [Lemma 3](#) becomes

$$\begin{aligned}
\frac{J_\beta(\pi^\dagger) - J_\beta(\pi)}{\beta} &= \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H ((\mathcal{T}_\beta^\pi Q)(s_h, a_h) - Q(s_h, a_h)) \right] + \mathbb{E}_\pi \left[\sum_{h=1}^H (Q(s_h, a_h) - (\mathcal{T}_\beta^\pi Q)(s_h, a_h)) \right] \\
&\quad + \mathbb{E}_{\pi^\dagger} \left[\sum_{h=1}^H (Q(s_h, a_h) - \log \pi^\dagger(a_h | s_h) - \mathbb{E}_{a_h \sim \pi(\cdot | s_h)} [Q(s_h, a_h) - \log \pi(a_h | s_h)]) \right].
\end{aligned}$$

As corollary, we have the following upper bound on the sub-optimality of any policy induced by a value function.

Corollary 4. Suppose that the policy $\hat{\pi} = \pi_Q$ is induced by a value function Q , i.e.,

$$\hat{\pi}(a | s) = \frac{\exp(Q(s, a))}{\sum_{a' \in \mathcal{A}} \exp(Q(s, a'))}.$$

Then it holds that

$$J_\beta(\pi^\star) - J_\beta(\hat{\pi}) \leq 2\beta \max_{\pi \in \{\pi^\star, \hat{\pi}\}} \left| \mathbb{E}_\pi \left[\sum_{h=1}^H (Q(s_h, a_h) - (\mathcal{T}_\beta Q)(s_h, a_h)) \right] \right|$$

Proof of Corollary 4. By [Lemma 3](#), we have

$$J_\beta(\pi^\star) - J_\beta(\hat{\pi})$$

$$\begin{aligned}
&= \beta \mathbb{E}_{\pi^*} \left[\sum_{h=1}^H \left((\mathcal{T}_{\hat{\beta}} Q)(s_h, a_h) - Q(s_h, a_h) \right) \right] + \beta \mathbb{E}_{\hat{\pi}} \left[\sum_{h=1}^H \left(Q(s_h, a_h) - (\mathcal{T}_{\hat{\beta}} Q)(s_h, a_h) \right) \right] \\
&\quad + \beta \mathbb{E}_{\pi^*} \left[\sum_{h=1}^H \left(Q(s_h, a_h) - \log \pi^*(a_h | s_h) - \mathbb{E}_{a'_h \sim \hat{\pi}(\cdot | s_h)} [Q(s_h, a'_h) - \log \hat{\pi}(a'_h | s_h)] \right) \right] \\
&\leq \beta \mathbb{E}_{\pi^*} \left[\sum_{h=1}^H ((\mathcal{T}_{\beta} Q)(s_h, a_h) - Q(s_h, a_h)) \right] + \beta \mathbb{E}_{\hat{\pi}} \left[\sum_{h=1}^H (Q(s_h, a_h) - (\mathcal{T}_{\beta} Q)(s_h, a_h)) \right] \\
&\leq 2\beta \max_{\pi} \left| \mathbb{E}_{\pi} \left[\sum_{h=1}^H (Q(s_h, a_h) - (\mathcal{T}_{\beta} Q)(s_h, a_h)) \right] \right|,
\end{aligned}$$

where the first inequality follows from the definition that

$$\hat{\pi}(\cdot | s_h) = \operatorname{argmax}_{p \in \Delta(\mathcal{A})} \mathbb{E}_{a'_h \sim p} [Q(s_h, a'_h) - \log p(a'_h)], \quad \forall s_h \in \mathcal{S}_h,$$

and hence

$$\mathbb{E}_{a_h \sim \pi^*(\cdot | s_h)} [Q(s_h, a_h) - \log \pi^*(a_h | s_h)] \leq \mathbb{E}_{a'_h \sim \hat{\pi}(\cdot | s_h)} [Q(s_h, a'_h) - \log \hat{\pi}(a'_h | s_h)].$$

□

C Change of Trajectory Measure with Concentrability

We first state the following lemma, which follows immediately from the definition of χ^2 -divergence, see e.g. (Polyanskiy and Wu, 2025, Eq. (7.91))

Lemma 5. *For distribution $P, Q \in \Delta(\mathcal{X})$ and function $F : \mathcal{X} \rightarrow \mathbb{R}$ such that $\mathbb{E}_Q[F(X)] = 0$, it holds that*

$$(\mathbb{E}_P F(X))^2 \leq \chi^2(P \parallel Q) \cdot \mathbb{E}_Q F(X)^2.$$

Proof of Lemma 1. We only need to prove that for any function $f : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$,

$$\left(\mathbb{E}_{\pi} \left[\sum_{h=1}^H f(s_h, a_h) \right] \right)^2 \leq \left(1 + \sum_{h=1}^H \chi^2(d_h^{\pi} \parallel d_h^{\pi_{\text{ref}}}) \right) \cdot \mathbb{E}_{\pi_{\text{ref}}} \left[\left(\sum_{h=1}^H f(s_h, a_h) \right)^2 \right]. \quad (7)$$

First of all, we construct function $\bar{f} : \mathcal{S} \rightarrow \mathbb{R}$ as

$$\bar{f}(s_h) := \mathbb{E}_{\mu} \left[\sum_{h'=h}^H f(s_{h'}, a_{h'}) \mid s_h \right], \quad \forall h \in [H], s_h \in \mathcal{S}_h. \quad (8)$$

Then, it is direct to verify that for $h \in [H], s_h \in \mathcal{S}_h$,

$$\bar{f}(s_h) := \mathbb{E}_{\mu} [f(s_h, a_h) + \bar{f}(s_{h+1}) \mid s_h],$$

where we adopt the notation that s_{H+1} is a deterministic terminal state and $\bar{f}(s_{H+1}) = 0$. Then, we expand

$$\begin{aligned}
&\mathbb{E}_{\mu} \left[\left(\sum_{h=1}^H f(s_h, a_h) \right)^2 \right] \\
&= \mathbb{E}_{\mu} \left[\left(\bar{f}(s_1) + \sum_{h=1}^H [f(s_h, a_h) + \bar{f}(s_{h+1}) - \bar{f}(s_h)] \right)^2 \right] \\
&= \mathbb{E}_{\mu} \left[\bar{f}(s_1)^2 + \sum_{h=1}^H (f(s_h, a_h) + \bar{f}(s_{h+1}) - \bar{f}(s_h))^2 \right]
\end{aligned}$$

$$\begin{aligned}
& + \mathbb{E}_\mu \left[\sum_{1 \leq h \leq H} \bar{f}(s_1) (f(s_h, a_h) + \bar{f}(s_{h+1}) - \bar{f}(s_h)) \right] \\
& + \mathbb{E}_\mu \left[\sum_{1 \leq h' < h \leq H} (f(s_{h'}, a_{h'}) + \bar{f}(s_{h'+1}) - \bar{f}(s_{h'})) (f(s_h, a_h) + \bar{f}(s_{h+1}) - \bar{f}(s_h)) \right].
\end{aligned}$$

Therefore, using the Markov property, it holds that for any $h \in [H]$,

$$\mathbb{E}_\mu [f(s_h, a_h) + \bar{f}(s_{h+1}) - \bar{f}(s_h) | s_1, a_1, \dots, s_h] = \mathbb{E}_\mu [f(s_h, a_h) + \bar{f}(s_{h+1}) | s_h] - \bar{f}(s_h) = 0,$$

and hence we can deduce that

$$\mathbb{E}_\mu \left[\left(\sum_{h=1}^H f(s_h, a_h) \right)^2 \right] = \mathbb{E}_\mu \left[\bar{f}(s_1)^2 + \sum_{h=1}^H (f(s_h, a_h) + \bar{f}(s_{h+1}) - \bar{f}(s_h))^2 \right]. \quad (9)$$

Next, for every $h \in [H]$, we apply [Lemma 5](#) on the function $(s_h, a_h, s_{h+1}) \mapsto f(s_h, a_h) - \bar{f}(s_h) + \bar{f}(s_{h+1})$ to derive

$$\begin{aligned}
& (\mathbb{E}_\pi [f(s_h, a_h) - \bar{f}(s_h) + \bar{f}(s_{h+1})])^2 \\
& \leq \chi^2 (d_h^\pi \parallel d_h^\mu) \cdot \mathbb{E}_\mu \left[(f(s_h, a_h) - \bar{f}(s_h) + \bar{f}(s_{h+1}))^2 \right],
\end{aligned} \quad (10)$$

where we again use the fact that $\mathbb{E}_\mu [f(s_h, a_h) + \bar{f}(s_{h+1}) - \bar{f}(s_h)] = 0$.

Furthermore, we note that $\mathbb{E}_\pi \bar{f}(s_1) = \mathbb{E}_{s_1 \sim \rho} \bar{f}(s_1) = \mathbb{E}_\mu \bar{f}(s_1)$, and hence

$$(\mathbb{E}_\pi \bar{f}(s_1))^2 = (\mathbb{E}_\mu \bar{f}(s_1))^2 \leq \mathbb{E}_\mu \bar{f}(s_1)^2. \quad (11)$$

Therefore, combining the inequalities above, we have

$$\begin{aligned}
& \left(\mathbb{E}_\pi \left[\sum_{h=1}^H f(s_h, a_h) \right] \right)^2 \\
& = \left(\bar{f}(s_1) + \sum_{h=1}^H \mathbb{E}_\pi [f(s_h, a_h) - \bar{f}(s_h) + \bar{f}(s_{h+1})] \right)^2 \\
& \stackrel{(i)}{\leq} \left(\sqrt{\mathbb{E}_\mu \bar{f}(s_1)^2} + \sum_{h=1}^H \sqrt{\chi^2 (d_h^\pi \parallel d_h^\mu) \cdot \mathbb{E}_\mu [(f(s_h, a_h) - \bar{f}(s_h) + \bar{f}(s_{h+1}))^2]} \right)^2 \\
& \stackrel{(ii)}{\leq} \left(1 + \sum_{h=1}^H \chi^2 (d_h^\pi \parallel d_h^\mu) \right) \left(\mathbb{E}_\mu \bar{f}(s_1)^2 + \sum_{h=1}^H \mathbb{E}_\mu (f(s_h, a_h) - \bar{f}(s_h) + \bar{f}(s_{h+1}))^2 \right) \\
& \stackrel{(iii)}{=} \left(1 + \sum_{h=1}^H \chi^2 (d_h^\pi \parallel d_h^\mu) \right) \cdot \mathbb{E}_\mu \left[\left(\sum_{h=1}^H f(s_h, a_h) \right)^2 \right],
\end{aligned}$$

where (i) uses [Eq. \(10\)](#) and [Eq. \(11\)](#), (ii) uses Cauchy-Schwarz inequality, and (iii) uses [Eq. \(9\)](#). Hence [Eq. \(7\)](#) is verified. $\square$

Recall the definition of state-action concentrability (e.g. [Eq. \(5\)](#) in [Jia et al. \(2025\)](#)):

$$C_{\text{sa}}(\pi; \mu) = \max_{h \in [H]} \sup_{s_h \in \mathcal{S}_h, a_h \in \mathcal{A}} \frac{d_h^\pi(s_h, a_h)}{d_h^\mu(s_h, a_h)}.$$

We have following direct corollary of [Lemma 1](#), after noticing that the χ^2 divergence can be upper bounded by the state concentrability. However, we remark that the state concentrability might yield a much more pessimistic bound compared to the upper bound of [Lemma 1](#).

Corollary 6. Given MDP $M = (H, \mathcal{S}, \mathcal{A}, \mathcal{P}, r, \rho)$. We use $d_h^\pi(\cdot), d_h^\mu(\cdot) \in \Delta(\mathcal{S}_h \times \mathcal{A})$ to denote the occupancy measure of the MDP under policy π and μ . Then we have

$$\sup_f \frac{\left(\mathbb{E}_\pi \left[\sum_{h=1}^H f(s_h, a_h) \right] \right)^2}{\mathbb{E}_\mu \left[\left(\sum_{h=1}^H f(s_h, a_h) \right)^2 \right]} \leq H \cdot C_{\text{sa}}(\pi; \mu),$$

where the supremum is over all functions $f : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$.

Proof of Corollary 6. In view of Lemma 1, we only need to verify that

$$1 + \chi^2(d_h^\pi \parallel d_h^\mu) \leq C_{\text{sa}}. \quad (12)$$

We have

$$\begin{aligned} 1 + \chi^2(d_h^\pi \parallel d_h^\mu) &= \sum_{s_h \in \mathcal{S}_h, a_h \in \mathcal{A}} \frac{d_h^\pi(s_h, a_h)^2}{d_h^\mu(s_h, a_h)} \\ &\leq \sum_{s_h \in \mathcal{S}_h, a_h \in \mathcal{A}} d_h^\pi(s_h, a_h) \cdot C_{\text{sa}}(\pi; \mu) = C_{\text{sa}}(\pi; \mu), \end{aligned}$$

and Eq. (12) is verified. $\square$

D Proof of Theorem 2

Recall that the TBRM loss defined in Eq. (4) is given by

$$\mathcal{L}_{\mathcal{D}}^{\text{TBRM}}(\theta) := \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(\text{logit}_\theta(s_1, a_1) - \frac{r(\tau)}{\beta} - \log \pi_{\text{ref}}(\tau) + \sum_{h=2}^H \log \pi_\theta(a_h \mid s_h) \right)^2. \quad (13)$$

In the following, to simplify presentation, we define

$$\begin{aligned} f_{\theta,1}(s_1, a_1) &= \text{logit}_\theta(s_1, a_1) - \frac{r(s_1, a_1)}{\beta} - \log \pi_{\text{ref}}(a_1 \mid s_1), \\ f_{\theta,h}(s_h, a_h) &= \log \pi_\theta(a_h \mid s_h) - \frac{r(s_h, a_h)}{\beta} - \log \pi_{\text{ref}}(a_h \mid s_h), \quad \forall h > 1. \end{aligned}$$

and

$$\begin{aligned} f_\theta(\tau) &:= \text{logit}_\theta(s_1, a_1) - \frac{r(\tau)}{\beta} - \log \pi_{\text{ref}}(\tau) + \sum_{h=2}^H \log \pi_\theta(a_h \mid s_h) \\ &= \sum_{h=1}^H f_{\theta,h}(s_h, a_h), \quad \forall \tau = (s_1, a_1, \dots, s_H, a_H). \end{aligned}$$

Uniform convergence. Before applying Lemma 1, we need to first relate the empirical loss $\mathcal{L}_{\mathcal{D}}^{\text{TBRM}}$ to the population loss. We introduce the following assumption on the parametrization.

Assumption 2 (Bounded and smooth parametrization). *There exists constant $C_\Theta \geq 1$ and parameter L_Θ such that for any $\theta \in \Theta$, it holds that $\forall s \in \mathcal{S}, a \in \mathcal{A}$,*

$$|\text{logit}_\theta(s, a)| \leq \frac{C_\Theta}{\beta}, \quad \|\nabla \text{logit}_\theta(s, a)\|_* \leq L_\Theta,$$

where $\|\cdot\|_*$ is the dual norm of $\|\cdot\|$. We also assume $|\log \pi_{\text{ref}}(a \mid s)| \leq \frac{C_\Theta}{\beta} \forall s \in \mathcal{S}, a \in \mathcal{A}$.

Lemma 7. Fix $\delta \in (0, 1)$. Suppose that Assumption 2 holds. Then with probability at least $1 - \delta$ (over the randomness of the dataset $\mathcal{D}$),

$$\mathbb{E}_{\tau \sim \mu} f_\theta(\tau)^2 \leq 2\mathcal{L}_{\mathcal{D}}^{\text{TBRM}}(\theta) + \frac{\varepsilon_{\text{stat}}(N)}{\beta^2},$$

where the statistical error $\varepsilon_{\text{stat}}(N)$ is defined as

$$\varepsilon_{\text{stat}}(N) := c \cdot C_\Theta^2 H^2 \left(\frac{\log(\mathcal{N}(\Theta, \alpha)/\delta)}{N} + H L_\Theta \beta \alpha \right)$$

where $c > 0$ is a large absolute constant and $\alpha \geq 0$ is a fixed parameter.

The above upper bound can be further simplified by the standard assumption on the covering number of Θ .

Assumption 3 (Parametric function class). *The parameter space $\Theta \subseteq \{\theta \in \mathbb{R}^d : \|\theta\| \leq R\}$. In this case, we write $\dim(\Theta) = d$.*

Under [Assumption 3](#), it is clear that $\log \mathcal{N}(\Theta, \alpha) \leq O(d \log(1/\alpha))$ for all $\alpha > 0$ (see e.g., [Wainwright \(2019\)](#)). Therefore, [Lemma 7](#) implies that $\varepsilon_{\text{stat}}(N) \asymp \frac{H^2}{N}$ (up to poly-logarithmic factors).

Bounding the sub-optimality. Under [Assumption 1](#), it is clear that $f_{\theta^*} \equiv 0$ and hence $\mathcal{L}_{\mathcal{D}}^{\text{TBRM}}(\theta^*) = 0$ (as we have argued in [Section 3.1](#)). Therefore, using the condition that $\mathcal{L}_{\mathcal{D}}^{\text{TBRM}}(\hat{\theta}) - \inf_{\theta \in \Theta} \mathcal{L}_{\mathcal{D}}^{\text{TBRM}}(\theta) \leq \varepsilon_{\text{opt}}$, we have

$$\mathbb{E}_{\tau \sim \mu} f_{\hat{\theta}}(\tau)^2 \leq 2\varepsilon_{\text{opt}} + \beta^{-2} \varepsilon_{\text{stat}}(N).$$

In the following, we denote $\hat{\pi} := \pi_{\hat{\theta}}$, $\hat{Q}(s, a) := \text{logit}_{\hat{\theta}}(s, a)$, and $\hat{V}$ be the corresponding value function. Then, by [Corollary 4](#), we have

$$J_{\beta}(\pi^*) - J_{\beta}(\hat{\pi}) \leq 2\beta \max_{\pi \in \{\pi^*, \hat{\pi}\}} \left| \mathbb{E}_{\pi} \left[\sum_{h=1}^H \left(\hat{Q}(s_h, a_h) - (\mathcal{T}_{\beta} \hat{Q})(s_h, a_h) \right) \right] \right|.$$

Note that the MDP is deterministic, and hence $(\mathcal{T}_{\beta} \hat{Q})(s_h, a_h) = R_{\beta}(s_h, a_h) + \hat{V}(s_{h+1})$ holds deterministically. Therefore, for any fixed policy π , we have

$$\begin{aligned} & \mathbb{E}_{\pi} \left[\sum_{h=1}^H \left(\hat{Q}(s_h, a_h) - (\mathcal{T}_{\beta} \hat{Q})(s_h, a_h) \right) \right] \\ &= \mathbb{E}_{\pi} \left[\sum_{h=1}^H \left(\hat{Q}(s_h, a_h) - R_{\beta}(s_h, a_h) - \hat{V}(s_{h+1}) \right) \right] \\ &= \mathbb{E}_{\pi} \left[\hat{Q}(s_1, a_1) - \sum_{h=1}^H R_{\beta}(s_h, a_h) + \sum_{h=1}^H [\hat{Q}(s_h, a_h) - \hat{V}(s_h)] \right] \\ &= \mathbb{E}_{\pi} \left[\hat{Q}(s_1, a_1) - R_{\beta}(\tau) + \sum_{h=1}^H \log \pi_{\hat{\theta}}(a_h | s_h) \right] = \mathbb{E}_{\pi} [f_{\hat{\theta}}(\tau)]. \end{aligned}$$

Further, by [Lemma 1](#), it holds that for any $\pi = \pi_{\theta}$,

$$(\mathbb{E}_{\pi} [f_{\hat{\theta}}(\tau)])^2 \leq HC_{\text{conc}}(\mu) \cdot \mathbb{E}_{\tau \sim \mu} f_{\hat{\theta}}(\tau)^2 \leq HC_{\text{conc}}(\mu) (2\varepsilon_{\text{opt}} + \beta^{-2} \varepsilon_{\text{stat}}(N)).$$

Therefore, we can conclude that

$$J_{\beta}(\pi^*) - J_{\beta}(\hat{\pi}) \leq 2\beta \max_{\pi \in \{\pi^*, \hat{\pi}\}} \sqrt{\mathbb{E}_{\pi} [f_{\hat{\theta}}(\tau)]} \leq 2\sqrt{HC_{\text{conc}}(\mu) (2\beta^2 \varepsilon_{\text{opt}} + \varepsilon_{\text{stat}}(N))}.$$

This is the desired upper bound. $\square$

D.1 Proof of [Lemma 7](#)

By [Assumption 2](#), it holds that $|f_{\theta}(\tau)| \leq B := \frac{2C_{\Theta}+1}{\beta}$ for any $\theta \in \Theta$ and any trajectory τ . Using Freedman's inequality with the standard union bound, we have the following: with probability at least $1 - \delta$ (over the randomness of the dataset $\mathcal{D}$), for all $\theta \in \Theta$,

$$\mathbb{E}_{\tau \sim \mu} f_{\theta}(\tau)^2 \leq 2\mathcal{L}_{\mathcal{D}}^{\text{TBRM}}(\theta) + c_0 B^2 \left(\frac{\log \mathcal{N}(\mathcal{F}, \alpha) + \log(1/\delta)}{N} + B^{-1} \alpha \right),$$

where $\mathcal{F} = \{f_{\theta} : \theta \in \Theta\}$ is the function class induced by Θ , and $c_0 > 0$ is an absolute constant.

Next, for any fixed trajectory τ , it holds that

$$\nabla_{\theta} f_{\theta}(\tau) = \nabla_{\theta} \text{logit}_{\theta}(s_1, a_1) + \sum_{h=2}^H \nabla_{\theta} \log \pi_{\theta}(a_h | s_h)$$

$$= \nabla_{\theta} \text{logit}_{\theta}(s_1, a_1) + \sum_{h=2}^H \left[\nabla_{\theta} \text{logit}_{\theta}(s_h, a_h) - \mathbb{E}_{a'_h \sim \pi_{\theta}(\cdot | s_h)} \nabla_{\theta} \text{logit}_{\theta}(s_h, a'_h) \right].$$

Hence, we can upper bound $\|\nabla_{\theta} f_{\theta}(\tau)\|_* \leq 2HL_{\Theta}$. This immediately implies that

$$\|f_{\theta} - f_{\theta'}\|_{\infty} = \sup_{\tau} |f_{\theta}(\tau) - f_{\theta'}(\tau)| \leq 2HL_{\Theta} \|\theta - \theta'\|, \quad \forall \theta, \theta' \in \Theta.$$

Therefore, we have

$$\mathcal{N}(\mathcal{F}, \alpha) \leq \mathcal{N}\left(\Theta, \frac{\alpha}{2HL_{\Theta}}\right).$$

Combining the inequalities above and rescaling $\alpha \leftarrow 2HL_{\Theta}\alpha$ completes the proof. $\square$

E Comparison with Related Algorithms

In this section, we compare TBRM with other related algorithms in detail. We group the related algorithms into the following categories:

- **Soft Q-Learning (SQL) Style:** SQL (Haarnoja et al., 2017; Schulman et al., 2017a; Guo et al., 2022)
- **Soft Actor-Critic (SAC) Style:** PCL (Nachum et al., 2017), SAC (Haarnoja et al., 2018), DQO (Ji et al., 2024), OREO (Wang et al., 2024)
- **Point-Wise Direct Alignment Algorithms (DAA-pt):** KTO (Ethayarajh et al., 2024), DRO (Richemond et al., 2024), online policy mirror decent variant (Team et al., 2025)
- **Pair-Wise Direct Alignment Algorithms (DAA-pair):** DPO (Rafailov et al., 2023, 2024), IPO (Azar et al., 2024)

It is important to note that the present paper primarily addresses LLM reasoning in environments where the state space is tokenized and the base model operates autoregressively. Several algorithms mentioned above were initially developed for continuous control domains such as robotics; however, our analysis considers only their adaptation to the discrete, tokenized setting relevant to language models, as in this paper. Given space constraints, we restrict our discussion to the fundamental principles underlying each algorithmic category rather than providing exhaustive implementation details.

We first present Table 2 to summarize the key differences between TBRM and other algorithms. Note that this comparison is only for algorithm design; the consequences of these differences for theoretical guarantees are likely to be more significant. However, given that TBRM is the only algorithm here with established finite-sample guarantees under the more general MDP setting, we will leave the theoretical comparison to future work.

Algorithms	Optimization	Single Rollout	Single Model Training	Traj. Reward Allowed
SQL	✗	✓	✓	✗
SAC	✓	✓	✗	✗
DAA-pt	✓	✗	✓	✓
DAA-pair	✓	✗	✓	✓
TBRM	✓	✓	✓	✓

Table 2: Comparison between TBRM and related algorithms in terms of algorithm design.

For the ease of comparison, we rewrite the loss function of TBRM as follows, by the definition of the autoregressive function approximation:

$$\begin{aligned} \mathcal{L}_{\mathcal{D}}^{\text{TBRM}}(\theta) &= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(\text{logit}_{\theta}(s_1, a_1) - \frac{r(\tau)}{\beta} - \log \pi_{\text{ref}}(\tau) + \sum_{h=2}^H \log \pi_{\theta}(a_h | s_h) \right)^2 \\ &= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(V_{\theta}(s_1) - \frac{r(\tau)}{\beta} + \sum_{h=1}^H \log \frac{\pi_{\theta}(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right)^2. \end{aligned}$$

Comparing TBRM with Soft Q-Learning Style Algorithms. The soft Q-learning based algorithms are typically iterative algorithms, formulated as two different versions: single-step case and multi-step case. We consider the loss used in Guo et al. (2022) which is motivated by *path consistency learning* (PCL; Nachum et al., 2017). The single-step case is then formulated as

$$\theta_{t+1} \leftarrow \underset{\theta}{\operatorname{argmin}} \mathcal{L}_{\mathcal{D}}^{\text{SQL-s}}(\theta; \theta_t) := \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \sum_{h=1}^H \left(V_{\theta_t}(s_h) - \frac{r(s_h, a_h)}{\beta} + \log \frac{\pi_{\theta}(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - V_{\theta_t}(s_{h+1}) \right)^2,$$

while the multi-step case is formulated as

$$\theta_{t+1} \leftarrow \underset{\theta}{\operatorname{argmin}} \mathcal{L}_{\mathcal{D}}^{\text{SQL-m}}(\theta; \theta_t) := \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \sum_{h=1}^H \left(V_{\theta_t}(s_h) - \sum_{h'=h}^H \frac{r(s_{h'}, a_{h'})}{\beta} + \sum_{h'=h}^H \log \frac{\pi_{\theta}(a_{h'} | s_{h'})}{\pi_{\text{ref}}(a_{h'} | s_{h'})} \right)^2.$$

If we want to exactly match the original loss of soft Q-learning (Haarnoja et al., 2017; Schulman et al., 2017a), then these should be rewritten as

$$\begin{aligned} \theta_{t+1} \leftarrow \underset{\theta}{\operatorname{argmin}} \tilde{\mathcal{L}}_{\mathcal{D}}^{\text{SQL-s}}(\theta; \theta_t) &:= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \sum_{h=1}^H \left(\text{logit}_{\theta}(s_h) - \frac{r(s_h, a_h)}{\beta} - \log \pi_{\text{ref}}(a_h | s_h) - V_{\theta_t}(s_{h+1}) \right)^2, \\ \theta_{t+1} \leftarrow \underset{\theta}{\operatorname{argmin}} \tilde{\mathcal{L}}_{\mathcal{D}}^{\text{SQL-m}}(\theta; \theta_t) &:= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \sum_{h=1}^H \left(\text{logit}_{\theta}(s_h) - \sum_{h'=h}^H \frac{r(s_{h'}, a_{h'})}{\beta} - \log \pi_{\text{ref}}(a_h | s_h) + \sum_{h'=h+1}^H \log \frac{\pi_{\theta_t}(a_{h'} | s_{h'})}{\pi_{\text{ref}}(a_{h'} | s_{h'})} \right)^2. \end{aligned}$$

There can be a more general multi-step version, which blends the multi-step return in $\mathcal{L}_{\mathcal{D}}^{\text{SQL-m}}$ and the value bootstrap in $\mathcal{L}_{\mathcal{D}}^{\text{SQL-s}}$, but we omit it here for brevity as our existing argument would directly extend to this case.

From the derivation above, we can identify two key distinctions between TBRM and soft Q-learning based algorithms: 1) TBRM employs direct optimization rather than an iterative approach, and 2) TBRM’s loss function operates on complete trajectories rather than summing losses over individual timesteps within trajectories, hence eliminating the need for per-step reward.

Comparing TBRM with Soft Actor-Critic Style Algorithms. The soft actor-critic style algorithms for LLMs are similar to SQL, but they 1) introduce a separate V model; 2) operate as optimization rather than iteration. In particular, in the single-step case,

$$\underset{\theta, \phi}{\operatorname{argmin}} \mathcal{L}_{\mathcal{D}}^{\text{SAC-s}}(\theta, \phi) := \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \sum_{h=1}^H \left(V_{\phi}(s_h) - \frac{r(s_h, a_h)}{\beta} + \log \frac{\pi_{\theta}(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - V_{\phi}(s_{h+1}) \right)^2,$$

while the multi-step case is

$$\underset{\theta, \phi}{\operatorname{argmin}} \mathcal{L}_{\mathcal{D}}^{\text{SAC-m}}(\theta, \phi) := \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \sum_{h=1}^H \left(V_{\phi}(s_h) - \sum_{h'=h}^H \frac{r(s_{h'}, a_{h'})}{\beta} + \sum_{h'=h}^H \log \frac{\pi_{\theta}(a_{h'} | s_{h'})}{\pi_{\text{ref}}(a_{h'} | s_{h'})} \right)^2.$$

Comparing TBRM with Point-Wise Direct Alignment Algorithms. Perhaps surprisingly, among all four categories of algorithms, the point-wise direct alignment algorithms appear to be the most similar to TBRM, although they are derived from a different perspective (mostly from bandit formulation). We view the core objective of these algorithms as optimizing the following loss

$$\underset{\theta}{\operatorname{argmin}} \mathcal{L}_{\mathcal{D}}^{\text{DAA-pt}}(\theta) = \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(\hat{V}_{\beta}^*(s_1) - \frac{r(\tau)}{\beta} + \sum_{h=1}^H \log \frac{\pi_{\theta}(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right)^2,$$

where $\hat{V}_{\beta}^*$ is an estimate of V_{β}^* (the optimal value function for KL-regularized MDP). Note that V_{β}^* is exactly the same as the partition function in the bandit formulation. One popular way to estimate $\hat{V}_{\beta}^*$ is by using the softmax of returns from multiple rollouts for each question,⁴ for example (Team et al., 2025),

$$\hat{V}_{\beta}^*(s_1) \leftarrow \beta \log \sum_{\tau \sim \pi_{\theta} | s_1} \exp \left(\frac{r(\tau)}{\beta} \right).$$

⁴ $\hat{V}_{\beta}^*$ can also be estimated using a separate model (Richemond et al., 2024), similar to our discussion of SAC-style algorithms.

However, it is unclear whether this estimate is accurate enough, particularly when the rollout policy π_θ differs significantly from the optimal policy π_β^* . In contrast, TBRM leverages 1) Bellman equation in KL-regularized RL (see, e.g., [Section 2.3](#)) and 2) recent advances in change of trajectory measure, which allows us to directly use V_θ instead of requiring $\widehat{V}_\beta^*$. This approach enables TBRM to provably converge to a near-optimal policy using only a single rollout per prompt, while maintaining the advantages of a direct optimization algorithm.

Comparing TBRM with Pair-Wise Direct Alignment Algorithms. Under the perspective above, we can view the motivation of pair-wise direct alignment algorithms as using a pair of responses from the same question s_1 to cancel the need for $V_\theta(s_1)$, which leads to the following objective:

$$\operatorname{argmin}_{\theta} \mathcal{L}_{\mathcal{D}}^{\text{DAA-pair}}(\theta) = \frac{1}{|\mathcal{D}|} \sum_{(\tau, \tau') \in \mathcal{D}} \left(\sum_{h=1}^H \log \frac{\pi_\theta(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} - \sum_{h=1}^H \log \frac{\pi_\theta(a'_h | s'_h)}{\pi_{\text{ref}}(a'_h | s'_h)} - \frac{r(\tau)}{\beta} + \frac{r(\tau')}{\beta} \right)^2.$$

Comparing with TBRM, the pair-wise direct alignment algorithms are basically optimizing the difference between the Bellman residuals of two trajectories.

F Hard Instances for Iterative Algorithms

In this section, we demonstrate the advantages of direct optimization (TBRM) over its iterative variant using a simple but illustrative hard instance.

By the autoregressive function approximation definition, we can rewrite the loss function of TBRM as follows:

$$\begin{aligned} \mathcal{L}_{\mathcal{D}}^{\text{TBRM}}(\theta) &= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(\text{logit}_{\theta}(s_1, a_1) - \frac{r(\tau)}{\beta} - \log \pi_{\text{ref}}(\tau) + \sum_{h=2}^H \log \pi_\theta(a_h | s_h) \right)^2 \\ &= \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(V_\theta(s_1) - \frac{r(\tau)}{\beta} + \sum_{h=1}^H \log \frac{\pi_\theta(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right)^2. \end{aligned}$$

A typical iterative variant of this approach can be formulated as:

$$\theta_{t+1} \leftarrow \operatorname{argmin}_{\theta} \frac{1}{|\mathcal{D}|} \sum_{\tau \in \mathcal{D}} \left(V_{\theta_t}(s_1) - \frac{r(\tau)}{\beta} + \sum_{h=1}^H \log \frac{\pi_\theta(a_h | s_h)}{\pi_{\text{ref}}(a_h | s_h)} \right)^2, \quad (14)$$

where V_{θ_t} is fixed from the previous iteration while optimizing for θ .

To illustrate the difference between these approaches, we consider a simple 2-arm bandit problem where $r(a_1) = 1$ and $r(a_2) = 0$. We will show that, even at the population level, the iterative algorithm becomes trapped at a suboptimal solution, whereas TBRM converges to the globally optimal solution.

For this example, our Q-function class contains only two elements: $Q^\dagger = (0, 0)$ and $Q^* = (10, 0)$, corresponding to the uniform policy and optimal policy, respectively.

Suppose at iteration t we have $Q_t = Q^\dagger = (0, 0)$, with temperature parameter $\beta = 0.1$, and data uniformly distributed over actions. The loss for the next iteration becomes:

$$\ell_t(Q) := (1 - Q(a_1) + V_Q - V_t)^2 + (-Q(a_2) + V_Q - V_t)^2,$$

where $V_t = \beta \log \sum_a \exp(Q_t(a)/\beta)$ represents the value function from the current iteration.

In this setting, we can verify that $\ell_t(Q^\dagger) < \ell_t(Q^*)$, meaning the iterative algorithm will select $Q_{t+1} = Q_t = Q^\dagger = (0, 0)$ and remain stuck at this suboptimal solution. In contrast, Q^* is the global minimizer of the TBRM loss by definition, demonstrating the advantage of direct optimization over the iterative approach.

G Implementation Details

G.1 Training Details

To ensure rigorous and reproducible experimentation, we employ standardized and universally adopted hyperparameter settings, as detailed in [Section 4.2](#). For baselines, we adhere closely to the recommended hyperparameter configurations as presented in VERL. Specifically, PPO training utilizes a learning rate of 1×10^{-6} for the actor policy and 1×10^{-5} for the critic policy. We incorporate a KL divergence coefficient and an entropy regularization coefficient of 0.001 for PPO. The clip ratio for the actor loss function is set to 0.2. For the GRPO baseline, we maintain the same KL divergence coefficient as PPO for the KL regularization term. To balance computational efficiency and performance, we utilize a prompt batch size of 128 and a maximum response length of 2048 tokens per training iteration. All generated responses are sampled using a temperature parameter of 1.0. For the TBRM method, the parameter β is consistently set to 0.002 across all experimental conditions. The learning rate for TBRM experiments is 2.5×10^{-6} , with the exception of the TBRM_{n=1} with Qwen2.5-Math-7B model adopting a learning rate of 2×10^{-6} . All models are trained for a total of 760 steps.

All experiments are conducted on the same platform featuring 4x H100 80GB GPUs.

G.2 Training Efficiency of TBRM

We compare the resource cost of TBRM, GRPO, and PPO in [Section 4.3](#) by examining wall-clock time and maximal GPU memory usage. The labels for the time segments used in [Figure 2\(b\)](#) are detailed in [Appendix G.2](#). To ensure a fair comparison, all experiments were conducted on the same platform featuring 4x H100 GPUs, and all configurations were standardized. Specifically, we employed vllm as the rollout backend and set `gpu_memory_utilization` to 0.4. For policy updates, we set `micro_batch_size_per_gpu` to 1, and for calculating log probabilities for both $\pi_{\theta_{\text{old}}}$ and π_{ref} , we also used a `micro_batch_size_per_gpu` of 1. The value function model update for PPO also utilized a `micro_batch_size_per_gpu` of 1. We use Qwen2.5-Math-7B as the base model, with a prompt batch size of 128.

Name	Description	Involved Algorithms
ref_prob	Computing $\pi_{\text{ref}}(a_t s_t)$	TBRM, GRPO, PPO
old_prob	Computing $\pi_{\theta_{\text{old}}}(a_t s_t)$	GRPO, PPO
update_policy	Updating the policy parameter θ	TBRM, GRPO, PPO
ppo_update_value	Updating the value function model V_ϕ	PPO
ppo_value	Computing $V_\phi(s_t)$	PPO
rollout	Sampling trajectories from the prompt set	TBRM, GRPO, PPO
other	Miscellaneous computations, e.g., rule-based reward $r(\tau)$, advantage (for GRPO and PPO), etc.	

Table 3: Description of labels of time segments in [Figure 2\(b\)](#).

G.3 Prompt Templates

Qwen-Math Prompt Template. We use the default prompt template of Qwen2.5-Math in the main experiments ([Section 4.2](#)).

```
<|im_start|>system
Please reason step by step, and put your final answer within \boxed{\}.<|im_end|>
<|im_start|>user
question<|im_end|>
<|im_start|>assistant
```

DeepSeek-R1 Prompt Template. We use DeepSeek-R1 prompt template in the experiment discussed in [Section 4.3](#).

A conversation between User and Assistant. The user asks a question, and the Assistant solves it. The Assistant first thinks about the reasoning process in the mind and then provides the user

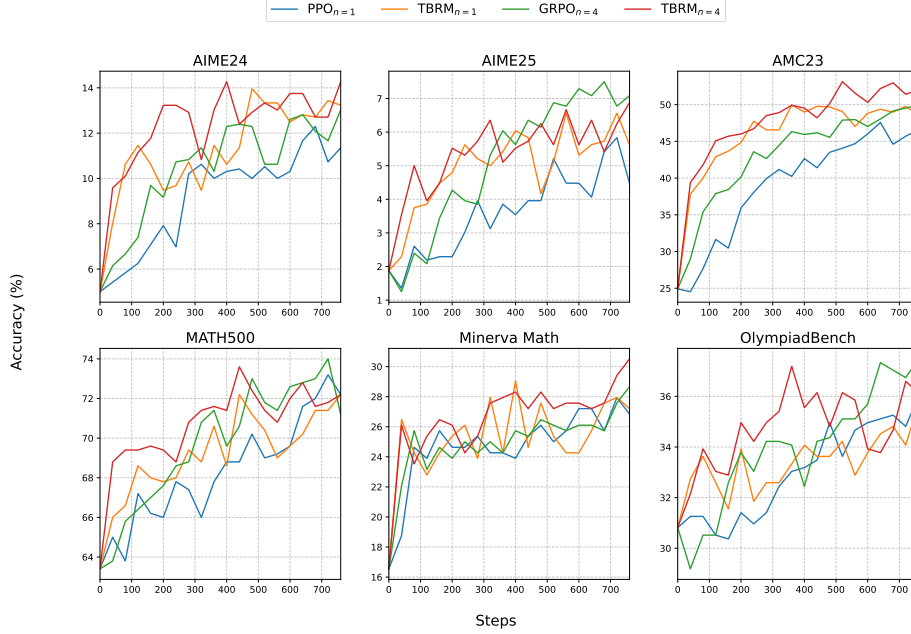


Figure 3: Learning accuracy across benchmarks with model Qwen2.5-Math-1.5B.

with the answer. The reasoning process and answer are enclosed within `<think> ... </think>` and `<answer> ... </answer>` tags, respectively, i.e., `<think>` reasoning process here `</think>` `<answer>` answer here `</answer>`.

User: **question**

Assistant:

H Additional Experimental Results

H.1 Training-Time Performance

Following the setup described in [Section 4.1](#), we train and evaluate models of various sizes. Their performance on benchmarks during training is presented in [Figure 3](#) and [Figure 4](#).

H.2 Qualitative Analysis

We delve into the responses of the TBRM model and observe intriguing reasoning behaviors. In this subsection, we highlight emergent reasoning patterns that the model develops during training, which contribute to its improved reasoning capabilities. Unless otherwise noted, all responses presented here are sampled from Qwen2.5-Math-7B-TBRM_{n=4}.

Decomposition. The model demonstrates the ability to break down complex problems into smaller, manageable subproblems, solving each sequentially to construct the final solution. This hierarchical strategy reflects a structured and efficient approach to reasoning.

Decomposition

(MATH500) For what real values of x is $-4 < x^4 + 4x^2 < 21$ satisfied? Express your answer in interval notation.

To solve the inequality $-4 < x^4 + 4x^2 < 21$, we need to break it into two parts and solve each part separately:

1. Solve the inequality $x^4 + 4x^2 > -4$.
2. Solve the inequality $x^4 + 4x^2 < 21$.

Step 1: Solve $x^4 + 4x^2 > -4$

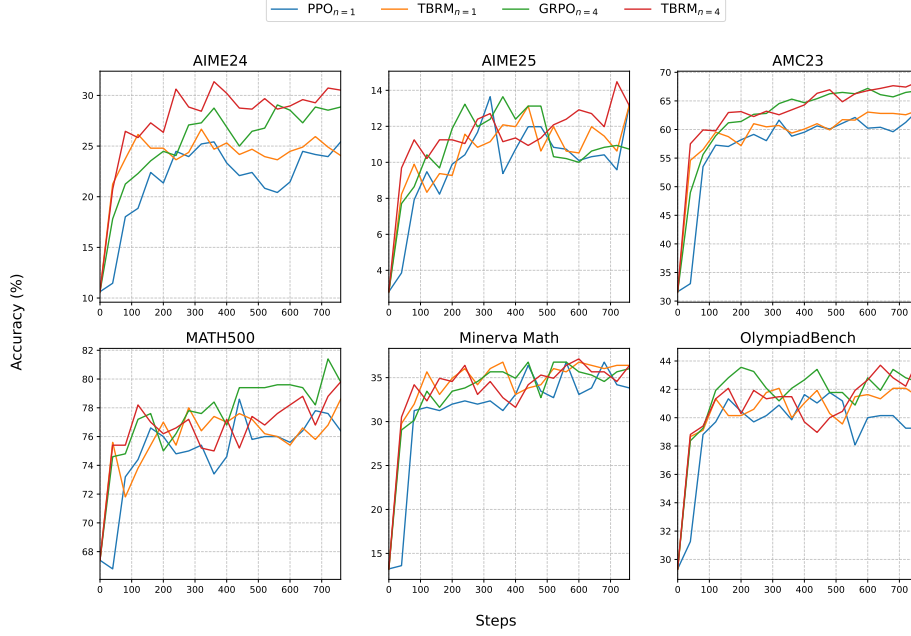


Figure 4: Learning accuracy across benchmarks with model Qwen2.5-Math-7B.

The expression $x^4 + 4x^2$ is always non-negative for all real numbers x because both x^4 and $4x^2$ are non-negative. Therefore, $x^4 + 4x^2 \geq 0 > -4$ for all x . So, the inequality $x^4 + 4x^2 > -4$ is satisfied for all real numbers x .

Step 2: Solve $x^4 + 4x^2 < 21$

To solve $x^4 + 4x^2 < 21$, let's make a substitution. Let $y = x^2$. Then the inequality becomes $y^2 + 4y < 21$. Rearranging terms, we get $y^2 + 4y - 21 < 0$. We need to find the roots of the equation $y^2 + 4y - 21 = 0$ to determine the interval for y .

The quadratic equation $y^2 + 4y - 21 = 0$ can be factored as $(y + 7)(y - 3) = 0$. So, the roots are $y = -7$ and $y = 3$. Since $y = x^2$ and x^2 is always non-negative, we only consider the positive root $y = 3$. Thus, we need to solve $0 \leq y < 3$.

Substituting back $y = x^2$, we get $0 \leq x^2 < 3$. This means $-\sqrt{3} < x < \sqrt{3}$.

Combining the Results

Since the first inequality $x^4 + 4x^2 > -4$ is satisfied for all real numbers x , the solution to the original inequality $-4 < x^4 + 4x^2 < 21$ is determined by the second inequality. Therefore, the solution is $-\sqrt{3} < x < \sqrt{3}$.

The final answer in interval notation is:

$$(-\sqrt{3}, \sqrt{3})$$

Enumeration. The model adopts an exhaustive strategy by systematically iterating through possibilities or repeatedly applying a specific operation until the correct answer emerges. This approach is particularly useful in combinatorial or search-based problems.

Enumeration

(AIME24) Alice chooses a set A of positive integers. Then Bob lists all finite nonempty sets B of positive integers with the property that the maximum element of B belongs to A . Bob's list has 2024 sets. Find the sum of the elements of A .

To solve this problem, we need to understand the relationship between the set A and the number of nonempty subsets B of positive integers where the maximum element of B belongs to A .

If $A = \{a_1, a_2, \dots, a_n\}$, then for each element $a_i \in A$, the number of subsets B where the maximum element is a_i is given by the number of nonempty subsets of $\{1, 2, \dots, a_i - 1\}$, which is $2^{a_i-1} - 1$. Summing over all elements in A , the total number of such subsets B is:

$$(2^{a_1} - 1) + (2^{a_2} - 1) + \dots + (2^{a_n} - 1) = 2^{a_1} + 2^{a_2} + \dots + 2^{a_n} - n$$

Given that this total is 2024, we can write:

$$2^{a_1} + 2^{a_2} + \dots + 2^{a_n} - n = 2024$$

or

$$2^{a_1} + 2^{a_2} + \dots + 2^{a_n} = 2024 + n$$

To find the appropriate values for $a_1, a_2, \dots, a_n$, we need to find a combination of powers of 2 that sum up to just over 2024. We can start by finding the largest power of 2 less than or equal to 2024. The largest power of 2 less than 2024 is $2^{10} = 1024$.

By subtracting 1024 from 2024, we get 1000. The next largest power of 2 less than or equal to 1000 is $2^9 = 512$. Subtracting 512 from 1000 gives 488. The next largest power of 2 less than or equal to 488 is $2^8 = 256$, leaving us with 232. The next largest power of 2 less than or equal to 232 is $2^7 = 128$, leaving us with 104. The next largest power of 2 less than or equal to 104 is $2^6 = 64$, leaving us with 40. The next largest power of 2 less than or equal to 40 is $2^5 = 32$, leaving us with 8. The next largest power of 2 less than or equal to 8 is $2^3 = 8$, leaving us with 0. So, the sum of the powers of 2 is:

$$1024 + 512 + 256 + 128 + 64 + 32 + 8 = 2024$$

This means that we have $n = 7$ elements in set A , and these elements are $\{11, 10, 9, 8, 7, 6, 4\}$ because the powers of 2 correspond to $2^{10}, 2^9, 2^8, 2^7, 2^6, 2^5, 2^3$.

So, the set $A = \{11, 10, 9, 8, 7, 6, 4\}$. The sum of the elements of A is:

$$11 + 10 + 9 + 8 + 7 + 6 + 4 = 55$$

The final answer is:

55

Verification. The model exhibits self-monitoring behavior by verifying intermediate results during problem solving. This helps catch early mistakes and reinforces correct reasoning steps, enhancing both accuracy and reliability.

Verification

(MATH500) What is the length, in units, of the radius of a sphere whose volume and surface area, in cubic units and square units, respectively, are numerically equal?

To find the radius of a sphere whose volume and surface area are numerically equal, we need to set up the equations for the volume and surface area of a sphere and then solve for the radius r .

The volume V of a sphere is given by: $V = \frac{4}{3}\pi r^3$

The surface area A of a sphere is given by: $A = 4\pi r^2$

We are given that the volume and surface area are numerically equal, so we can set the two equations equal to each other: $\frac{4}{3}\pi r^3 = 4\pi r^2$

To solve for r , we can simplify the equation by dividing both sides by $4\pi r^2$ (assuming $r \neq 0$): $\frac{r}{3} = 1$ $r = 3$

So, the radius of the sphere is $r = 3$. Let's verify this by substituting $r = 3$ back into the volume and surface area formulas to ensure they are indeed equal.

The volume when $r = 3$ is: $V = \frac{4}{3}\pi(3)^3 = \frac{4}{3}\pi \cdot 27 = 36\pi$

The surface area when $r = 3$ is: $A = 4\pi(3)^2 = 4\pi \cdot 9 = 36\pi$

Since the volume and surface area are both 36π , our solution is correct. The length of the radius of the sphere is $\boxed{3}$.

Backtracking. Upon encountering contradictions or errors, the model is able to reassess its prior steps, revise incorrect assumptions, and pursue an alternative line of reasoning. This reflective capability mirrors human-like problem-solving resilience.

Backtracking

(MATH500) The asymptotes of a hyperbola are $y = 2x - 3$ and $y = 17 - 2x$. Also, the hyperbola passes through the point $(4, 7)$. Find the distance between the foci of the hyperbola.

To find the distance between the foci of the hyperbola, we need to follow these steps:

1. **Determine the center of the hyperbola:** The center of the hyperbola is the intersection point of the asymptotes. To find this point, we solve the system of equations given by the asymptotes:

$$\begin{aligned} y &= 2x - 3 \\ y &= 17 - 2x \end{aligned}$$

Setting the two equations equal to each other:

$$2x - 3 = 17 - 2x$$

Solving for x :

$$4x = 20 \implies x = 5$$

Substituting $x = 5$ back into one of the equations to find y :

$$y = 2(5) - 3 = 10 - 3 = 7$$

So, the center of the hyperbola is $(5, 7)$.

2. **Determine the slopes of the asymptotes:** The slopes of the asymptotes are 2 and -2 . This indicates that the hyperbola is oriented with its transverse axis vertical or horizontal. Since the slopes are ± 2 , the hyperbola is of the form:

$$\frac{(x - h)^2}{b^2} - \frac{(y - k)^2}{a^2} = -1 \quad \text{or} \quad \frac{(y - k)^2}{a^2} - \frac{(x - h)^2}{b^2} = 1$$

where the slopes of the asymptotes are $\pm \frac{a}{b}$ or $\pm \frac{b}{a}$. In this case, $\frac{a}{b} = 2$ or $\frac{b}{a} = 2$. Let's assume the hyperbola is of the form:

$$\frac{(y - 7)^2}{a^2} - \frac{(x - 5)^2}{b^2} = 1$$

with $\frac{a}{b} = 2$, so $a = 2b$.

3. **Use the point $(4, 7)$ to find a and b :** The hyperbola passes through the point $(4, 7)$. Substituting $x = 4$ and $y = 7$ into the equation of the hyperbola:

$$\frac{(7 - 7)^2}{a^2} - \frac{(4 - 5)^2}{b^2} = 1$$

Simplifying, we get:

$$0 - \frac{1}{b^2} = 1 \implies -\frac{1}{b^2} = 1 \implies b^2 = -1$$

This is incorrect, so let's assume the hyperbola is of the form:

$$\frac{(x - 5)^2}{b^2} - \frac{(y - 7)^2}{a^2} = -1$$

with $\frac{a}{b} = 2$, so $a = 2b$. Substituting $x = 4$ and $y = 7$:

$$\frac{(4 - 5)^2}{b^2} - \frac{(7 - 7)^2}{a^2} = -1$$

Simplifying, we get:

$$\frac{1}{b^2} - 0 = -1 \implies \frac{1}{b^2} = -1 \implies b^2 = 1$$

So, $b = 1$ and $a = 2b = 2$.

4. **Find the distance between the foci:** The distance between the foci of a hyperbola is $2c$, where $c = \sqrt{a^2 + b^2}$. In this case:

$$c = \sqrt{a^2 + b^2} = \sqrt{2^2 + 1^2} = \sqrt{4 + 1} = \sqrt{5}$$

So, the distance between the foci is:

$$2c = 2\sqrt{5}$$

The final answer is:

$$\boxed{2\sqrt{5}}$$

H.3 Ablation Study: Classical BRM on LLMs

As discussed in Section 3, in the context of LLMs, directly minimizing BRM may not be possible because the token-level reward signal is either unavailable (e.g. if we assign the outcome reward to the final token) or very sparse. To highlight the advantages of the trajectory-level approach, we implemented classical BRM, whose loss function is defined in Eq. (3), using the same hyperparameters as in our main setup on Qwen2.5-Math-1.5B.

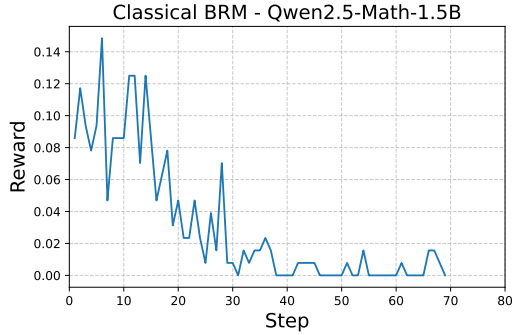


Figure 5: Training reward with classical BRM on model Qwen2.5-Math-1.5B.

Figure 5 shows that training reward quickly collapses, and we observe that the model outputs become random and meaningless. Intuitively, this degradation occurs because BRM has to propagate the sparse reward signal, which only receives at the final token, back through multiple token-wise regressions, whereas the TBRM provides a better implicit credit assignment through a single trajectory-level objective.

H.4 Tasks Beyond Mathematical Problems

To demonstrate the generalizability of our method beyond mathematical tasks, we evaluate TBRM on five tasks from the reasoning-gym (Stojanovski et al., 2025) under the *graphs* category: *course_schedule*, *family_relationships*, *largest_island*, *quantum_lock*, and *shortest_path*. These tasks are naturally represented as graphs, consisting of nodes and edges, and typically require traversing connections to identify relationships, compute optimal paths, or determine reachable components. They involve reasoning patterns that differ significantly from those in mathematical tasks.

Method	<i>course_schedule</i>	<i>family_relationships</i>	<i>largest_island</i>	<i>quantum_lock</i>	<i>shortest_path</i>	Average
Qwen2.5-Math-1.5B	29.5	3.0	11.0	5.5	0.0	9.8
GRPO $n = 4$	54.0	84.0	34.0	30.5	26.0	45.7
TBRM $n = 4$	60.0	80.0	38.0	27.0	31.0	47.2

Table 4: Performance of GRPO and TBRM on various tasks from reasoning-gym, category *graphs*.

We construct a training set of 10,000 problems, with 2,000 questions per task, and a test set of 500 problems, comprising 100 questions from each task. For both training and evaluation, we use the official verifiers provided by reasoning-gym to compute rewards. Our experiments are conducted on

Qwen2.5-Math-1.5B using both TBRM and GRPO, with a prompt batch size of 1024 and 4 sampled responses per question ($n = 4$). Models are trained for 100 steps. All evaluations are conducted using greedy decoding. Results in Table 4 demonstrate that TBRM generalizes well to diverse reasoning tasks and performs on par with GRPO.

H.5 TBRM with More Rollouts

To demonstrate that TBRM scales effectively with increasing number of sampled responses per prompt, we rerun GRPO and TBRM using most hyperparameters from DAPO (Yu et al., 2025). Specifically, we used a prompt batch size of 512 and generated $n = 16$ responses per prompt. For GRPO, we set the microbatch size to 512, resulting in 16 updates per training step. The experiments were conducted on the Qwen2.5-Math-7B model, following the same evaluation pipeline described in our paper. Both algorithms were trained for 100 steps. Table 5 shows that TBRM remains comparable to GRPO under these aligned settings.

Method	AIME24	AIME25	AMC23	MATH500	Minerva Math	OlympiadBench
	Avg@32	Avg@32	Avg@32	Avg@1	Avg@1	Avg@1
Qwen2.5-Math-7B	10.6	2.8	31.6	67.4	13.2	29.3
GRPO $n = 16$	26.6	11.0	61.8	77.8	32.7	40.4
TBRM $n = 16$	27.9	10.9	62.8	76.4	33.5	39.9

Table 5: Performance of GRPO and TBRM with $n = 16$.

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: Claims made in the abstract and introduction clearly reflect our main contributions, which are later elaborated in [Section 3](#) and [Section 4](#).

Guidelines:

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

2. Limitations

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

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

We explicitly discuss our work’s limitations in [Section 3](#) and their proofs in [Appendix D](#). Specifically, we note the reliance on repeated inference runs to mitigate LLM randomness on small test sets ([Section 4.1](#)), and we analyze our training pipeline’s computational efficiency and scalability in [Section 4.3](#).

Guidelines:

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

judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

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

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

Justification: In our theoretical results (e.g., [Theorem 2](#)), we clearly state the assumption and postpone the complete proof to [Appendix D](#).

Guidelines:

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

4. Experimental result reproducibility

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

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

Justification: We thoroughly discuss the experimental setup in [Section 4.1](#) and [Appendix G](#).

Guidelines:

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

- (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
- (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

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

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

Justification: The codebase for our experiments is publicly available at <https://github.com/rlx-lab/TBRM>. All datasets used in this work are public.

Guidelines:

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

6. Experimental setting/details

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

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

Justification: [Section 4.1](#) and [Appendix G](#) includes all information needed to understand and replicate our results.

Guidelines:

- The answer NA means that the paper does not include experiments.

- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

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

Answer: [No]

Justification: Due to time and resource constraints, we did not perform repeated training runs, and therefore do not report error bars or confidence intervals based on training variability. For evaluation, we include repeated inference runs to mitigate intrinsic LLM randomness on small test sets, which enhances reproducibility. However, we do not report formal error bars, statistical significance tests, or confidence intervals in our results.

Guidelines:

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

8. Experiments compute resources

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

Answer: [Yes]

Justification: We discuss the platform used for our experiments in [Appendix G](#).

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: The research conducted in the paper 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: [No]

Justification: This paper introduces a fundamental, generic algorithm for LLM training that, in itself, has no direct societal impact, as any such effects are tied to the independent applications of LLM technology, not this specific algorithmic work.

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: We utilize public datasets that only involve math tasks and we do not see the potential risks for misuse.

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 clarify the required information in [Section 4.1](#).

Guidelines:

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

13. New assets

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

Answer: [Yes]

Justification: The codebase for TBRM is publicly available at <https://github.com/r1x-lab/TBRM> along with detailed documentation.

Guidelines:

- The answer NA means that the paper does not release new assets.

- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

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

Answer: [NA]

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

Guidelines:

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

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

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

Answer: [NA]

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

Guidelines:

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

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.