

UNREGULARIZED LINEAR CONVERGENCE IN ZERO-SUM GAME FROM PREFERENCE FEEDBACK

Anonymous authors

Paper under double-blind review

ABSTRACT

Aligning large language models (LLMs) with human preferences has proven effective for enhancing model capabilities, yet standard preference modeling using the Bradley-Terry model assumes transitivity, overlooking the inherent complexity of human population preferences. Nash learning from human feedback (NLHF) addresses this by framing non-transitive preferences as a two-player zero-sum game, where alignment reduces to finding the Nash equilibrium (NE). However, existing algorithms typically rely on regularization, incurring unavoidable bias when computing the duality gap in the original game. In this work, we provide the first convergence guarantee for Optimistic Multiplicative Weights Update (OMWU) in NLHF, showing that it achieves last-iterate linear convergence after a burn-in phase whenever an NE with full support exists, with an instance-dependent linear convergence rate to the original NE, measured by duality gaps. Compared to prior results in Wei et al. (2020), we do not require the assumption of NE uniqueness. Our analysis identifies a novel marginal convergence behavior, where the probability of rarely played actions grows exponentially from exponentially small values, enabling exponentially better dependence on instance-dependent constants than prior results. Experiments corroborate the theoretical strengths of OMWU in both tabular and neural policy classes, demonstrating its potential for LLM applications.

1 INTRODUCTION

The emergence of large language models (LLMs) raises the challenge of aligning model outputs with human preferences. Traditional reinforcement learning (RL) methods require explicit reward functions, which are often difficult to obtain directly. The Bradley-Terry (BT) model (Bradley & Terry, 1952) addresses this by assigning a scalar reward $r(x, y)$ to a response y given a prompt x , based on collected preference data. Specifically, the probability of preferring y over y' on prompt x is modeled as $\mathcal{P}(y > y' \mid x) = \sigma(r(x, y) - r(x, y'))$, where $\sigma(t) = 1/(1 + e^{-t})$. Building on this framework, Direct Preference Optimization (DPO, Rafailov et al. (2023)) leverages the closed-form solution to fine-tune the policy.

However, the BT model implicitly assumes that human preferences are transitive. Empirical studies (May, 1954) provide evidence of intransitivity at the population level. To address this limitation, Munos et al. (2023) introduced *Nash Learning from Human Feedback* (NLHF), which formulates alignment as finding a Nash equilibrium (NE) of the preference game. Since $\mathcal{P}(y > y' \mid x) + \mathcal{P}(y' > y \mid x) = 1$, the problem naturally reduces to a two-player constant-sum game.

For NLHF to be deployable in the LLM setting, the algorithm must satisfy additional requirements:

First, it must achieve *last-iterate convergence*, meaning the final policy produced by training is close to a Nash equilibrium. In contrast, *average-iterate convergence* guarantees only that the time-averaged policy approximates equilibrium, which is impractical for deployment.

Second, the algorithm should be implementable with neural network parameterizations. This rules out certain methods such as Optimistic Gradient Descent Ascent (OGDA, Wei et al. (2020)), which requires orthogonal projections onto the probability simplex.

Table 1: Comparison of OMWU results in the NLHF setting. The comparison between "exponential dependence" and "polynomial dependence" is explained towards end of Section 1.1.

Assumptions	Unique Equilibrium	Multiple Equilibria Full-Support Equilibrium
Full-Support Equilibrium Exists	Linear convergence exponential dependence (Wei et al., 2020) polynomial dependence (ours)	Linear convergence polynomial dependence (ours)

Third, the algorithm should avoid nested optimization, e.g., $\theta^{(t+1)} = \arg \min_{\theta} \mathcal{L}_{\text{inner}}(\theta; \theta^{(t)})$, which is widely adopted by almost all prior works in NLHF (Munos et al., 2023; Swamy et al., 2024; Wu et al., 2024; Shani et al., 2024; Zhang et al., 2024; Wang et al., 2024; Zhang et al., 2025). In practice, there is a trade-off between the number of gradient descent steps on $\mathcal{L}_{\text{inner}}$ to approximate $\hat{\theta}^{(t+1)}$ and the approximation error, while more steps incur significantly higher computational overhead.

Motivated by these requirements, we revisit the *Optimistic Multiplicative Weights Update* (OMWU, Daskalakis & Panageas (2018)) algorithm to assess its potential for more complex tasks such as NLHF. Wei et al. (2020) proved that under the assumption of a unique NE, OMWU achieves last-iterate linear convergence after a burn-in period, with an instance-dependent convergence rate for two-player zero-sum games. However, general preference matrices typically have infinitely many NEs. Furthermore, their analysis yields burn-in times and convergence rates with *exponential* dependence on instance-dependent constants, which may be prohibitive for NLHF tasks.

1.1 MAIN CONTRIBUTIONS

This paper provides new theoretical guarantees for OMWU in the NLHF setting. Our contributions are as follows:

- **Improved efficiency under milder assumptions.** Under the natural assumption that a full-support Nash equilibrium exists (rather than the uniqueness of NE), we establish polynomial (rather than exponential) dependence on instance-dependent constants for both convergence rate and burn-in time, while maintaining last-iterate linear convergence over the number of updates. This substantially reduces concerns about the practicality of OMWU for NLHF.

- **New analytical framework for escaping behavior.** We introduce a novel method for analyzing how OMWU escapes undesirable regions of the strategy space. To our knowledge, this is the first work to formalize such escaping dynamics, even beyond OMWU, opening a new direction for understanding dynamics in game-theoretical problems.

We compare our results with those of Wei et al. (2020) in Table 1.

Under the unique equilibrium assumption, the orders of the burn-in time and convergence rate in Wei et al. (2020) are

$$O\left(\frac{\ln(A)}{\eta^4 C_P^2 \exp(-4 \ln(A)/\varepsilon)}\right) \text{ and } O(\eta^2 C_P^2 \exp(-3 \ln(A)/\varepsilon))$$

when initialization is uniform, which is in stark contrast with our orders

$$O\left(\frac{D_{\text{KL}}(\pi^* \|\hat{\pi}^{(1)})^3 (\eta L)^2 A^2}{(1 - 4\eta^2 L^2) C_P^4 \varepsilon^6} \cdot \max\left\{1, \frac{C_P^2 \varepsilon^4 A}{(\eta L)^2}\right\}^3\right) \text{ and } O(\eta^2 \varepsilon^3 C_P^2),$$

which does not depend exponentially on $\ln(A)/\varepsilon$. Here, A is the size of action space, $\hat{\pi}^{(1)}$ is the initialization, π^* is the equilibrium, η is the learning rate, and readers may refer to Section 3 for the precise definitions of ε , L , and C_P .

1.2 PAPER OVERVIEW

We begin by introducing the NLHF framework and the OMWU algorithm in Section 3. In Section 4, we present our main theoretical result, together with an overview of the convergence behavior and the analytical techniques used in our proofs. We then validate our theory with empirical results in Section 5.

Table 2: Comparison of algorithms with convergence guarantees to δ -NE in NLHF. **Convergence to δ -NE**: the number of steps required to find a policy with duality gap at most δ . $\tilde{O}$ hides $\text{poly}(\log(|\mathbb{A}|/\eta))$ factors. When $\text{poly}(\delta^{-1})$ terms exist, $\text{poly}(\log(\delta^{-1}))$ factors are also hidden. **Last-iterate convergence**: “Yes” indicates that the convergence rate applies to the final policy; “No” indicates that it applies only to the average policy. **Efficient update**: “Yes” indicates a single-step gradient update suffices, while “No” indicates nested optimization is required.

Algorithm	Convergence to δ -NE	Last-iterate Convergence	Efficient Update
OMD	$\tilde{O}(\delta^{-2})$	No	Yes
SPPO (Wu et al., 2024)	$\tilde{O}(\delta^{-2})$	No	No
MPO (Wang et al., 2024)	$\tilde{O}(\delta^{-2})$	Yes	No
ONPO (Zhang et al., 2025)	$\tilde{O}(\delta^{-1})$	No	No
EGPO (Zhou et al., 2025)	$\tilde{O}(\delta^{-1})$	Yes	Yes
OMWU	$\tilde{O}(\log(\delta^{-1}))$ (linear)	Yes	Yes

2 RELATED WORKS

NLHF. Since the introduction of the NLHF framework by Munos et al. (2023), a growing body of work has studied algorithms for solving NLHF, most of which rely on regularization. The Nash-MD algorithm proposed in Munos et al. (2023) achieves linear convergence in the regularized setting. The MPO algorithm Wang et al. (2024) provides an $\tilde{O}(\delta^{-2})$ last-iterate convergence to a δ -NE. More recently, the EGPO algorithm of Zhou et al. (2025) achieves an improved $\tilde{O}(\delta^{-1})$ convergence rate compared to MPO, while eliminating the need for nested optimization. Additional studies of NLHF include Wu et al. (2024); Swamy et al. (2024); Ye et al. (2024); Rosset et al. (2024); Calandriello et al. (2024); Zhang et al. (2024; 2025). Table 2 summarizes the theoretical guarantees of several algorithms in the NLHF setting.

Computing Nash equilibria in two-player zero-sum games. Computing Nash equilibria in two-player zero-sum games was a central research topic long before NLHF was proposed. Online Mirror Descent (OMD, Cesa-Bianchi & Lugosi (2006); Lattimore & Szepesvári (2020)), originally developed for online convex learning, naturally applies to this setting but guarantees only average-iterate convergence. In contrast, OMWU and Optimistic Gradient Descent Ascent (OGDA) have been shown to enjoy instance-dependent linear convergence (Wei et al., 2020). Among these methods, OGDA uses direct parametrization (using $\theta_{x,y}$ directly as the probability), making it less applicable; however, its convergence holds even when the equilibrium is non-unique.

OMWU. The Optimistic Multiplicative Weights Update (OMWU) algorithm was first proposed by Daskalakis & Panageas (2018), who established its last-iterate convergence under the uniqueness assumption, though without a convergence rate. Later, Wei et al. (2020) proved linear convergence at an instance-dependent rate for general saddle-point optimization problems under the same assumption. On the negative side, Cai et al. (2024) showed that last-iterate convergence must be arbitrarily slow for a broad class of algorithms including OMWU, even in simple two-action settings. Other studies of OMWU include Lee et al. (2021); Daskalakis et al. (2021); Anagnostides et al. (2022).

3 PRELIMINARIES

3.1 MULTI-ARMED BANDITS

A multi-armed bandit has an action space $\mathbb{A}$ (the reward function is irrelevant in our setting). A contextual bandit additionally has a context space $\mathbb{X}$, where the agent chooses an action $a \in \mathbb{A}$ for each context $x \in \mathbb{X}$. In our fine-tuning setting, $\mathbb{X}$ corresponds to the prompt space and $\mathbb{A}$ to the response space. For clarity, we state our theory in the multi-armed bandit setting.

A policy π is a probability distribution over $\mathbb{A}$. For computational convenience, we parametrize π with a vector θ such that

$$\pi_a = \frac{\exp(\theta_a)}{\sum_{a' \in \mathbb{A}} \exp(\theta_{a'})}.$$

Note that θ is a valid parametrization of π if and only if it differs from $\log \pi$ by a constant shift.

3.2 NLHF

In the NLHF problem (see Munos et al. (2023)), each pair of actions $a, a' \in \mathbb{A}$ is associated with a probability $\mathcal{P}(a > a')$, representing the probability that a is preferred over a' . Clearly, $\mathcal{P}(a > a') + \mathcal{P}(a' > a) = 1$ when $a \neq a'$, and we set $\mathcal{P}(a > a) = \frac{1}{2}$ so that this relation also holds for $a = a'$.

We define the preference matrix P as

$$P_{a,a'} = \mathcal{P}(a > a') - \frac{1}{2}.$$

By construction, P is skew-symmetric ($P + P^\top = 0$)¹.

The goal of NLHF is to find a policy π^* that maximizes its probability of being preferred against an adversarial policy π' , i.e.,

$$\pi^* = \arg \max_{\pi} \min_{\pi'} P(\pi > \pi'),$$

where

$$P(\pi > \pi') = \mathbb{E}_{a \sim \pi, a' \sim \pi'} P(a > a') = \pi^\top P \pi' + \frac{1}{2}.$$

By von Neumann's minimax theorem v. Neumann (1928),

$$\max_{\pi} \min_{\pi'} \pi^\top P \pi' = \min_{\pi'} \max_{\pi} \pi^\top P \pi',$$

and since P is skew-symmetric,

$$\max_{\pi} \min_{\pi'} \pi^\top P \pi' = \min_{\pi'} \max_{\pi} \pi^\top P \pi' = 0 = a^\top P a \quad \forall a.$$

Thus, if π^* is an optimizer of the minimax objective, then

$$\min_{\pi'} (\pi^*)^\top P \pi' = \max_{\pi} \pi^\top P \pi^* = 0.$$

In this case, we call π^* a Nash equilibrium of P^2 .

We denote the set of all Nash equilibria by $\mathbb{M}$. For any policy π , the duality gap is defined as

$$\text{DualGap}(\pi) = \max_{\pi'} (\pi')^\top P \pi - \min_{\pi'} \pi^\top P \pi' = 2 \max_{a \in \mathbb{A}} (P \pi)_a.$$

The duality gap is always nonnegative and equals zero if and only if π is a Nash equilibrium. We say a policy π is δ -NE when $\text{DualGap}(\pi) \leq \delta$.

3.3 OPTIMISTIC MULTIPLICATIVE WEIGHTS UPDATE (OMWU)

The Optimistic Multiplicative Weights Update (OMWU) algorithm was introduced by Daskalakis & Panageas (2018) for solving equilibria in zero-sum games and later extended to saddle-point problems. In the NLHF setting, the algorithm simplifies to

$$\pi^{(t)} = \arg \min_{\pi} \left\{ \eta \langle \pi, P \pi^{(t-1)} \rangle + D_{\text{KL}}(\pi \| \hat{\pi}^{(t)}) \right\},$$

¹Some works define the preference matrix without subtracting $\frac{1}{2}$. This choice has little impact in practice, but in our case, the current definition ensures skew-symmetry.

²Formally, a Nash equilibrium of a zero-sum game matrix P is a pair (π_1, π_2) such that $\pi_1^\top P \pi_2 = \min_{\pi'} \pi_1^\top P \pi' = \max_{\pi'} (\pi')^\top P \pi_2$. For skew-symmetric P , (π_1, π_2) is a Nash equilibrium if and only if (π_1, π_1) and (π_2, π_2) are Nash equilibria. Hence, our notational simplification is justified.

$$\hat{\pi}^{(t+1)} = \arg \min_{\pi} \left\{ \eta \langle \pi, P\pi^{(t)} \rangle + D_{\text{KL}}(\pi \| \hat{\pi}^{(t)}) \right\},$$

where η is the learning rate, and $\pi^{(0)} = \hat{\pi}^{(1)}$ are initializations.

In the parametrized form, OMWU reduces to

$$\theta^{(t)} = \theta^{(t-1)} + \eta P\hat{\pi}^{(t)}, \quad (1)$$

$$\hat{\theta}^{(t+1)} = \theta^{(t)} + \eta P\hat{\pi}^{(t)}. \quad (2)$$

3.4 REGULARIZATION IN NLHF

Although regularization is not part of OMWU, it appears in related algorithms. For a reference policy π_{ref} , define

$$\begin{aligned} \text{Dualgap}_{\beta}(\pi) = & \max_{\pi'} \left((\pi')^{\top} P\pi - \beta D_{\text{KL}}(\pi' \| \pi_{\text{ref}}) + \beta D_{\text{KL}}(\pi \| \pi_{\text{ref}}) \right) \\ & - \min_{\pi''} \left(\pi^{\top} P\pi'' - \beta D_{\text{KL}}(\pi \| \pi_{\text{ref}}) + \beta D_{\text{KL}}(\pi'' \| \pi_{\text{ref}}) \right). \end{aligned}$$

This transforms NLHF into a convex optimization problem in $\text{Dualgap}_{\beta}(\pi)$ at the cost of introducing regularization error.

3.5 ASSUMPTIONS

Throughout, we impose the following assumption on P :

Assumption 1. For every $a \in \mathbb{A}$, there exists a Nash equilibrium π of P such that $\pi_a > 0$.

Remark 1. This is equivalent to the existence of a Nash equilibrium π with $\pi_a > 0$ for all $a \in \mathbb{A}$. It is well known that the set $\mathbb{M}$ of Nash equilibria is convex.

We define the KL projection of a policy π onto $\mathbb{M}$ as

$$p(\pi) = \arg \min_{\pi' \in \mathbb{M}} D_{\text{KL}}(\pi' \| \pi).$$

Under Assumption 1, we obtain the following:

Lemma 1. If π satisfies $\pi_a > 0$ for all $a \in \mathbb{A}$, then $p(\pi)$ is well-defined, unique, and satisfies $p(\pi)_a > 0$ for all a .

Lemma 2. Under Assumption 1, the OMWU algorithm satisfies

$$p(\hat{\pi}^{(1)}) = p(\hat{\pi}^{(2)}) = \dots = p(\hat{\pi}^{(t)}) = \dots$$

Proofs are deferred to Appendix A. By Lemma 2, we define

$$\pi^* = p(\hat{\pi}^{(t)}).$$

Remark 2. If Lemma 2 holds and the sequence $\hat{\pi}^{(t)}$ converges as $t \rightarrow \infty$, then the limit must be π^* . This plays the role of the uniqueness assumption in prior work, allowing us to predict the convergence point without taking infinite steps. Moreover, with uniform initialization, π^* corresponds to the equilibrium with minimal negative entropy.

Finally, we introduce constants used in later proofs:

Definition 1 (Instance-dependent constants).

$$\varepsilon = \min_{a \in \mathbb{A}} \pi_a^*, \quad L = \max_{a, a' \in \mathbb{A}} |P_{a, a'}|, \quad C_P = \min_{\pi \in \Delta(\mathbb{A}) \setminus \mathbb{M}} \frac{\|P\pi\|_{\infty}}{\|\pi - p(\pi)\|_1}.$$

Clearly $\varepsilon > 0$ by Assumption 1, and the proof of $C_P > 0$ is given in Appendix B.

The constant C_P is novel but crucial. This conveys the idea that if some policy π is far from the predicted convergence point $p(\pi)$, then there must be a large update following OMWU algorithm in the parametrized space at some coordinate.

4 MAIN THEOREM AND PROOF SKETCH

4.1 STATEMENT OF MAIN THEOREM

Our main contribution is to establish a last-iterate linear convergence guarantee for OMWU in zero-sum games with a full-support Nash equilibrium. This extends the uniqueness-based result of Wei et al. (2020), providing a broader characterization of OMWU dynamics.

Theorem 2. *For any skew-symmetric matrix $\mathbf{P}$ satisfying Assumption 1, any initialization $\hat{\pi}^{(1)}$, and any learning rate η such that $\eta L < \frac{1}{2}$, there exists a burn-in time*

$$T = O\left(\frac{D_{\text{KL}}(\pi^* \|\hat{\pi}^{(1)})^3 (\eta L)^2 |\mathbb{A}|^2}{(1 - 4\eta^2 L^2) C_{\mathbf{P}}^4 \varepsilon^6} \cdot \max\left\{1, \frac{C_{\mathbf{P}}^2 \varepsilon^4 |\mathbb{A}|}{(\eta L)^2}\right\}^3\right)$$

such that

$$D_{\text{KL}}(\pi^* \|\hat{\pi}^{(t)}) \leq \varepsilon \exp\left(-O\left(\frac{\eta^2 \varepsilon^3 C_{\mathbf{P}}^2 (t - T)}{(1 - 4\eta^2 L^2)}\right)\right)$$

for all $t \geq T$.

A complete proof is given in Appendix C. Here we highlight the key ideas behind the analysis.

Remark 3. Several observations follow from our theorem:

- Although the result is expressed in terms of $D_{\text{KL}}(\pi^* \|\hat{\pi}^{(t)})$, the guarantee directly implies convergence in terms of the duality gap via Pinsker’s inequality (Lemma 7):

$$D_{\text{KL}}(\pi^* \|\hat{\pi}^{(t)}) \geq \frac{1}{4L} \text{DualGap}(\hat{\pi}^{(t)})^2.$$

- The burn-in time T simplifies under mild assumptions. For uniform initialization, $D_{\text{KL}}(\pi^* \|\hat{\pi}^{(1)}) = O(\log |\mathbb{A}|)$. Moreover, since $\varepsilon < |\mathbb{A}|$ under Assumption 1 and $C_{\mathbf{P}} \leq 1$, the term $\frac{C_{\mathbf{P}}^2 \varepsilon^4 |\mathbb{A}|}{(\eta L)^2}$ is typically smaller than 1, except when η is chosen extremely small.

4.2 PROOF SKETCH

Non-increasing KL-projection. We begin by showing that the quantity $D_{\text{KL}}(\pi^* \|\hat{\pi}_t)$ decreases over time, up to a small perturbation. Specifically, define

$$\Theta_t = D_{\text{KL}}(\pi^* \|\hat{\pi}^{(t)}) + 4\eta^2 L^2 D_{\text{KL}}(\hat{\pi}^{(t)} \|\pi^{t-1}).$$

It can be shown that

$$\Theta_t - \Theta_{t+1} \geq (1 - 4\eta^2 L^2) \left(D_{\text{KL}}(\pi^{(t)} \|\hat{\pi}^{(t)}) + D_{\text{KL}}(\hat{\pi}^{(t+1)} \|\pi^{(t)}) \right).$$

This inequality, adapted from Lemma 10 of Wei et al. (2020), ensures that Θ_t strictly decreases whenever $\eta < 1/(2L)$. The main task is therefore to quantify how fast Θ_t converges to 0.

Figure 1 illustrates the trajectory of $\Theta_t - \Theta_{t+1}$ on a cyclic-preference matrix with initialization near $(1, 0, 0)$. The dynamics naturally split into two phases:

- a *burn-in stage*, with oscillatory behavior;
- a *convergence stage*, with nearly linear decay.

Burn-in Stage: Subgame Case. The subgame case arises when some action a satisfies both: (i) $\hat{\pi}_a^{(t)}$ (or $\hat{\pi}_a^{(t+1)}$) is large, and (ii) the update $|\eta(\mathbf{P}\pi^{(t)})_a|$ is also large. Here the policy shifts significantly. Define

$$K^{(t+1)} = \max_{a \in \mathbb{A}} \hat{\pi}_a^{(t+1)} |\eta(\mathbf{P}\pi^{(t)})_a|, \quad \hat{K}^{(t+1)} = \max_{a \in \mathbb{A}} \hat{\pi}_a^{(t)} |\eta(\mathbf{P}\pi^{(t)})_a|.$$

One can then show that there exists a constant $C' > 0$ such that

$$D_{\text{KL}}(\pi^{(t)} \|\hat{\pi}^{(t)}) + D_{\text{KL}}(\hat{\pi}^{(t+1)} \|\pi^{(t)}) \geq C' \max\{\hat{K}^{(t+1)}, K^{(t+1)}\}^2,$$

following the argument of Appendix D.3 in Wei et al. (2020).

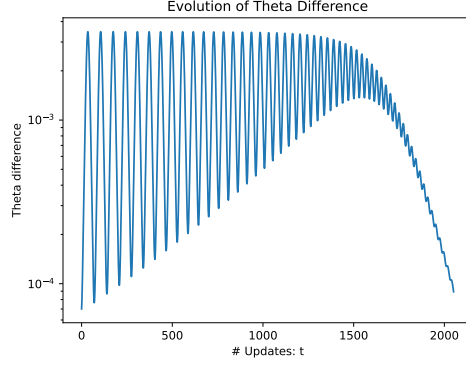


Figure 1: Evolution of $\Theta_t - \Theta_{t+1}$ for $\mathbf{P} = \begin{pmatrix} 0 & 1/2 & -1/2 \\ -1/2 & 0 & 1/2 \\ 1/2 & -1/2 & 0 \end{pmatrix}$ with initialization near $(1, 0, 0)$.

Burn-in Stage: Marginal Case. In contrast, the marginal case occurs when neither $K^{(t+1)}$ nor $\hat{K}^{(t+1)}$ is large, so updates in probability space are small. However, if $D_{\text{KL}}(\pi^* \parallel \hat{\pi}^{(t)})$ is still large, at least one coordinate of $\hat{\pi}^{(t)}$ must drift in its *logarithmic* value, ensuring continued progress. To capture this, we introduce the potential

$$\Phi_t = \frac{\langle \log \hat{\pi}^{(t)} - \log \pi^*, \eta \mathbf{P} \hat{\pi}^{(t)} \rangle}{(\Theta_t + 2e^{-1})^2}$$

and analyze $\Phi_t - \Phi_{t+1}$. The leading term reduces to

$$\langle \eta \mathbf{P} \pi^{(t)}, \eta \mathbf{P} \pi^{(t)} \rangle \geq \|\eta \mathbf{P} \pi^{(t)}\|_\infty^2,$$

using the update rule and the closeness of $\pi^{(t)}$ and $\hat{\pi}^{(t+1)}$ (guaranteed by small $\hat{K}^{(t+1)}$). After controlling approximation and residual terms, we obtain

$$\Phi_t - \Phi_{t+1} \geq \frac{\|\eta \mathbf{P} \pi^{(t)}\|_\infty^2}{(\Theta_t + 2e^{-1})^2} - O\left(\frac{\eta L |\mathbb{A}| (K^{(t)} + \hat{K}^{(t+1)} + K^{(t+1)})}{\varepsilon (\Theta_t + 2e^{-1})^2}\right) - (\text{minor terms}).$$

Remark 4. The dependence on ε arises naturally from Assumption 1. Intuitively, if $\hat{\pi}^{(t)}$ and $\pi^{(t)}$ concentrate near a non-equilibrium policy π' with restricted support, then escaping requires roughly

$$\min_{a: (\mathbf{P}\pi')_a > 0} \frac{-\log \hat{\pi}_a^{(t)}}{(\mathbf{P}\pi')_a}$$

steps. If $\pi_a^* = 0$, then $-\log \hat{\pi}_a^{(t)} \rightarrow \infty$, making such bounds impossible without the assumption $\varepsilon > 0$.

Burn-in Stage: Combined Analysis. To unify both subcases, we introduce an augmented expression

$$\bar{\Theta}_t = \Theta_{t-1} + c_1 \Theta_t + c_2 \Phi_t,$$

for constants $c_1, c_2 > 0$. One can show

$$\bar{\Theta}_t - \bar{\Theta}_{t+1} \geq O\left(\frac{(1 - 4\eta^2 L^2) C_{\mathbf{P}}^4 \varepsilon^6}{\eta^2 L^2 |\mathbb{A}|^2 (\Theta_t + 2e^{-1})^2}\right)$$

whenever $\Theta_t \geq \varepsilon$, while ensuring $\bar{\Theta}_t > \Theta_t$. Consequently, we obtain

$$\Theta_T < \varepsilon \quad \text{for} \quad T = O\left(\frac{\bar{\Theta}_1^3 (\eta L)^2 |\mathbb{A}|^2}{(1 - 4\eta^2 L^2) C_{\mathbf{P}}^4 \varepsilon^6}\right).$$

Convergence Stage. Once $\Theta_T < \varepsilon$, we can use the subgame bound together with lower bounds on $\min_a \hat{\pi}_a^{(t+1)}$ and $\max_a |(\mathbf{P}\boldsymbol{\pi}^{(t)})_a|$ to establish exponential convergence:

$$D_{\text{KL}}(\boldsymbol{\pi}^* \parallel \hat{\boldsymbol{\pi}}^{(t)}) \leq \varepsilon \exp(-O(\eta^2 \varepsilon^3 C_P^2 (t - T))).$$

Remark 5. Neither stage is analyzed tightly. In the burn-in phase, the lower bound on $\bar{\Theta}_t - \bar{\Theta}_{t+1}$ stems mainly from the transition region, which may not dominate in practice. In the convergence stage, our rate can be loose when $\arg \min_a \hat{\pi}_a^{(t+1)} \neq \arg \max_a |(\mathbf{P}\boldsymbol{\pi}^{(t)})_a|$. Improving both bounds is left for future work.

5 EXPERIMENTS

We present simulation results comparing the performance of OMWU against several existing NLHF algorithms (OMD, SPPO, MPO, ONPO, and EGPO).

5.1 GAME MATRICES

We construct $\mathcal{P}$ as follows: fix n (the matrix size) and an integer $0 \leq m < n$ which roughly denotes the rank of the NE polytope. Generate m random policies, and let $\mathbf{M}$ be the $n \times (n - m)$ matrix whose columns span their orthogonal complement. Then generate a random $(n - m) \times (n - m)$ skew-symmetric matrix $\mathbf{A}$ and set $\mathbf{P} = \mathbf{M}\mathbf{A}\mathbf{M}^\top$. Finally, normalize $\mathbf{P}$ so that $\mathbf{P} + \frac{1}{2}$ is nonnegative. Refer to Algorithm 1 for more details. We choose $n = 10$ for **tabular** and $n = 100$ for **neural** policy setting, respectively, and $0 \leq m < n/2$. Note that Assumption 1 can fail when $m = 0$.

5.2 IMPLEMENTATION

Using the notation in Zhou et al. (2025), OMWU update can be written as: $\boldsymbol{\theta}^{(-1/2)} = \boldsymbol{\theta}^{(0)} = \mathbf{0}$, for $t \geq 0$,

$$\begin{aligned}\boldsymbol{\theta}^{(t+1/2)} &= \boldsymbol{\theta}^{(t)} + \eta \mathbf{P} \boldsymbol{\pi}^{(t-1/2)}, \\ \boldsymbol{\theta}^{(t+1)} &= \boldsymbol{\theta}^{(t)} + \eta \mathbf{P} \boldsymbol{\pi}^{(t+1/2)}.\end{aligned}$$

Define a generalized IPO (Azar et al., 2023) loss using separate distributions for (y, y') and y'' (here we assume they are independent of θ , and simply choose $\beta = 1$):

$$\mathcal{L}_{\text{IPO}}(\boldsymbol{\theta}; \boldsymbol{\pi}_{\text{ref}}, \rho, \mu) = \mathbb{E}_{(y, y') \sim \rho} \left[\left(\log \frac{\boldsymbol{\pi}_{\boldsymbol{\theta}}(y) \boldsymbol{\pi}_{\text{ref}}(y')}{\boldsymbol{\pi}_{\boldsymbol{\theta}}(y') \boldsymbol{\pi}_{\text{ref}}(y)} - \mathbb{E}_{y'' \sim \mu} [\mathcal{P}(y > y'') - \mathcal{P}(y' > y'')] \right)^2 \right].$$

The result in Section 4.2 of Zhou et al. (2025) gives us the update:

$$\begin{aligned}\boldsymbol{\theta}^{(t+1/2)} &= \boldsymbol{\theta}^{(t)} - \underbrace{\frac{\eta_{\text{theory}} |\mathbb{A}|}{4}}_{=: \eta_{\text{optimizer}}} \nabla_{\boldsymbol{\theta}} \mathcal{L}_{\text{IPO}}(\boldsymbol{\theta}^{(t)}; \text{sg}[\boldsymbol{\pi}^{(t)}], \text{Uniform}, \boldsymbol{\pi}^{(t-1/2)}), \\ \boldsymbol{\theta}^{(t+1)} &= \boldsymbol{\theta}^{(t)} - \frac{\eta_{\text{theory}} |\mathbb{A}|}{4} \nabla_{\boldsymbol{\theta}} \mathcal{L}_{\text{IPO}}(\boldsymbol{\theta}^{(t)}; \text{sg}[\boldsymbol{\pi}^{(t)}], \text{Uniform}, \boldsymbol{\pi}^{(t+1/2)}),\end{aligned}$$

where Uniform is the uniform distribution over $\mathbb{A} \times \mathbb{A}$, and $\text{sg}[\cdot]$ means stopping-gradient.

The choice of codebase and implementation of the baselines are detailed in Appendix D.2. The hyperparameters when running the experiments are listed in Appendix D.3.

5.3 RESULTS

We select one matrix from the tabular and neural policy setting, respectively. They are shown in Figure 2, while the full experiment results are shown in Appendix D.4. In the title of each experiment, we report ε and $\lambda_{\min}$, the smallest positive singular value of $\mathcal{P}$, which is related to the constant C_P . Refer to Definition 1 for the definitions of ε and C_P .

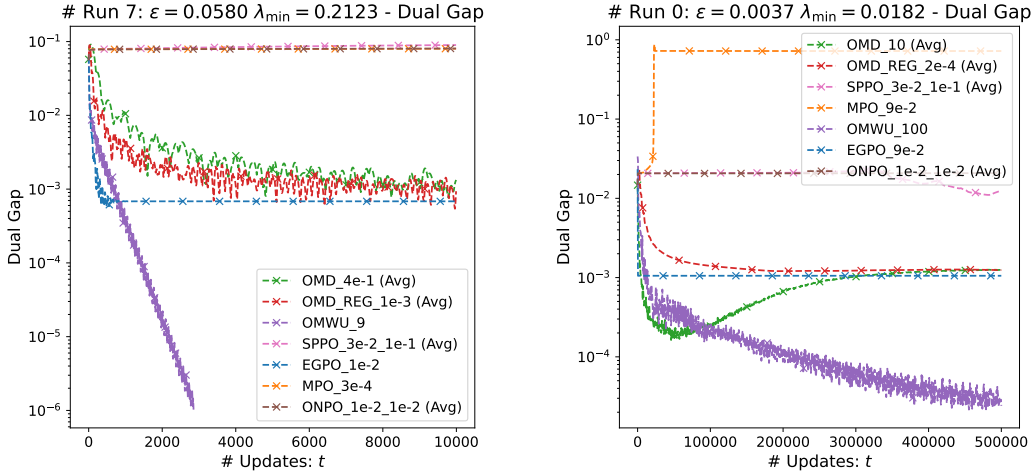


Figure 2: Selected results under the **tabular** (left) and **neural** (right) policy setting.

5.4 REMARKS

We make several remarks based on results in this section as well as those in Appendix D.4.

- In tabular setting, OMWU consistently achieves **linear** convergence in terms of the duality gap across all matrices, corroborating our main theorem. In neural setting, OMWU still outperforms other baselines. Meanwhile, regularized algorithms can only converge to the value related to the regularization coefficient β . Convergence of OMWU accelerates when $\lambda_{\min}$ is large.
- Figure 3 illustrates the stark contrast between the **last-iterate** and **average-iterate** behavior of OMD (regularized). This result highlights the advantage of algorithms with last-iterate convergence guarantees.
- In our constructed game matrices, **all** the algorithms requiring **nested-optimization** (SPPO, MPO, ONPO) fail to converge even with extensive hyperparameter search. This might be a direct result of insufficient inner optimization steps (we choose 10 steps). However, more inner optimization steps are not acceptable, as currently these algorithms already consume 3 to 10 more running time compared to OMD, EGPO, and OMWU.
- The duality gap of OMWU (and occasionally of other algorithms) exhibits noise. This highlights the necessity of analyzing $D_{\text{KL}}(\pi^* \parallel \hat{\pi}^{(t)})$ rather than relying solely on the duality gap.

6 CONCLUSION

We have analyzed the OMWU algorithm in the context of NLHF. Compared with Wei et al. (2020), our results relax the uniqueness assumption, improve the convergence rate and burn-in characterization, and eliminate the exponential dependence on $D_{\text{KL}}(\pi^* \parallel \hat{\pi}_1)/\epsilon$.

This work provides a theoretical guarantee for non-regularized preference-based learning, highlighting the potential of OMWU as an alternative to regularizer-dependent methods. However, we are currently not able to reproduce our result on a fine-tuning problem of large language models due to resource constraints.

However, one limitation of our work is that the result still does not hold for any preference matrix. In addition, the burn-in time and convergence rate are not proved to be of tightest order given the related terms (and we believe they are not tight). We leave the generalization and order problem to future research.

REFERENCES

- Ioannis Anagnostides, Constantinos Daskalakis, Gabriele Farina, Maxwell Fishelson, Noah Golowich, and Tuomas Sandholm. Near-optimal no-regret learning for correlated equilibria in multi-player general-sum games. In *Proceedings of the 54th Annual ACM SIGACT Symposium on Theory of Computing*, pp. 736–749, 2022.
- Mohammad Gheshlaghi Azar, Mark Rowland, Bilal Piot, Daniel Guo, Daniele Calandriello, Michal Valko, and Rémi Munos. A general theoretical paradigm to understand learning from human preferences. *ArXiv*, abs/2310.12036, 2023.
- Ralph Allan Bradley and Milton E Terry. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952.
- Yang Cai, Gabriele Farina, Julien Grand-Clément, Christian Kroer, Chung-Wei Lee, Haipeng Luo, and Weiqiang Zheng. Fast last-iterate convergence of learning in games requires forgetful algorithms. *Advances in Neural Information Processing Systems*, 37:23406–23434, 2024.
- Daniele Calandriello, Daniel Guo, Remi Munos, Mark Rowland, Yunhao Tang, Bernardo Avila Pires, Pierre Harvey Richemond, Charline Le Lan, Michal Valko, Tianqi Liu, et al. Human alignment of large language models through online preference optimisation. *arXiv preprint arXiv:2403.08635*, 2024.
- Nicolo Cesa-Bianchi and Gábor Lugosi. *Prediction, learning, and games*. Cambridge university press, 2006.
- Constantinos Daskalakis and Ioannis Panageas. Last-iterate convergence: Zero-sum games and constrained min-max optimization. *arXiv preprint arXiv:1807.04252*, 2018.
- Constantinos Daskalakis, Maxwell Fishelson, and Noah Golowich. Near-optimal no-regret learning in general games. *Advances in Neural Information Processing Systems*, 34:27604–27616, 2021.
- Xavier Glorot and Yoshua Bengio. Understanding the difficulty of training deep feedforward neural networks. In *Proceedings of the thirteenth international conference on artificial intelligence and statistics*, pp. 249–256. JMLR Workshop and Conference Proceedings, 2010.
- Tor Lattimore and Csaba Szepesvári. *Bandit algorithms*. Cambridge University Press, 2020.
- Chung-Wei Lee, Christian Kroer, and Haipeng Luo. Last-iterate convergence in extensive-form games. *Advances in Neural Information Processing Systems*, 34:14293–14305, 2021.
- Kenneth O May. Intransitivity, utility, and the aggregation of preference patterns. *Econometrica: Journal of the Econometric Society*, pp. 1–13, 1954.
- Rémi Munos, Michal Valko, Daniele Calandriello, Mohammad Gheshlaghi Azar, Mark Rowland, Zhaohan Daniel Guo, Yunhao Tang, Matthieu Geist, Thomas Mesnard, Andrea Michi, et al. Nash learning from human feedback. *arXiv preprint arXiv:2312.00886*, 18, 2023.
- 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.
- Corby Rosset, Ching-An Cheng, Arindam Mitra, Michael Santacrose, Ahmed Awadallah, and Tengyang Xie. Direct nash optimization: Teaching language models to self-improve with general preferences. *arXiv preprint arXiv:2404.03715*, 2024.
- Lior Shani, Aviv Rosenberg, Asaf Cassel, Oran Lang, Daniele Calandriello, Avital Zipori, Hila Noga, Orgad Keller, Bilal Piot, Idan Szpektor, et al. Multi-turn reinforcement learning from preference human feedback. *arXiv preprint arXiv:2405.14655*, 2024.
- Gokul Swamy, Christoph Dann, Rahul Kidambi, Zhiwei Steven Wu, and Alekh Agarwal. A minimaximalist approach to reinforcement learning from human feedback. *arXiv preprint arXiv:2401.04056*, 2024.

- J v. Neumann. Zur theorie der gesellschaftsspiele. *Mathematische annalen*, 100(1):295–320, 1928.
- Mingzhi Wang, Chengdong Ma, Qizhi Chen, Linjian Meng, Yang Han, Jiancong Xiao, Zhaowei Zhang, Jing Huo, Weijie J Su, and Yaodong Yang. Magnetic preference optimization: Achieving last-iterate convergence for language model alignment. *arXiv preprint arXiv:2410.16714*, 2024.
- Chen-Yu Wei, Chung-Wei Lee, Mengxiao Zhang, and Haipeng Luo. Linear last-iterate convergence in constrained saddle-point optimization. *arXiv preprint arXiv:2006.09517*, 2020.
- Yue Wu, Zhiqing Sun, Huizhuo Yuan, Kaixuan Ji, Yiming Yang, and Quanquan Gu. Self-play preference optimization for language model alignment. *arXiv preprint arXiv:2405.00675*, 2024.
- Chenlu Ye, Wei Xiong, Yuheng Zhang, Hanze Dong, Nan Jiang, and Tong Zhang. Online iterative reinforcement learning from human feedback with general preference model. *Advances in Neural Information Processing Systems*, 37:81773–81807, 2024.
- Yuheng Zhang, Dian Yu, Baolin Peng, Linfeng Song, Ye Tian, Mingyue Huo, Nan Jiang, Haitao Mi, and Dong Yu. Iterative nash policy optimization: Aligning llms with general preferences via no-regret learning. *arXiv preprint arXiv:2407.00617*, 2024.
- Yuheng Zhang, Dian Yu, Tao Ge, Linfeng Song, Zhichen Zeng, Haitao Mi, Nan Jiang, and Dong Yu. Improving llm general preference alignment via optimistic online mirror descent. *arXiv preprint arXiv:2502.16852*, 2025.
- Runlong Zhou, Maryam Fazel, and Simon Shaolei Du. Extragradients preference optimization (EGPO): Beyond last-iterate convergence for nash learning from human feedback. In *Second Conference on Language Modeling*, 2025. URL <https://openreview.net/forum?id=EP7mAqx2BO>.

Appendix

USE STATEMENT OF LARGE LANGUAGE MODELS

Large language models have been used to polish writing, including correcting grammatical errors and making content presentation more cohesive.

A PROOFS OF LEMMAS ON KL PROJECTION

This appendix establishes several basic properties of the KL projection $p(\pi)$.

Lemma 3. *Under Assumption 1, we have $\pi \in \mathbb{M}$ if and only if $P\pi = 0$.*

Proof. We only prove the “only if” direction, as the reverse is immediate.

By Assumption 1, there exists some $\pi' \in \mathbb{M}$ such that $\pi'_a > 0$ for all $a \in \mathbb{A}$. Since both π and π' are equilibria, we have

$$\pi'^\top P\pi = 0.$$

Expanding, we obtain

$$\sum_{a \in \mathbb{A}} \pi'_a (P\pi)_a = 0.$$

Because $\pi \in \mathbb{M}$, each $(P\pi)_a \leq 0$. Moreover, $\pi'_a > 0$ for all $a \in \mathbb{A}$ by our choice of π' . These two facts force $(P\pi)_a = 0$ for all $a \in \mathbb{A}$, i.e., $P\pi = 0$. $\square$

Lemma 4. *Under Assumption 1, for any policy π with full support, there exists a unique equilibrium $\pi' \in \mathbb{M}$ such that*

$$\pi' = \arg \min_{\pi'' \in \mathbb{M}} D_{\text{KL}}(\pi'' \| \pi).$$

Moreover, the minimizer satisfies $\pi'_a > 0$ for all $a \in \mathbb{A}$.

Proof. By Lemma 3, the set $\mathbb{M}$ is compact. Since $\pi' \mapsto D_{\text{KL}}(\pi' \| \pi)$ is continuous on $\mathbb{M}$, existence follows.

To show positivity of all coordinates, note that there exists some $\pi'' \in \mathbb{M}$ with $\pi''_a > 0$ for every $a \in \mathbb{A}$. Consider

$$f(\lambda) = D_{\text{KL}}(\lambda\pi'' + (1-\lambda)\pi' \| \pi), \quad 0 \leq \lambda \leq 1.$$

Since $\mathbb{M}$ is convex, $\lambda\pi'' + (1-\lambda)\pi' \in \mathbb{M}$. Minimality of π' implies that $f(\lambda)$ attains its minimum at $\lambda = 0$. Differentiating,

$$f'(\lambda) = \sum_{a \in \mathbb{A}} (\pi''_a - \pi'_a) \log \frac{\lambda\pi''_a + (1-\lambda)\pi'_a}{\pi_a}.$$

If $\pi'_a > 0$, the limit as $\lambda \rightarrow 0^+$ is finite. If $\pi'_a = 0$, then $\pi''_a - \pi'_a > 0$, and the limit becomes $-\infty$, contradicting minimality. Hence $\pi'_a > 0$ for all a .

For uniqueness, suppose π' and π'' are two minimizers with $D_{\text{KL}}(\pi' \| \pi) = D_{\text{KL}}(\pi'' \| \pi)$. Define

$$f(\lambda) = D_{\text{KL}}(\lambda\pi' + (1-\lambda)\pi'' \| \pi).$$

Differentiating,

$$f'(\lambda) = \sum_{a \in \mathbb{A}} (\pi'_a - \pi''_a) \log \frac{\lambda\pi'_a + (1-\lambda)\pi''_a}{\pi_a},$$

$$f''(\lambda) = \sum_{a \in \mathbb{A}} \frac{(\pi'_a - \pi''_a)^2}{\lambda\pi'_a + (1-\lambda)\pi''_a} \geq 0.$$

Thus f is convex. Since $f(0) = f(1)$ achieves the minimum, f must be constant, forcing $f''(\lambda) \equiv 0$. Therefore $\pi' = \pi''$. $\square$

Lemma 5. *Under Assumption 1, we have*

$$p(\hat{\pi}^{(1)}) = p(\hat{\pi}^{(2)}) = \dots = p(\hat{\pi}^{(t)}) = p(\hat{\pi}^{(t+1)}) = \dots$$

Proof. From the update rule in Equation (2), there exists a constant $\hat{M}^{(t+1)}$ such that

$$\log \hat{\pi}^{(t+1)} = \log \hat{\pi}^{(t)} + \eta P\pi^{(t)} + \hat{M}^{(t+1)}, \quad t \geq 1.$$

Let $\pi^* \in \mathbb{M}$ be any Nash equilibrium. Then

$$\begin{aligned} D_{\text{KL}}(\pi^* \parallel \hat{\pi}^{(t+1)}) &= \langle \pi^*, \log \pi^* \rangle - \langle \pi^*, \log \hat{\pi}^{(t+1)} \rangle \\ &= D_{\text{KL}}(\pi^* \parallel \hat{\pi}^{(t)}) - \eta \langle \pi^*, P\pi^{(t)} \rangle - \langle \pi^*, \hat{M}^{(t+1)} \rangle. \end{aligned}$$

Since $\langle \pi^*, P\pi^{(t)} \rangle = -\langle P\pi^*, \pi^{(t)} \rangle = 0$, the difference reduces to $-\hat{M}^{(t+1)}$, a normalization term. Thus the projection target does not change with t , i.e.,

$$p(\hat{\pi}^{(t+1)}) = p(\hat{\pi}^{(t)}).$$

□

B PROOF OF MATRIX-DEPENDENT CONSTANT

This section establishes that C_P is well defined and strictly positive.

Lemma 6. *Let $\mathbb{N} \subseteq \mathbb{R}^n$ be a subspace. Then there exists a constant $C(\mathbb{N}) < 1$ such that the following holds: if $\mathbf{a}, \mathbf{b} \in \mathbb{R}^n$ and $\mathbf{r} \in \mathbb{N}$ satisfy $\text{sgn}(\mathbf{a}) = \text{sgn}(\mathbf{b})$ and $\mathbf{b} \in \mathbb{N}^\perp$, then*

$$|\langle \mathbf{r}, \mathbf{a} \rangle| \leq C(\mathbb{N}) \|\mathbf{r}\|_2 \|\mathbf{a}\|_2.$$

Proof. Since $\text{sgn}(\mathbf{a}) \in \{-1, 0, 1\}^n$ has only finitely many possibilities, it suffices to fix any nonzero sign pattern $\mathbf{s} \in \{-1, 0, 1\}^n$ and show the claim holds with some $C(\mathbb{N}, \mathbf{s}) < 1$ under $\text{sgn}(\mathbf{a}) = \text{sgn}(\mathbf{b}) = \mathbf{s}$. By scaling, we may assume $\|\mathbf{r}\|_2 = \|\mathbf{a}\|_2 = 1$.

If no $\mathbf{b} \in \mathbb{N}^\perp$ has sign pattern $\mathbf{s}$, then the claim is vacuous and we may take $C(\mathbb{N}, \mathbf{s}) = 0$. Otherwise, fix such a vector $\mathbf{b} \in \mathbb{N}^\perp$ with $\text{sgn}(\mathbf{b}) = \mathbf{s}$.

Define

$$\mathbb{S} = \left\{ \mathbf{a} \in \mathbb{R}^n : \|\mathbf{a}\|_2 = 1, a_i \geq 0 \text{ if } s_i = 1, a_i = 0 \text{ if } s_i = 0, a_i \leq 0 \text{ if } s_i = -1 \right\}.$$

By construction, $\langle \mathbf{a}, \mathbf{b} \rangle \geq 0$ for all $\mathbf{a} \in \mathbb{S}$. Moreover, equality cannot hold because $\mathbf{a} \neq 0$ and every nonzero coordinate of $\mathbf{a}$ agrees in sign with $\mathbf{b}$. Thus $\mathbb{S}$ does not intersect $\mathbb{N}$ (since $\mathbf{b} \in \mathbb{N}^\perp$).

It follows that for every $\mathbf{r} \in \mathbb{N}$ and $\mathbf{a} \in \mathbb{S}$ with $\|\mathbf{r}\|_2 = 1$, we must have $|\langle \mathbf{r}, \mathbf{a} \rangle| < 1$. Compactness of $\mathbb{S}$ then ensures

$$C(\mathbb{N}, \mathbf{s}) = \max_{\substack{\mathbf{r} \in \mathbb{N}, \|\mathbf{r}\|_2=1 \\ \mathbf{a} \in \mathbb{S}}} |\langle \mathbf{r}, \mathbf{a} \rangle| < 1.$$

Taking $C(\mathbb{N}) = \max_{\mathbf{s}} C(\mathbb{N}, \mathbf{s})$ proves the claim. □

We now show that

$$C_P = \min_{\pi} \frac{\|P\pi\|_\infty}{\|\pi - p(\pi)\|_1}$$

is well defined and positive under Assumption 1.

—

Let

$$\mathbb{N} = \left\{ \mathbf{r} \in \mathbb{R}^{|\mathbb{A}|} : P\mathbf{r} = \mathbf{0}, \sum_{a \in \mathbb{A}} r_a = 0 \right\},$$

and let π' denote the orthogonal projection of π onto $p(\pi) + \mathbb{N}$, i.e.,

$$\pi' = \arg \min_{\pi' \in p(\pi) + \mathbb{N}} \|\pi - \pi'\|_2,$$

where π' need not be a distribution. Since $p(\pi)_a > 0$ for all $a \in \mathbb{A}$ by Lemma 1, for any $\mathbf{r} \in \mathbb{N}$ the perturbed strategy $p(\pi) + \lambda \mathbf{r}$ is also a Nash equilibrium for sufficiently small $|\lambda|$ by Lemma 3. By first-order optimality of $p(\pi)$, we obtain

$$\langle \mathbf{r}, \log \pi - \log p(\pi) \rangle = 0, \quad \forall \mathbf{r} \in \mathbb{N},$$

which implies

$$\log \pi - \log p(\pi) \in \mathbb{N}^\perp.$$

Since

$$\text{sgn}(\log \pi - \log p(\pi)) = \text{sgn}(\pi - p(\pi)),$$

Lemma 6 guarantees the existence of $C(\mathbb{N}) < 1$ such that

$$|\langle \pi' - p(\pi), \pi - p(\pi) \rangle| \leq C(\mathbb{N}) \|\pi' - p(\pi)\|_2 \|\pi - p(\pi)\|_2.$$

Since π' is the projection of π onto $p(\pi) + \mathbb{N}$, we also have

$$\langle \pi' - p(\pi), \pi - \pi' \rangle = 0.$$

Combining these facts yields

$$\|\pi' - p(\pi)\|_2 \leq C(\mathbb{N}) \|\pi - p(\pi)\|_2,$$

and by the Pythagorean theorem,

$$\|\pi - \pi'\|_2 = \sqrt{\|\pi - p(\pi)\|_2^2 - \|\pi' - p(\pi)\|_2^2} \geq \sqrt{1 - C(\mathbb{N})^2} \|\pi - p(\pi)\|_2.$$

Next, we relate $\|\pi - \pi'\|_2$ to $\|P\pi\|_\infty$. Note that $\pi - \pi' \in \mathbb{N}' \cap \mathbb{N}^\perp$, where $\mathbb{N}' = \{\mathbf{r} : \sum_{a \in \mathbb{A}} r_a = 0\}$. Since $\mathbb{N}$ is the null space of $P|_{\mathbb{N}'}$, elementary linear algebra gives

$$\|P(\pi - \pi')\|_2 \geq \lambda_{\min} \|\pi - \pi'\|_2,$$

where $\lambda_{\min}$ is the smallest positive singular value of $P|_{\mathbb{N}'}$.

Finally, combining all inequalities:

$$\begin{aligned} \|P\pi\|_\infty &\geq \sqrt{|\mathbb{A}|} \|P\pi\|_2 \\ &\geq \lambda_{\min} \sqrt{|\mathbb{A}|} \|\pi - \pi'\|_2 \\ &\geq \lambda_{\min} \sqrt{|\mathbb{A}|(1 - C(\mathbb{N})^2)} \|\pi - p(\pi)\|_2 \\ &\geq \lambda_{\min} |\mathbb{A}| \sqrt{1 - C(\mathbb{N})^2} \|\pi - p(\pi)\|_1. \end{aligned}$$

Thus we may take

$$C_P = \lambda_{\min} |\mathbb{A}| \sqrt{1 - C(\mathbb{N})^2} > 0.$$

C PROOF OF THE MAIN THEOREM

C.1 BASIC LEMMAS

Lemma 7 (Pinsker's Inequality). *If π and π' are probability distributions, then*

$$D_{\text{KL}}(\pi \| \pi') \geq \frac{1}{2} \|\pi - \pi'\|_1^2.$$

Lemma 8. *Let $\{x_n\}_{n=0}^t$ be a decreasing sequence. Suppose there exists a nonnegative, non-increasing, continuously differentiable function f on $[x_t, x_0]$ such that*

$$x_n - x_{n+1} \geq f(x_n)$$

for all $n = 0, 1, \dots, t-1$. Then

$$\int_{x_t}^{x_0} \frac{1}{f(x)} dx \geq t.$$

Proof. Since f is non-increasing, we have

$$1 \leq \frac{x_n - x_{n+1}}{f(x_n)} = \int_{x_{n+1}}^{x_n} \frac{dx}{f(x_n)} \leq \int_{x_{n+1}}^{x_n} \frac{dx}{f(x)}$$

for all $n = 0, 1, \dots, t-1$. Summing over n gives

$$t \leq \int_{x_t}^{x_0} \frac{dx}{f(x)}.$$

□

Lemma 9. Let $\{x_n\}_{n=0}^t$ be a decreasing sequence. Suppose there exists a nonnegative, non-decreasing, continuously differentiable function f on $[x_t, x_0]$ such that

$$x_n - x_{n+1} \geq f(x_{n+1})$$

for all $n = 0, 1, \dots, t-1$. Then

$$\ln \frac{f(x_0)}{f(x_t)} + \int_{x_t}^{x_0} \frac{1}{f(x)} dx \geq t.$$

Proof. Let $F(x) = x + f(x)$. Then $F'(x) = 1 + f'(x) > 0$, so F is strictly increasing and has an inverse F^{-1} on $[F(x_t), F(x_0)]$. Denote $y_n = F(x_n)$. Then

$$y_{n+1} \leq y_n - f(x_n) = y_n - f(F^{-1}(y_n))$$

for $n = 0, 1, \dots, t-1$. Applying Lemma 8 to the sequence $\{y_n\}$ yields

$$t \leq \int_{F(x_t)}^{F(x_0)} \frac{dy}{f(F^{-1}(y))}.$$

Now substitute $y = F(x) = x + f(x)$. Then $dy = (1 + f'(x))dx$, so

$$\int_{F(x_t)}^{F(x_0)} \frac{dy}{f(F^{-1}(y))} = \int_{x_t}^{x_0} \frac{1 + f'(x)}{f(x)} dx = \int_{x_t}^{x_0} \frac{1}{f(x)} dx + \int_{f(x_t)}^{f(x_0)} \frac{dt}{t}.$$

The second term equals $\ln \frac{f(x_0)}{f(x_t)}$, giving the claim. □

C.2 MONOTONICITY OF Θ_t

Define

$$\Theta_t = D_{\text{KL}}(\pi^* \|\hat{\pi}^{(t)}) + 4\eta^2 L^2 D_{\text{KL}}(\hat{\pi}^{(t)} \|\pi^{(t-1)}),$$

where $L = \max_{a,b \in \mathbb{A}} P_{a,b}$.

Lemma 10.

$$\Theta_t - \Theta_{t+1} \geq (1 - 4\eta^2 L^2) \left(D_{\text{KL}}(\pi^{(t)} \|\hat{\pi}^{(t)}) + D_{\text{KL}}(\hat{\pi}^{(t+1)} \|\pi^{(t)}) \right).$$

In particular, if $\eta L < \frac{1}{2}$, then Θ_t is strictly decreasing.

Proof. From the update rule Equation (2),

$$\langle \pi^* - \pi^{(t)}, \log \hat{\pi}^{(t+1)} \rangle = \langle \pi^* - \pi^{(t)}, \log \hat{\pi}^{(t)} + \eta P \pi^{(t)} \rangle.$$

Since π^* is an equilibrium and P is skew-symmetric,

$$\langle \pi^*, P \pi^{(t)} \rangle = -\langle P \pi^*, \pi^{(t)} \rangle \geq 0, \quad \langle \pi^{(t)}, P \pi^{(t)} \rangle = 0,$$

hence

$$\langle \pi^* - \pi^{(t)}, \log \hat{\pi}^{(t+1)} \rangle = \langle \pi^* - \pi^{(t)}, \log \hat{\pi}^{(t)} \rangle.$$

Therefore,

$$D_{\text{KL}}(\pi^* \|\hat{\pi}^{(t)}) - D_{\text{KL}}(\pi^* \|\hat{\pi}^{(t+1)}) = \langle \pi^*, \log \hat{\pi}^{(t+1)} \rangle - \langle \pi^*, \log \hat{\pi}^{(t)} \rangle$$

$$\geq \langle \pi^{(t)}, \log \hat{\pi}^{(t+1)} - \log \hat{\pi}^{(t)} \rangle.$$

Subtracting $D_{\text{KL}}(\hat{\pi}^{(t+1)} \parallel \pi^{(t)})$ and rearranging gives

$$\begin{aligned} D_{\text{KL}}(\pi^* \parallel \hat{\pi}^{(t)}) - D_{\text{KL}}(\pi^* \parallel \hat{\pi}^{(t+1)}) - D_{\text{KL}}(\hat{\pi}^{(t+1)} \parallel \pi^{(t)}) \\ \geq \langle \pi^{(t)} - \hat{\pi}^{(t+1)}, \log \hat{\pi}^{(t+1)} - \log \pi^{(t)} \rangle + D_{\text{KL}}(\pi^{(t)} \parallel \hat{\pi}^{(t)}). \end{aligned}$$

Meanwhile, by Lemma 7,

$$\begin{aligned} \|\pi^{(t)} - \hat{\pi}^{(t+1)}\|_1^2 &\leq D_{\text{KL}}(\pi^{(t)} \parallel \hat{\pi}^{(t+1)}) + D_{\text{KL}}(\hat{\pi}^{(t+1)} \parallel \pi^{(t)}) \\ &= \langle \pi^{(t)} - \hat{\pi}^{(t+1)}, \log \pi^{(t)} - \log \hat{\pi}^{(t+1)} \rangle \\ &= \langle \pi^{(t)} - \hat{\pi}^{(t+1)}, \eta P(\pi^{(t-1)} - \pi^{(t)}) \rangle \\ &\leq \eta \|\pi^{(t)} - \hat{\pi}^{(t+1)}\|_1 \|P(\pi^{(t)} - \pi^{(t-1)})\|_\infty \\ &\leq \eta L \|\pi^{(t)} - \hat{\pi}^{(t+1)}\|_1 \|\pi^{(t)} - \pi^{(t-1)}\|_1. \end{aligned}$$

Thus,

$$\|\pi^{(t)} - \hat{\pi}^{(t+1)}\|_1 \leq \eta L \|\pi^{(t)} - \pi^{(t-1)}\|_1.$$

Plugging back,

$$\begin{aligned} |\langle \pi^{(t)} - \hat{\pi}^{(t+1)}, \log \hat{\pi}^{(t+1)} - \log \pi^{(t)} \rangle| &\leq \eta^2 L^2 \|\pi^{(t)} - \pi^{(t-1)}\|_1^2 \\ &\leq 2\eta^2 L^2 (\|\pi^{(t)} - \hat{\pi}^{(t)}\|_1^2 + \|\hat{\pi}^{(t)} - \pi^{(t-1)}\|_1^2) \\ &\leq 4\eta^2 L^2 (D_{\text{KL}}(\pi^{(t)} \parallel \hat{\pi}^{(t)}) + D_{\text{KL}}(\hat{\pi}^{(t)} \parallel \pi^{(t-1)})). \end{aligned}$$

Combining with the earlier inequality proves the claim. $\square$

C.3 OTHER IMPORTANT INEQUALITIES

Let

$$M^{(t)} = \log \langle \hat{\pi}^{(t)}, \exp(\eta P \pi^{(t-1)}) \rangle, \quad \hat{M}^{(t+1)} = \log \langle \hat{\pi}^{(t)}, \exp(\eta P \pi^{(t)}) \rangle$$

be the normalizing constants chosen so that

$$\log \pi^{(t)} = \log \hat{\pi}^{(t)} + \eta P \pi^{(t-1)} - M^{(t)}, \quad \log \hat{\pi}^{(t+1)} = \log \hat{\pi}^{(t)} + \eta P \pi^{(t)} - \hat{M}^{(t+1)}.$$

Also, denote

$$K^{(t)} = \max_{a \in \mathbb{A}} \hat{\pi}_a^{(t)} |\eta P \pi^{(t-1)}|_a, \quad \hat{K}^{(t+1)} = \max_{a \in \mathbb{A}} \hat{\pi}_a^{(t)} |\eta P \pi^{(t)}|_a.$$

Lemma 11.

$$\|\log \hat{\pi}^{(t+1)} - \log \hat{\pi}^{(t)}\|_\infty \leq 2\eta L.$$

Proof. Recall the update rule Equation (2):

$$\log \hat{\pi}^{(t+1)} = \log \hat{\pi}^{(t)} + \eta P \pi^{(t)} - \hat{M}^{(t+1)}. \quad (3)$$

From the definition of $\hat{M}^{(t+1)}$, we have $|\hat{M}^{(t+1)}| \leq \|P \pi^{(t)}\|_\infty$. Hence

$$\|\log \hat{\pi}^{(t+1)} - \log \hat{\pi}^{(t)}\|_\infty \leq 2\eta \|P \pi^{(t)}\|_\infty \leq 2\eta L.$$

$\square$

Lemma 12.

$$|\hat{M}^{(t+1)}| \leq e^{\eta L} |\mathbb{A}| \hat{K}^{(t+1)}, \quad |M^{(t)}| \leq e^{\eta L} |\mathbb{A}| K^{(t)}.$$

Proof. From the convexity of the exponential function, we have

$$1 + x \leq e^x \leq 1 + \frac{e^{\eta L} - 1}{\eta L} x \quad \text{for all } x \in [0, \eta L].$$

Since $\eta(\mathbf{P}\boldsymbol{\pi}^{(t)})_a \in [-\eta L, \eta L]$, it follows that

$$1 + \eta(\mathbf{P}\boldsymbol{\pi}^{(t)})_a \leq \exp(\eta(\mathbf{P}\boldsymbol{\pi}^{(t)})_a) \leq 1 + \frac{e^{\eta L} - 1}{\eta L} \max\{\eta(\mathbf{P}\boldsymbol{\pi}^{(t)})_a, 0\}.$$

Recall that

$$\exp(\hat{M}^{(t+1)}) = \sum_{a \in \mathbb{A}} \hat{\pi}_a^{(t)} \exp(\eta(\mathbf{P}\boldsymbol{\pi}^{(t)})_a).$$

Hence

$$1 + |\mathbb{A}| \eta \min_{a \in \mathbb{A}} \hat{\pi}_a^{(t)} (\mathbf{P}\boldsymbol{\pi}^{(t)})_a \leq e^{\hat{M}^{(t+1)}} \leq 1 + \frac{e^{\eta L} - 1}{\eta L} |\mathbb{A}| \max\{\eta \max_{a \in \mathbb{A}} \hat{\pi}_a^{(t)} (\mathbf{P}\boldsymbol{\pi}^{(t)})_a, 0\}.$$

Thus,

$$\hat{K}^{(t+1)} \geq \frac{\hat{M}^{(t+1)}}{e^{\eta L} |\mathbb{A}|}.$$

The other inequality follows analogously. $\square$

Lemma 13.

$$\|\hat{\boldsymbol{\pi}}^{(t+1)} - \hat{\boldsymbol{\pi}}^{(t)}\|_1 \leq \frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)|\mathbb{A}| \hat{K}^{(t+1)}}{2\eta L}.$$

Also,

$$\|\boldsymbol{\pi}^{(t)} - \hat{\boldsymbol{\pi}}^{(t)}\|_1 \leq \frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)|\mathbb{A}| K^{(t)}}{2\eta L}.$$

Proof. From the update rule Equation (2),

$$\hat{\pi}_a^{(t+1)} - \hat{\pi}_a^{(t)} = \hat{\pi}_a^{(t)} \exp((\eta \mathbf{P}\boldsymbol{\pi}^{(t)})_a - \hat{M}^{(t+1)}) - \hat{\pi}_a^{(t)}.$$

When $(\eta \mathbf{P}\boldsymbol{\pi}^{(t)})_a - \hat{M}^{(t+1)} \geq 0$, convexity and the trivial bound $|(\eta \mathbf{P}\boldsymbol{\pi}^{(t)})_a - \hat{M}^{(t+1)}| \leq 2\eta L$ give

$$\exp((\eta \mathbf{P}\boldsymbol{\pi}^{(t)})_a - \hat{M}^{(t+1)}) \leq 1 + \frac{e^{2\eta L} - 1}{2\eta L} ((\eta \mathbf{P}\boldsymbol{\pi}^{(t)})_a - \hat{M}^{(t+1)}).$$

Hence,

$$0 \leq \hat{\pi}_a^{(t+1)} - \hat{\pi}_a^{(t)} \leq \frac{(e^{2\eta L} - 1)(\hat{\pi}_a^{(t)} |\hat{M}^{(t+1)}| + \hat{K}^{(t+1)})}{2\eta L}.$$

When $(\eta \mathbf{P}\boldsymbol{\pi}^{(t)})_a - \hat{M}^{(t+1)} \leq 0$, we use

$$\exp((\eta \mathbf{P}\boldsymbol{\pi}^{(t)})_a - \hat{M}^{(t+1)}) \geq 1 + (\eta \mathbf{P}\boldsymbol{\pi}^{(t)})_a - \hat{M}^{(t+1)},$$

which implies

$$0 \geq \hat{\pi}_a^{(t+1)} - \hat{\pi}_a^{(t)} \geq -\hat{\pi}_a^{(t)} |\hat{M}^{(t+1)}| - \hat{K}^{(t+1)}.$$

Combining the two cases and noting that $e^{2\eta L} \geq 1 + 2\eta L$, we obtain

$$\|\hat{\boldsymbol{\pi}}^{(t+1)} - \hat{\boldsymbol{\pi}}^{(t)}\|_1 \leq \frac{e^{2\eta L} - 1}{2\eta L} (|\hat{M}^{(t+1)}| + |\mathbb{A}| \hat{K}^{(t+1)}).$$

Applying Lemma 12 yields the claimed bound. The second inequality follows by the same reasoning. $\square$

Recall that $\varepsilon = \min_{a \in \mathbb{A}} \pi_a^*$.

Lemma 14. If $\pi_a > 0$ for all $a \in \mathbb{A}$, then

$$\varepsilon \|\log \boldsymbol{\pi}^* - \log \boldsymbol{\pi}\|_1 \leq D_{\text{KL}}(\boldsymbol{\pi}^* \|\boldsymbol{\pi}) + 2e^{-1}.$$

Proof. We first note

$$\varepsilon \|\log \pi^* - \log \pi\|_1 \leq \|\pi^*(\log \pi^* - \log \pi)\|_1.$$

Thus it suffices to prove

$$\|\pi^*(\log \pi^* - \log \pi)\|_1 \leq D_{\text{KL}}(\pi^* \|\pi) + 2e^{-1}.$$

Rearranging, this is equivalent to showing

$$\sum_{a \in \mathbb{A}: \pi_a^* < \pi_a} \pi_a^*(\log \pi_a - \log \pi_a^*) \leq e^{-1}.$$

Define

$$S = \sum_{a \in \mathbb{A}: \pi_a^* < \pi_a} \pi_a, \quad S^* = \sum_{a \in \mathbb{A}: \pi_a^* < \pi_a} \pi_a^*.$$

By Jensen's inequality,

$$\sum_{a: \pi_a^* < \pi_a} \pi_a^*(\log \pi_a - \log \pi_a^*) \leq S^* \log(S/S^*).$$

Since $S \leq 1$, we further have

$$S^* \log(S/S^*) \leq -S^* \log S^* \leq e^{-1}.$$

This completes the proof. $\square$

Lemma 15.

$$\|\pi^* - \pi\|_1 \geq 2\varepsilon \sqrt{1 - \exp(-D_{\text{KL}}(\pi^* \|\pi)/\varepsilon)}.$$

Proof. Let

$$Q = \sum_{a \in \mathbb{A}} (\pi_a^* - \pi_a)^- = \sum_{a \in \mathbb{A}} (\pi_a^* - \pi_a)^+ = \frac{1}{2} \|\pi^* - \pi\|_1.$$

Note that for fixed $Q < \varepsilon$, $\pi_a > 0$ for all $a \in \mathbb{A}$. Hence in this regime $D_{\text{KL}}(\pi^* \|\pi)$ is finite, so it attains a maximum at some π . We restrict to such maximizers.

Pick $a \in \mathbb{A}$ with $\pi_a^* = \varepsilon$. If multiple exist, choose a minimizing π_a . We claim no $a' \neq a$ can satisfy $\pi_{a'} < \pi_{a'}^*$. Suppose such a' exists. Set

$$p = \pi_a^*, \quad q = \pi_{a'}^*, \quad x = \pi_a^* - \pi_a, \quad y = \pi_{a'}^* - \pi_{a'}.$$

Consider replacing $(\pi_a, \pi_{a'})$ with $(\pi_a - (q - \pi_{a'}), q)$. The KL contribution changes from

$$-p \log(p - x) - q \log(q - y)$$

to

$$-p \log(p - x - y) - q \log q.$$

Thus we need to check

$$q \log \frac{q}{q - y} < p \log \frac{p - x}{p - x - y}. \quad (4)$$

Define $f(q) = q \log \frac{q}{q - y}$. Then

$$f'(q) = \log \left(1 + \frac{y}{q - y} \right) - \frac{y}{q - y} < 0,$$

since $q \geq p > x + y > y$. Hence $f(q) \leq f(p)$. Moreover,

$$f(p) \leq p \log \frac{p - x}{p - x - y}$$

is equivalent to $p(p - x - y) \leq (p - x)(p - y)$, which holds.

Now, equality in Equation (4) would require simultaneously $f(q) = f(p)$ and $p(p - x - y) = (p - x)(p - y)$. The first condition holds only when $p = q$, and the second only when $x = 0$ (since by assumption $y > 0$). Thus equality would force $p = q$, $x = 0$, and $y > 0$, which contradicts

the choice of a : in that case $\pi_a^* = \pi_{a'}^*$, but $\pi_{a'} < \pi_a$. Therefore strict inequality holds, and the KL strictly increases, contradicting maximality.

This proves uniqueness of such a . A symmetric argument shows there is also a unique $b \in \mathbb{A}$ with $\pi_b > \pi_b^*$.

Hence

$$D_{\text{KL}}(\pi^* \| \pi) = \pi_a^* \log \frac{\pi_a^*}{\pi_a^* - Q} + \pi_b^* \log \frac{\pi_b^*}{\pi_b^* + Q} \leq \varepsilon \log \frac{\varepsilon^2}{\varepsilon^2 - Q^2}.$$

Solving for Q gives

$$Q \geq \varepsilon \sqrt{1 - \exp(-D_{\text{KL}}(\pi^* \| \pi)/\varepsilon)},$$

which implies the desired bound since $\|\pi^* - \pi\|_1 = 2Q$. $\square$

C.4 UNIFIED ANALYSIS OF ALL CASES

Lemma 16.

$$\Theta_t - \Theta_{t+1} \geq \frac{1 - 4\eta^2 L^2}{(\sqrt{2}\eta L + 2e^{2\eta L})^2} \max\{\hat{K}^{(t+1)}, K^{(t+1)}\}^2$$

Proof. From Equation (2), we have

$$\log \hat{\pi}_a^{(t+1)} - \langle \hat{\pi}^{(t+1)}, \log \hat{\pi}^{(t+1)} \rangle = \log \hat{\pi}_a^{(t)} + \eta(\mathbf{P}\pi^{(t)})_a - \langle \hat{\pi}^{(t+1)}, \log \hat{\pi}^{(t)} + \eta\mathbf{P}\pi^{(t)} \rangle.$$

Rearranging gives

$$\eta(\mathbf{P}\pi^{(t)})_a = \langle \eta\mathbf{P}\pi^{(t)}, \hat{\pi}^{(t+1)} \rangle + (\log \hat{\pi}_a^{(t+1)} - \log \hat{\pi}_a^{(t)}) - D_{\text{KL}}(\hat{\pi}^{(t+1)} \| \hat{\pi}^{(t)}).$$

Since $\langle \mathbf{P}\pi^{(t)}, \pi^{(t)} \rangle = 0$, the first term satisfies

$$\langle \eta\mathbf{P}\pi^{(t)}, \hat{\pi}^{(t+1)} \rangle = \langle \eta\mathbf{P}\pi^{(t)}, \hat{\pi}^{(t+1)} - \pi^{(t)} \rangle \leq \eta L \|\hat{\pi}^{(t+1)} - \pi^{(t)}\|_1.$$

For the second term, using

$$|\log a - \log b| \leq \frac{|a - b|}{\min\{a, b\}} \leq \frac{\exp(|\log a - \log b|)}{\max\{a, b\}} |a - b|,$$

together with Lemma 11, we obtain

$$|\log \hat{\pi}_a^{(t+1)} - \log \hat{\pi}_a^{(t)}| \leq \frac{e^{2\eta L}}{\max\{\hat{\pi}_a^{(t)}, \hat{\pi}_a^{(t+1)}\}} \|\hat{\pi}^{(t+1)} - \hat{\pi}^{(t)}\|_1.$$

Thus,

$$|\eta(\mathbf{P}\pi^{(t)})_a| \leq \frac{e^{2\eta L}}{\max\{\hat{\pi}_a^{(t)}, \hat{\pi}_a^{(t+1)}\}} \|\hat{\pi}^{(t+1)} - \hat{\pi}^{(t)}\|_1 + \eta L \|\hat{\pi}^{(t+1)} - \pi^{(t)}\|_1.$$

Multiplying both sides by $\max\{\hat{\pi}_a^{(t)}, \hat{\pi}_a^{(t+1)}\}$ and taking the maximum over $a \in \mathbb{A}$, we obtain

$$\max\{\hat{K}^{(t+1)}, K^{(t+1)}\} \leq e^{2\eta L} \|\hat{\pi}^{(t+1)} - \hat{\pi}^{(t)}\|_1 + \eta L \|\hat{\pi}^{(t+1)} - \pi^{(t)}\|_1.$$

By Lemma 7,

$$\begin{aligned} D_{\text{KL}}(\pi^{(t)} \| \hat{\pi}^{(t)}) + D_{\text{KL}}(\hat{\pi}^{(t+1)} \| \pi^{(t)}) &\geq \frac{1}{2} \|\pi^{(t)} - \hat{\pi}^{(t)}\|_1^2 + \frac{1}{2} \|\hat{\pi}^{(t+1)} - \pi^{(t)}\|_1^2 \\ &\geq \frac{1}{4} \|\hat{\pi}^{(t+1)} - \hat{\pi}^{(t)}\|_1^2. \end{aligned}$$

Hence,

$$\begin{aligned} &(\sqrt{2}\eta L + 2e^{2\eta L}) \sqrt{D_{\text{KL}}(\pi^{(t)} \| \hat{\pi}^{(t)}) + D_{\text{KL}}(\hat{\pi}^{(t+1)} \| \pi^{(t)})} \\ &\geq \eta L \|\hat{\pi}^{(t+1)} - \pi^{(t)}\|_1 + e^{2\eta L} \|\hat{\pi}^{(t+1)} - \hat{\pi}^{(t)}\|_1 \end{aligned}$$

$$\geq \max\{\hat{K}^{(t+1)}, K^{(t+1)}\}.$$

Thus, by Lemma 10,

$$\Theta_t - \Theta_{t+1} \geq \frac{1 - 4\eta^2 L^2}{(\sqrt{2}\eta L + 2e^{2\eta L})^2} \max\{\hat{K}^{(t+1)}, K^{(t+1)}\}^2.$$

□

Remark 6. For notational simplicity, define

$$C_1 = \frac{1 - 4\eta^2 L^2}{(\sqrt{2}\eta L + 2e^{2\eta L})^2}.$$

Recall that

$$\Phi_t = \frac{\langle \log \hat{\pi}^{(t)} - \log \pi^*, \eta P \hat{\pi}^{(t)} \rangle}{(\Theta_t + 2e^{-1})^2}.$$

Thus,

$$\begin{aligned} \Phi_t - \Phi_{t+1} &= \frac{\langle \log \hat{\pi}^{(t+1)} - \log \pi^*, \eta P \hat{\pi}^{(t)} - \eta P \hat{\pi}^{(t+1)} \rangle + \langle \log \hat{\pi}^{(t)} - \log \hat{\pi}^{(t+1)}, \eta P \hat{\pi}^{(t)} \rangle}{(\Theta_t + 2e^{-1})^2} \\ &\quad + \langle \log \hat{\pi}^{(t+1)} - \log \pi^*, \eta P \hat{\pi}^{(t+1)} \rangle \left(\frac{1}{(\Theta_t + 2e^{-1})^2} - \frac{1}{(\Theta_{t+1} + 2e^{-1})^2} \right) \\ &\geq \frac{\langle \log \hat{\pi}^{(t+1)} - \log \pi^*, \eta P \hat{\pi}^{(t)} - \eta P \hat{\pi}^{(t+1)} \rangle + \langle \log \hat{\pi}^{(t)} - \log \hat{\pi}^{(t+1)}, \eta P \hat{\pi}^{(t)} \rangle}{(\Theta_t + 2e^{-1})^2} \\ &\quad - \frac{\eta L \|\log \hat{\pi}^{(t+1)} - \log \pi^*\|_1 (\Theta_{t+1} + \Theta_t + 4e^{-1})}{(\Theta_t + 2e^{-1})^2 (\Theta_{t+1} + 2e^{-1})^2} (\Theta_t - \Theta_{t+1}). \end{aligned}$$

We bound each term separately.

Lemma 17.

$$\begin{aligned} \langle \log \hat{\pi}^{(t)} - \log \hat{\pi}^{(t+1)}, \eta P \hat{\pi}^{(t)} \rangle &\geq \frac{\eta^2}{4} \|P\pi^{(t)}\|_\infty^2 - \frac{|\mathbb{A}|^3 e^{2\eta L} (\hat{K}^{(t+1)})^2}{4} \\ &\quad - \frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)}{2\varepsilon} |\mathbb{A}| (\Theta_t + \Theta_{t+1} + 4e^{-1}) K^{(t)}. \end{aligned}$$

Proof. From the update rules,

$$\begin{aligned} \langle \log \hat{\pi}^{(t)} - \log \hat{\pi}^{(t+1)}, \eta P \hat{\pi}^{(t)} \rangle &= \langle \log \hat{\pi}^{(t)} - \log \hat{\pi}^{(t+1)}, \eta P \hat{\pi}^{(t)} - \eta P \pi^{(t)} \rangle \\ &\quad + \langle \eta P \pi^{(t)} + \hat{M}^{(t+1)}, \eta P \pi^{(t)} \rangle. \end{aligned}$$

Since $\eta \|P\pi^{(t)}\|_\infty \geq |\hat{M}^{(t+1)}|$, and for some $a \in \mathbb{A}$, $\eta (P\pi^{(t)})_a = \eta \|P\pi^{(t)}\|_\infty$, we obtain

$$\begin{aligned} \langle \eta P \pi^{(t)} + \hat{M}^{(t+1)}, \eta P \pi^{(t)} \rangle &\geq (\eta \|P\pi^{(t)}\|_\infty + \hat{M}^{(t+1)}) \cdot \eta \|P\pi^{(t)}\|_\infty - \frac{|\mathbb{A}| - 1}{4} (\hat{M}^{(t+1)})^2 \\ &= \left(\eta \|P\pi^{(t)}\|_\infty + \frac{1}{2} \hat{M}^{(t+1)} \right)^2 - \frac{|\mathbb{A}|}{4} (\hat{M}^{(t+1)})^2 \\ &\geq \frac{\eta^2}{4} \|P\pi^{(t)}\|_\infty^2 - \frac{|\mathbb{A}|^3 e^{2\eta L} (\hat{K}^{(t+1)})^2}{4}. \end{aligned}$$

Moreover,

$$\begin{aligned} \langle \log \hat{\pi}^{(t)} - \log \hat{\pi}^{(t+1)}, \eta P \hat{\pi}^{(t)} - \eta P \pi^{(t)} \rangle &\geq -\eta L (\|\log \hat{\pi}^{(t)} - \log \pi^*\|_1 + \|\log \hat{\pi}^{(t+1)} - \log \pi^*\|_1) \|\hat{\pi}^{(t)} - \pi^{(t)}\|_1 \\ &\geq -\frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)}{2\varepsilon} |\mathbb{A}| (\Theta_t + \Theta_{t+1} + 4e^{-1}) K^{(t)}, \end{aligned}$$

where the last step follows from Lemma 13 and Lemma 14. □

Lemma 18.

$$\langle \log \hat{\pi}^{(t+1)} - \log \pi^*, \eta \mathbf{P} \hat{\pi}^{(t)} - \eta \mathbf{P} \hat{\pi}^{(t+1)} \rangle \geq -\frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)}{2\varepsilon} |\mathbb{A}| (\Theta_{t+1} + 2e^{-1}) \hat{K}^{(t+1)}.$$

Proof. By Lemma 13 and Lemma 14,

$$\begin{aligned} \langle \log \hat{\pi}^{(t+1)} - \log \pi^*, \eta \mathbf{P} \hat{\pi}^{(t)} - \eta \mathbf{P} \hat{\pi}^{(t+1)} \rangle &\geq -\eta L \|\log \hat{\pi}^{(t+1)} - \log \pi^*\|_1 \|\hat{\pi}^{(t)} - \hat{\pi}^{(t+1)}\|_1 \\ &\geq -\frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)}{2\varepsilon} |\mathbb{A}| (\Theta_{t+1} + 2e^{-1}) \hat{K}^{(t+1)}. \end{aligned}$$

□

Lemma 19.

$$\frac{\Theta_{t+1} + \Theta_t + 4e^{-1}}{(\Theta_t + 2e^{-1})^2 (\Theta_{t+1} + 2e^{-1})^2} \leq \frac{1}{4} e^3.$$

Proof. Consider $f(u, v) = \frac{u+v}{u^2 v^2}$ for $u, v > 0$. Then

$$\partial_u f(u, v) = -\frac{u+2v}{u^3 v^2} < 0, \quad \partial_v f(u, v) = -\frac{2u+v}{u^2 v^3} < 0.$$

Thus,

$$f(\Theta_t + 2e^{-1}, \Theta_{t+1} + 2e^{-1}) \leq f(2e^{-1}, 2e^{-1}) = \frac{1}{4} e^3.$$

□

Combining Lemma 17, Lemma 18, Lemma 19, and $\Theta_t \geq \Theta_{t+1}$, we obtain

$$\begin{aligned} \Phi_t - \Phi_{t+1} &\geq \frac{\eta^2 \|\mathbf{P} \pi^{(t)}\|_\infty^2}{4(\Theta_t + 2e^{-1})^2} - \frac{|\mathbb{A}|^3 e^{2\eta L} (\hat{K}^{(t+1)})^2}{4(\Theta_t + 2e^{-1})^2} - \frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)}{\varepsilon(\Theta_t + 2e^{-1})} |\mathbb{A}| K^{(t)} \\ &\quad - \frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)}{2\varepsilon(\Theta_t + 2e^{-1})} |\mathbb{A}| \hat{K}^{(t+1)} - \frac{1}{4} \eta L e^3 (\Theta_t - \Theta_{t+1}). \end{aligned}$$

Define

$$\bar{\Theta}_t = \Theta_t + \frac{4|\mathbb{A}|(e^{2\eta L} - 1)(e^{\eta L} + 1)}{\varepsilon \eta^2} (\Theta_t + 2e^{-1}).$$

Lemma 20. For all $t \geq 0$,

$$\bar{\Theta}_t - \bar{\Theta}_{t+1} \geq \frac{\|\mathbf{P} \pi^{(t)}\|_\infty^2}{4(\Theta_t + 2e^{-1})^2} - \frac{|\mathbb{A}|^3 e^{2\eta L} (\hat{K}^{(t+1)})^2}{4\eta^2 (\Theta_t + 2e^{-1})^2} - \frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)}{2\varepsilon \eta^2 (\Theta_t + 2e^{-1})} |\mathbb{A}| \hat{K}^{(t+1)}.$$

Proof. From the definition of $\bar{\Theta}_t$ and the bound on $\Phi_t - \Phi_{t+1}$,

$$\begin{aligned} \bar{\Theta}_t - \bar{\Theta}_{t+1} &= (\Theta_t - \Theta_{t+1}) + \frac{4|\mathbb{A}|(e^{2\eta L} - 1)(e^{\eta L} + 1)}{\varepsilon \eta^2} (\Theta_t - \Theta_{t+1}) \\ &\geq (\Theta_t - \Theta_{t+1}) + \frac{4(\Theta_t + 2e^{-1})(\Phi_t - \Phi_{t+1})}{\eta^2}. \end{aligned}$$

Substituting the bound on $\Phi_t - \Phi_{t+1}$ completes the proof. □

Next, we show that $\bar{\Theta}_t$ decreases sufficiently fast in each iteration.

Lemma 21. For all $t \geq 0$,

$$\bar{\Theta}_t - \bar{\Theta}_{t+1} \geq \frac{1}{16(\Theta_t + 2e^{-1})^2} \max \left\{ \|\mathbf{P} \pi^{(t)}\|_\infty^2, (\hat{K}^{(t+1)})^2 \right\}.$$

Proof. From Lemma 20, it suffices to show

$$\frac{\|\mathbf{P}\boldsymbol{\pi}^{(t)}\|_\infty^2}{4(\Theta_t + 2e^{-1})^2} - \frac{|\mathbb{A}|^3 e^{2\eta L} (\hat{K}^{(t+1)})^2}{4\eta^2(\Theta_t + 2e^{-1})^2} - \frac{(e^{2\eta L} - 1)(e^{\eta L} + 1)}{2\varepsilon\eta^2(\Theta_t + 2e^{-1})} |\mathbb{A}| \hat{K}^{(t+1)}$$

is at least

$$\frac{1}{16(\Theta_t + 2e^{-1})^2} \max \left\{ \|\mathbf{P}\boldsymbol{\pi}^{(t)}\|_\infty^2, (\hat{K}^{(t+1)})^2 \right\}.$$

If $\|\mathbf{P}\boldsymbol{\pi}^{(t)}\|_\infty \geq 2\hat{K}^{(t+1)}$, the negative terms are dominated by the positive part, yielding

$$\bar{\Theta}_t - \bar{\Theta}_{t+1} \geq \frac{1}{8(\Theta_t + 2e^{-1})^2} \|\mathbf{P}\boldsymbol{\pi}^{(t)}\|_\infty^2.$$

Otherwise, $\hat{K}^{(t+1)} \geq \frac{1}{2} \|\mathbf{P}\boldsymbol{\pi}^{(t)}\|_\infty$. Substituting this relation, the inequality simplifies to

$$\bar{\Theta}_t - \bar{\Theta}_{t+1} \geq \frac{1}{16(\Theta_t + 2e^{-1})^2} (\hat{K}^{(t+1)})^2.$$

This completes the proof. $\square$

Combining Lemma 21 with Lemma 8 yields a global bound on the time to reach the region $\Theta_t < \varepsilon$.

Theorem 3. *There exists a constant*

$$T_1 = O\left(\frac{\eta^2 L^2 |\mathbb{A}|^2 (\bar{\Theta}_1^3 + 1)}{(1 - 4\eta^2 L^2) C_{\mathbf{P}}^4 \varepsilon^6}\right)$$

such that $\Theta_{T_1} < \varepsilon$.

Proof. If $\Theta_t \geq \varepsilon$ for all $t \leq T$, then Lemma 21 and Lemma 8 imply

$$T - 1 \leq \int_{\bar{\Theta}_T}^{\bar{\Theta}_1} \frac{3(e^{2\eta L} - 1)^2 (e^{\eta L} + 1)^2 (x + 2e^{-1})^2 |\mathbb{A}|^2}{(1 - e^{-1}) C_1 C_{\mathbf{P}}^4 \varepsilon^6} dx = O\left(\frac{\eta^2 L^2 |\mathbb{A}|^2 (\bar{\Theta}_1^3 + 1)}{(1 - 4\eta^2 L^2) C_{\mathbf{P}}^4 \varepsilon^6}\right).$$

Hence such a T_1 exists with the stated asymptotic bound. $\square$

Next we handle the phase after Θ_t enters the small region $\Theta_t < \varepsilon$.

When $\Theta_t < \varepsilon$ we have the lower bound

$$\Theta_t - \Theta_{t+1} \geq C_1 (\hat{K}^{(t+1)})^2$$

and the trivial estimate

$$\hat{K}^{(t+1)} \geq \eta \min_{a \in \mathbb{A}} \hat{\pi}_a^{(t)} \|\mathbf{P}\boldsymbol{\pi}^{(t)}\|_\infty.$$

Using

$$\pi_a^{(t)} \geq \varepsilon - \|\hat{\boldsymbol{\pi}}^{(t)} - \boldsymbol{\pi}^*\|_1 \geq \varepsilon(1 - \sqrt{1 - \exp(-\Theta_t/\varepsilon)}),$$

we obtain

$$\Theta_t - \Theta_{t+1} \geq 4\eta^2 C_1 C_{\mathbf{P}}^2 \varepsilon^4 (1 - \sqrt{1 - \exp(-\Theta_t/\varepsilon)}) (1 - \exp(-\Theta_t/\varepsilon)).$$

Set

$$w = \sqrt{1 - \exp(-\Theta_t/\varepsilon)}, \quad W = \sqrt{1 - e^{-1}}.$$

Applying Lemma 8 and the substitution $u = \sqrt{1 - \exp(-x/\varepsilon)}$ (so $dx = \frac{2u\varepsilon du}{1-u^2}$) gives

$$\begin{aligned} t - T_1 &\leq \ln \frac{W^2(1-W)}{w^2(1-w)} + \int_{\Theta_t}^{\varepsilon} \frac{dx}{4\eta^2 C_1 C_{\mathbf{P}}^2 \varepsilon^4 (1 - \sqrt{1 - \exp(-x/\varepsilon)}) (1 - \exp(-x/\varepsilon))} \\ &\leq \ln \frac{W^2(1-W)}{w^2(1-w)} + \frac{1}{2\eta^2 C_1 C_{\mathbf{P}}^2 \varepsilon^3} \int_w^W \frac{du}{u(1-u)(1-u^2)} \\ &= \ln \frac{W^2(1-W)}{w^2(1-w)} + \frac{1}{2\eta^2 C_1 C_{\mathbf{P}}^2 \varepsilon^3} \int_w^W \left(\frac{1}{u} + \frac{1+u-u^2}{(1-u)(1-u^2)} \right) du \end{aligned} \tag{5}$$

$$\leq 2 \ln \frac{W}{w} + \frac{1}{2\eta^2 C_1 C_P^2 \varepsilon^3} \left(\ln \frac{W}{w} + C_2 \right),$$

where $C_2 = \int_0^W \frac{1+x-x^2}{(1-x)(1-x^2)} dx$ is an absolute constant. The estimate in equation 5 used the substitution indicated above.

Finally, since

$$\Theta_t = -\varepsilon \ln(1 - w^2) \leq -\varepsilon \ln(1 - W^2) w^2 = \varepsilon w^2,$$

we deduce exponential decay:

$$\Theta_t \leq \varepsilon W \exp\left(-(4 + \eta^2 C_1^{-1} C_P^{-2} \varepsilon^{-3})(t - T_1 - C_2)\right).$$

This completes the analysis: after T_1 iterations to reach the small region, Θ_t decreases exponentially fast with the stated rate.

D EXPERIMENT DETAILS AND ADDITIONAL RESULTS

This sections provides additional experiment details and results.

D.1 GAME MATRIX CONSTRUCTION

We sample preference matrices using Algorithm 1.

Algorithm 1: Preference Matrix Sampling Algorithm

Require: Size of game n , rank of full-support equilibrium space m .

- 1: Sample m positive vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m$ and a random $(n - m) \times (n - m)$ matrix $\mathbf{A}$.
 - 2: Find an orthonormal basis $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{n-m}$ of the orthogonal complement of the subspace $\text{span}\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m\}$.
 - 3: Set $\mathbf{M} \leftarrow [\mathbf{u}_1 \quad \mathbf{u}_2 \quad \dots \quad \mathbf{u}_{n-m}]$ and $\mathbf{P} \leftarrow \mathbf{M}(\mathbf{A} - \mathbf{A}^\top)\mathbf{M}^\top$.
 - 4: **return** $\frac{1}{\max_{1 \leq i, j \leq n} |P_{i,j}|} \mathbf{P}$.
-

For $m \geq 1$, this guarantees that $\mathbf{P}\mathbf{v}_i = \mathbf{0}$, so $\frac{\mathbf{v}_i}{\|\mathbf{v}_i\|_1}$ is an equilibrium.

D.2 CODEBASE AND BASELINES

We chose the codebase³ open-sourced by Zhou et al. (2025) because OMD (regularized) and EGPO have been included. Necessary modifications and extensions have been made, and for completeness, we describe the implementation of baselines here. We use η as the step size and β as the regularization coefficient.

OMD. The update is

$$\boldsymbol{\theta}^{(t+1)} = \boldsymbol{\theta}^{(t)} + \eta \mathbf{P} \boldsymbol{\pi}^{(t)}.$$

The implementation uses the generalized online IPO (where we choose $\beta = 1$):

$$\boldsymbol{\theta}^{(t+1)} \leftarrow \boldsymbol{\theta}^{(t)} - \frac{\eta |\mathbb{A}|}{4} \nabla_{\boldsymbol{\theta}} \mathcal{L}_{\text{IPO}}(\boldsymbol{\theta}^{(t)}; \text{sg}[\boldsymbol{\pi}^{(t)}], \text{Uniform}, \text{sg}[\boldsymbol{\pi}^{(t)}]).$$

OMD (regularized). The update is shown as Equation (8) in Munos et al. (2023):

$$\boldsymbol{\pi}^{(t+1)} = \arg \max_{\boldsymbol{\pi}} \{\eta \boldsymbol{\pi}^\top \mathbf{P} \boldsymbol{\pi}^{(t)} - D_{\text{KL}}(\boldsymbol{\pi} \parallel \tilde{\boldsymbol{\pi}}^{(t)})\},$$

³<https://github.com/zhourunlong/EGPO>

where $\tilde{\pi}^{(t)}(y) \propto (\pi^{(t)}(y))^{1-\eta\beta} (\pi_{\text{ref}}(y))^{\eta\beta}$. In Appendix E.1 of Zhou et al. (2025), it is shown to be equivalent to

$$\theta^{(t+1)} = \theta^{(t)} - \eta\beta \left(\theta^{(t)} - \theta_{\text{ref}} - \frac{P\pi^{(t)}}{\beta} \right).$$

The implementation uses the generalized online IPO:

$$\theta^{(t+1)} \leftarrow \theta^{(t)} - \frac{\eta\beta |\mathbb{A}|}{4} \nabla_{\theta} \mathcal{L}_{\text{IPO}}(\theta^{(t)}; \pi_{\text{ref}}, \text{Uniform}, \text{sg}[\pi^{(t)}]).$$

SPPO. The update is shown as Equation (4.6) in Wu et al. (2024):

$$\pi^{(t+1)} = \arg \min_{\pi} \mathbb{E}_{y \sim \pi^{(t)}} \left(\log \frac{\pi(y)}{\pi^{(t)}(y)} - \eta \left(\mathcal{P}(y > \pi^{(t)}) - \frac{1}{2} \right) \right)^2.$$

Nested-optimization is used to update the parameters: for each iteration of t , we apply gradient descent using a learning rate η_{inner} independent of η for at most 10 steps, or until successive objects deviates by at most 10^{-5} .

MPO. The update is shown as Equation (8) in Wang et al. (2024):

$$\pi^{(t+1)} = \arg \max_{\pi} \mathbb{E}_{y_1 \sim \pi, y_2 \sim \pi_{\text{ref}}} \left[\mathcal{P}(y_1 > y_2) - \beta D_{\text{KL}}(\pi \| \pi_{\text{ref}}) - \frac{1}{\eta} D_{\text{KL}}(\pi \| \pi^{(t)}) \right].$$

As detailed in Algorithm 1 in Wang et al. (2024), each iteration of t is solved with nested-optimization similar as in SPPO. Algorithm 1 in Wang et al. (2024) incorporates two tricks: (1) step size annealing: $\eta_t = \max\{1 - t/T, 5 \times 10^{-4}\}$ (where t starts from 0); (2) reference policy refreshing: with every $\tau = 1000$ (our choice) steps, update $\pi_{\text{ref}} \leftarrow \pi^{(i\tau)}$.

ONPO. The update is shown in Section 4.2 of Zhang et al. (2025):

$$\begin{aligned} \pi^{(t)} &= \arg \max_{\pi} \left\{ \eta \pi^{\top} \mathcal{P} \pi^{(t-1)} - D_{\text{KL}}(\pi \| \hat{\pi}^{(t)}) \right\}, \\ \hat{\pi}^{(t+1)} &= \arg \max_{\pi} \left\{ \eta \pi^{\top} \mathcal{P} \pi^{(t)} - D_{\text{KL}}(\pi \| \hat{\pi}^{(t)}) \right\}. \end{aligned}$$

Each iteration of t is solved with nested-optimization twice, similar as in SPPO.

EGPO. The update is shown in Equations (2) and (3) in Zhou et al. (2025):

$$\begin{aligned} \theta^{(t+1/2)} &= (1 - \eta\beta) \theta^{(t)} + \eta\beta \left(\theta_{\text{ref}} + \frac{P\pi^{(t)}}{\beta} \right), \\ \theta^{(t+1)} &= (1 - \eta\beta) \theta^{(t)} + \eta\beta \left(\theta_{\text{ref}} + \frac{P\pi^{(t+1/2)}}{\beta} \right). \end{aligned}$$

The implementation uses the generalized online IPO:

$$\begin{aligned} \theta^{(t+1/2)} &= \theta^{(t)} - \frac{\eta\beta |\mathbb{A}|}{4} \nabla_{\theta} \mathcal{L}_{\text{IPO}}(\theta^{(t)}; \pi_{\text{ref}}, \text{Uniform}, \text{sg}[\pi^{(t)}]), \\ \theta^{(t+1)} &= \theta^{(t)} - \frac{\eta\beta |\mathbb{A}|}{4} \nabla_{\theta} \mathcal{L}_{\text{IPO}}(\theta^{(t)}; \pi_{\text{ref}}, \text{Uniform}, \pi^{(t+1/2)}). \end{aligned}$$

D.3 HYPERPARAMETERS

Neural network architecture. We use a 3-layer MLP with ReLU activation as the neural policy. The hidden dimension d is set to be 10. Since we consider multi-armed bandit environments, there is no input to this policy. Hence, we use a random Gaussian noise $\mathcal{N}(0, I_d)$ as input.

Reference policy. For tabular policies, we initialize all the parameters to be 0 to assign the reference policy with the uniform policy. For neural policies, we use Xavier normal initialization (Glorot & Bengio, 2010) in all the middle layers, and zero initialization in the output layer, also assigning the reference policy with the uniform policy while avoiding symmetry weights and homogeneous gradients.

Training arguments. For each algorithm under each scenario (tabular v.s. neural), we performed a grid search over all hyperparameters $\mathcal{V}_\eta \times \mathcal{V}_{\eta_{\text{inner}}}$:

$$\mathcal{V}_\eta = \{i \times 10^j : 1 \leq i \leq 10, -4 \leq j \leq 1\},$$

$$\mathcal{V}_{\eta_{\text{inner}}} = \{0.00003, 0.0001, 0.0003, 0.001, 0.003, 0.01, 0.03, 0.1\}.$$

For MPO, since η_t is defined, we did not search over $\mathcal{V}_\eta$; for OMD, OMD (regularized), EGPO and OMWU, since no nested-optimization is needed, we did not search over $\mathcal{V}_{\eta_{\text{inner}}}$.

We use Algorithm 1 to sample 100 game matrices, and select the hyperparameters by first ensuring convergence in the most games, then prioritizing the minimal final duality gap. The final chosen hyperparameters are listed in Table 3.

Table 3: Hyperparameters for different algorithms across tabular and neural settings.

Algorithm	Tabular LR	Tabular η	Neural LR	Neural η
OMD	0.4		10	
OMD (regularized)	0.001		0.0002	
SPPO	0.03	0.1	0.03	0.1
MPO	0.0003		0.09	
ONPO	0.01	0.01	0.01	0.01
EGPO	0.01		0.09	
OMWU	9		100	

D.4 FULL RESULTS

Here we present full experiment results. In all the mentioned figures, duality gap values are cut off below 10^{-6} due to floating point precision.

Last-iterate v.s. average-iterate convergence. In Figure 3, we display the duality gap of OMD (regularized) on tabular policies. These results illustrate the importance of last-iterate convergence guarantees.

Comparisons between algorithms. Figures 4 and 5 are results of different algorithms on both tabular and neural policies. For algorithms with only average-iterate convergence guarantees, we only display the duality gap curves of their average policies. Even with extensive hyperparameter search, SPPO, MPO, and ONPO all fail due to slow and unstable nested-optimization.

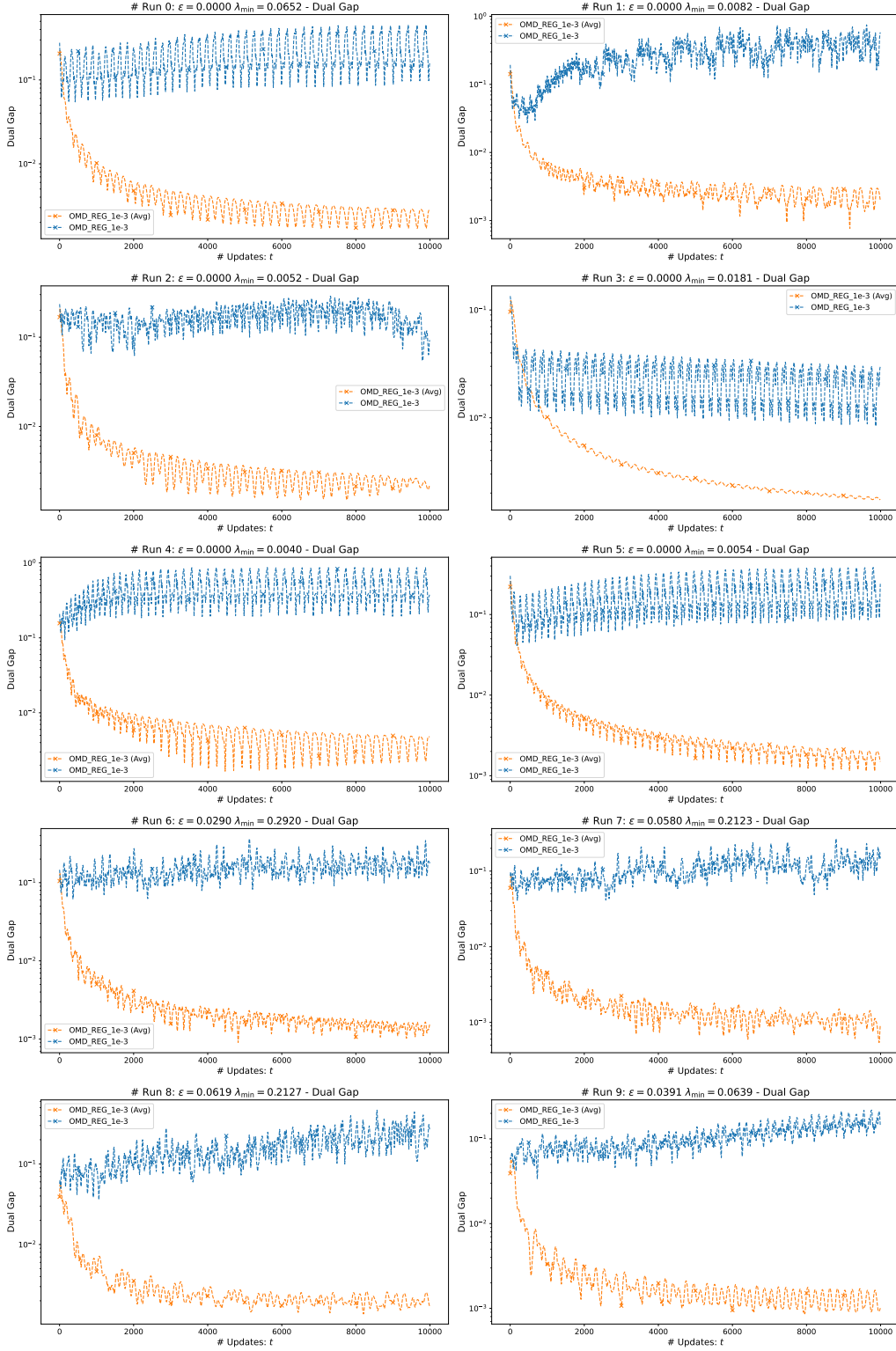
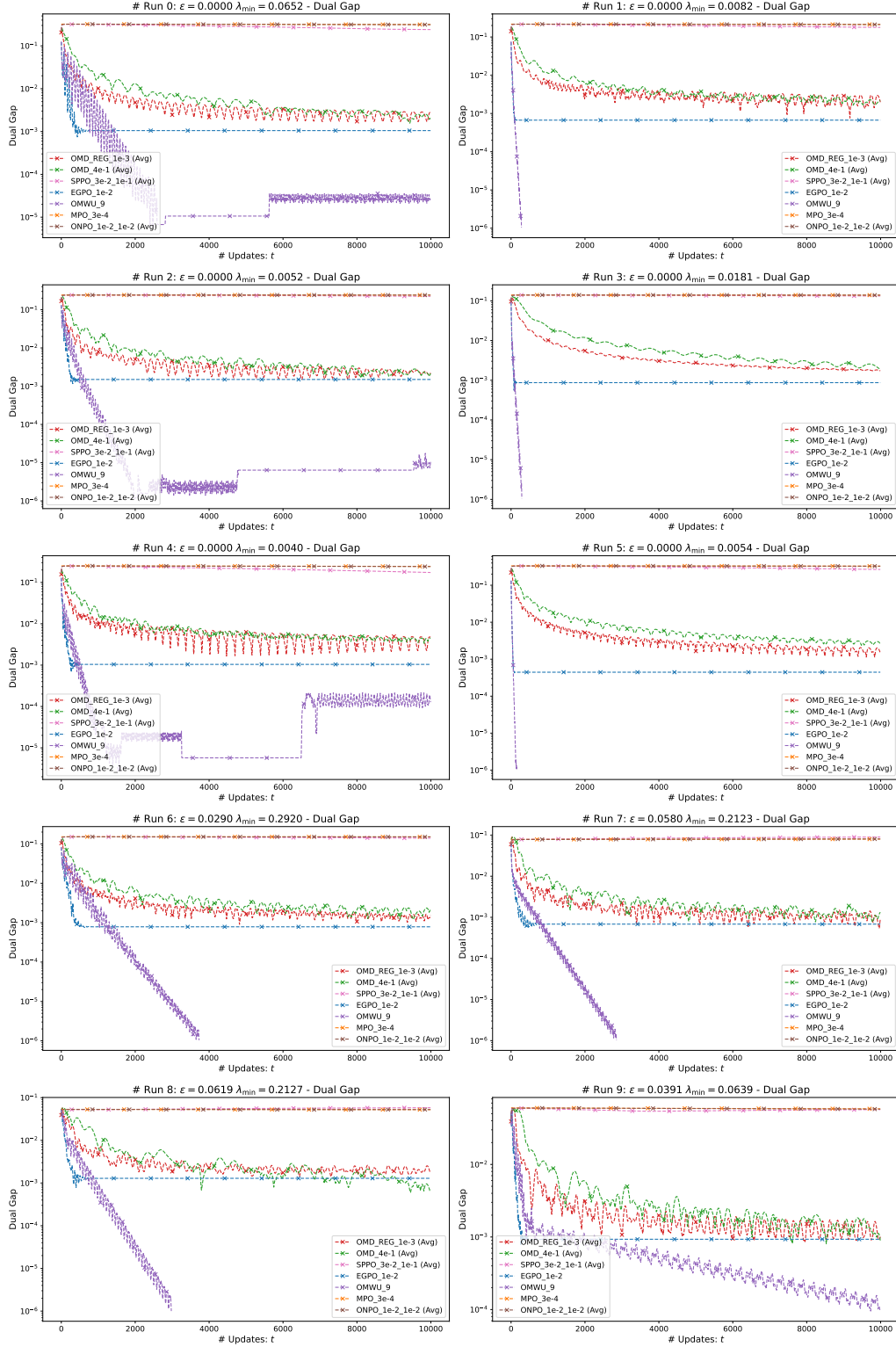
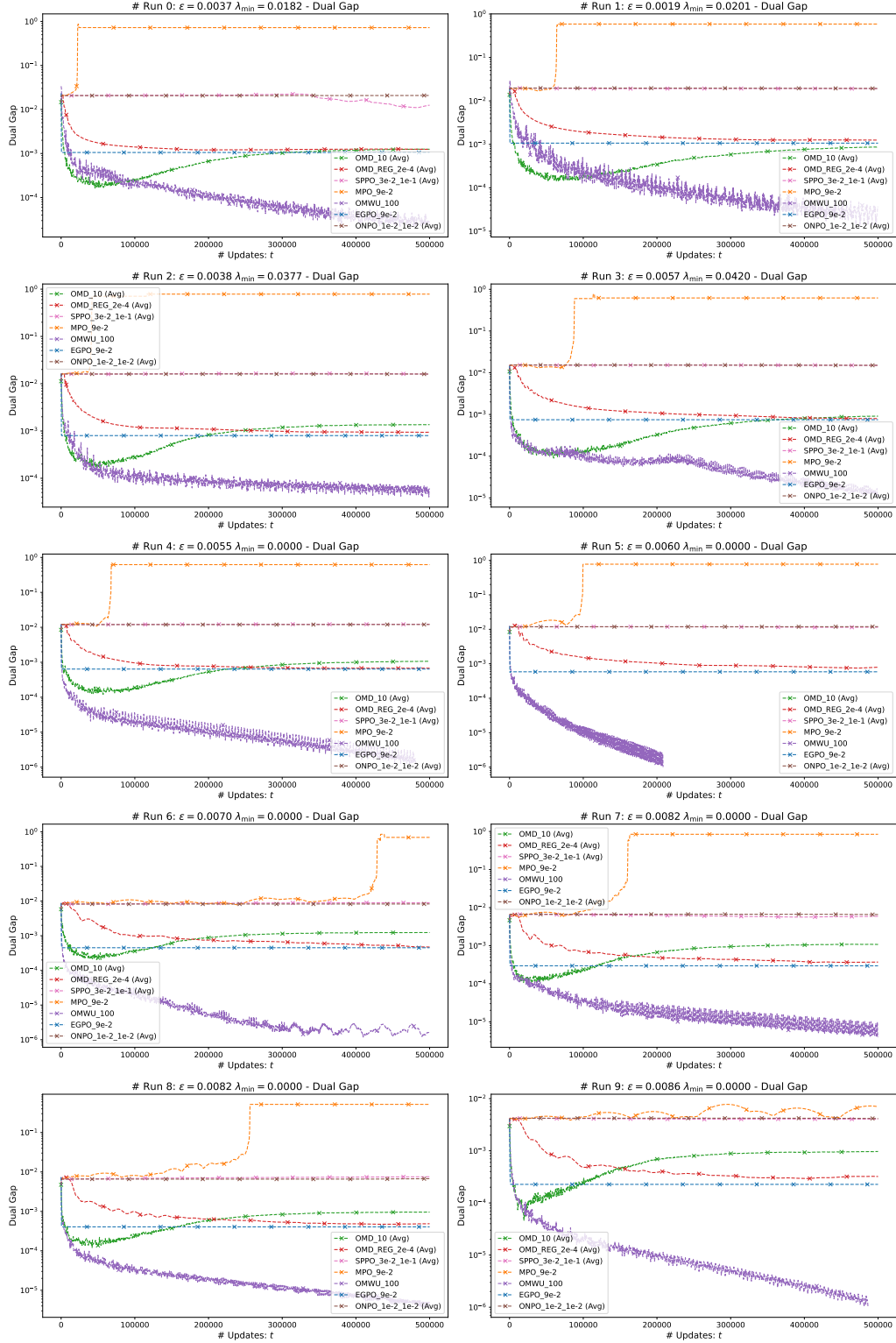


Figure 3: Duality gaps of OMD (regularized) applied to **tabular** policies, when evaluating the last-iterate policy $\pi^{(t)}$ and the average-iterate policy $\frac{1}{t} \sum_{i=1}^t \pi^{(i)}$.

Figure 4: Duality gaps of different algorithms applied to **tabular** policies.

Figure 5: Duality gaps of different algorithms applied to **neural** policies.