
Fine-tuning Behavioral Cloning Policies with Preference-Based Reinforcement Learning

Anonymous Author(s)

Affiliation

Address

email

Abstract

Deploying reinforcement-learning (RL) controllers in robotics, industry, and health care is blocked by two coupled obstacles: reward misspecification (informal goals are hard to encode as a safe numeric signal) and data-hungry exploration. We tackle these issues with a two-stage framework that begins from a reward-free dataset of expert demonstrations and refines the policy online using preference-based human feedback. We give the first principled analysis of this two-stage paradigm. In our work, we formulate a unified algorithm that (i) clones demonstrations offline to obtain a safe warm-start policy and (ii) fine-tunes it online with preference-based RL, integrating the two signals through an uncertainty-weighted objective. Then, we derive regret bounds that shrink with the demonstration counts and reflect reduced uncertainty.

1 Introduction

Deploying reinforcement-learning (RL) [Sutton and Barto, 2018] systems on physical robots, industrial processes, and health-care problems remains notoriously difficult for two intertwined reasons. First, reward misspecification: even experienced domain experts often find it hard to translate informal task goals into a numeric signal that is simultaneously accurate and safe [Leike et al., 2018]. Second, exploration is both risky and data-hungry [Dulac-Arnold et al., 2019]: a policy that begins from scratch can damage hardware, or user trust, long before it gathers enough experience to learn anything useful. Recent applied works [Nair et al., 2020, Kostrikov et al., 2022, Tang et al., 2025, Park et al., Tirinzoni et al., 2025] alleviate these issues by pre-training policies offline and fine-tuning them with online RL, but such methods assume direct access to (or observation of) the true reward, an assumption that rarely holds in practice. A more realistic strategy is to start with a reward-free corpus of expert behavior and allow the experts to refine that behavior online. The refinement signal can take several forms: explicit numeric rewards from an operator, pairwise comparisons of trajectories, or scalar ratings that train a reward model for RL from human feedback (RLHF). Variants of this idea already power modern dialogue agents such as ChatGPT, which first imitate curated demonstrations of desirable responses and are then fine-tuned via RLHF on preference-derived rewards [Ouyang et al., 2022]. Similar approaches reach near-expert scores in Atari and MuJoCo by combining brief expert play with thousands of comparison queries [Christiano et al., 2017], or recover usable rewards for real-robot manipulation by ranking tele-operated clips before on-hardware fine-tuning [Brown et al., 2020].

This paper formalizes a principled two-stage procedure: offline demonstrations supply a safe warm start, while lightweight online feedback repairs the blind spots of the behavioral-cloning policy. Merging these two information sources yields both higher sample efficiency and stronger safety guarantees. Demonstrations guide the agent away from dangerous or hard-to-reach regions of state

space, so exploration seldom visits unsafe states. Theoretically, when the offline data already explain many state–action pairs, regret bounds shrink because the set of plausible optimal policies is greatly reduced. Pragmatically, preference queries remain easy to elicit even when explicit rewards are not, which lets non-technical stakeholders participate in the corrective loop. We conclude by detailing our contributions:

- We develop a novel policy confidence set framework based on Hellinger distances between trajectory distributions. By separating the policy and transition components of the MLE objective, we extend Foster et al. [2024]’s framework to obtain distribution-level guarantees. This confidence set (Theorem 5) is geometrically interpretable as a Hellinger ball in trajectory distribution space while providing a corresponding constraint on allowable policies. Its radius decreases with the offline sample size, effectively leveraging demonstration data to restrict the policy search space. The approach generalizes to various policy and transition model classes through appropriate covering number arguments.
- We adapt the online preference-based learning framework to leverage our offline estimation components, resulting in BRIDGE (Algorithm 1). By constraining policy comparisons to our confidence set and initializing with the MLE transition estimate, our analysis yields a regret bound (Theorem 9) that exhibits optimal $\sqrt{T}$ dependence while demonstrating how offline data reduces online regret. The bound contains terms that explicitly diminish with increasing offline sample size n , and importantly, for any fixed horizon T , as $n \rightarrow \infty$, the regret approaches zero. This theoretical result confirms that high-quality offline demonstrations can dramatically improve online learning efficiency, creating a principled bridge between imitation learning and preference-based fine-tuning.

2 Related work

Behaviour Cloning (BC). BC reformulates RL as supervised learning on expert (state, action) pairs, pioneered by road-following systems like ALVINN [Pomerleau, 1988]. Recent theoretical advances by Foster et al. [2024] establish horizon-free sample complexity bounds under deterministic policies and sparse rewards. Other algorithms such as DAgger [Ross et al., 2011] mitigate the covariate shift during deployment through iterative expert corrections, achieving no-regret guarantees [Ross et al., 2011]. Our method inherits BC’s simplicity but circumvents DAgger’s need for ongoing expert availability through preference-based refinement.

Online RL with Offline datasets The paradigm of initializing policies through offline pre-training followed by online fine-tuning has gained considerable traction, mirroring successes in supervised learning. Early contributions in this domain include model-based algorithms tailored for hybrid settings, such as the work by Ross and Bagnell [2012]. After that, Xie et al. [2021] studied this hybrid RL setting and showing that offline data does not yield statistical improvements in tabular MDPs. This is different from our result, due to our expert’s data. Recently have been proposed empirical RL algorithms designed to be effective in both offline and online contexts, aiming to facilitate seamless offline-to-online fine-tuning [Rajeswaran et al., 2017, Hester et al., 2018, Nair et al., 2018, Vecerik et al., 2017, Lee et al., 2022, Ball et al., 2023]. On the more theoretical side Song et al. [2023], Wagenmaker and Pacchiano [2023], Tang et al. [2023] proposes statistical approaches to efficiently combine offline and online datasets. Although these methods are related to our work, they assume access to numeric rewards during fine-tuning. Our approach eliminates this, assuming access to an expert trajectories dataset and preference-based online feedback.

Prior imitation-only approaches lack robustness outside the demonstration manifold; offline RL fine-tuning usually demands ground-truth rewards. Our work bridges these gaps by (i) proving that demonstration coverage plus a modest preference-query budget yields sharper high-probability regret bounds, and (ii) showing empirically that preference-guided exploration fixes blind spots with far fewer risky interactions than pure online RL.

3 Problem formulation

We address the challenge of learning optimal policies by combining information from two complementary sources: offline expert demonstrations and online preference feedback. In this hybrid learning paradigm, we first leverage a dataset $\mathbb{D}$ of trajectories collected from an expert policy to

establish strong priors over the policy space. Then, we strategically utilize these priors to guide an online preference-based learning process, where an expert provides binary feedback comparing pairs of trajectories. This framework enables us to efficiently narrow the search space using offline demonstrations while refining our understanding of the expert’s underlying preference model through targeted online queries. We aim to quantify how knowledge from offline demonstrations translates to improved regret bounds in the online preference learning phase.

Finite MDP Setting (Reward Free). Consider a finite-horizon Markov Decision Process (MDP) defined by the tuple $\mathcal{M} = (\mathcal{S}, \mathcal{A}, P, H)$, where $\mathcal{S}$ is a finite state space, $\mathcal{A}$ is a finite action space, $H \in \mathbb{N}$ is the horizon length, and $P = \{P_h\}_{h \in [H]}$ represents the time-dependent transition dynamics, with $P_h(\cdot | s, a)$ denoting the probability distribution over next states given state-action pair (s, a) at step h . A policy $\pi = \{\pi_h\}_{h \in [H]}$ consists of a collection of mappings $\pi_h : \mathcal{S} \rightarrow \Delta(\mathcal{A})$, where $\Delta(\mathcal{A})$ is the probability simplex over actions. A trajectory $\tau = \{(s_h, a_h)\}_{h \in [H]}$ is a sequence of state-action pairs generated by executing a policy π in the environment following dynamics P . We denote the space of all possible trajectories as $\mathcal{T}$. We assume the trajectories have a continuous distribution with respect to counting or Lebesgue measure. For ease of notation, we will write $\mathbb{P}_P^\pi$ for the density function of the trajectory distribution induced by policy π and dynamics P .

Offline demonstrations. We assume access to an offline dataset $\mathbb{D}_n^H = \{\tau_i\}_{i \in [n]}$ consisting of n independent trajectories of length H , where each $\tau_i \sim \mathbb{P}_{P^*}^{\pi^*}$. This represents an imitation learning framework where trajectories are generated by an expert policy π^* interacting with the true environment dynamics P^* .

Online preference queries. We formalize preference-based learning through feature embeddings and a probabilistic preference model [Christiano et al., 2017, Saha et al., 2023]. For each trajectory $\tau \in \mathcal{T}$, we assume the existence of a trajectory embedding function $\phi : \mathcal{T} \rightarrow \mathbb{R}^d$ that is known to the learner. This creates a natural complementarity between our learning phases: while offline demonstrations provide raw trajectories that directly capture expert behavior, the embedding function transforms these complex sequences into a structured representation space that facilitates preference learning. The trajectory embedding function ϕ serves a critical purpose in our framework by enabling meaningful preference comparisons that would be difficult to perform on raw trajectories. This embedding approach provides a versatile framework that can accommodate various types of trajectory information. The flexibility of this representation allows our method to adapt to different domains and preference structures without changing the underlying learning algorithm. A policy π and dynamics P induce a distribution over trajectories, allowing us to define the policy embedding as the expected feature representation: $\phi^P(\pi) = \mathbb{E}_{\tau \sim \mathbb{P}_P^\pi} [\phi(\tau)]$.

In our work, we adopt two commonly used assumptions: bounded trajectory embeddings [Saha et al., 2023, Parker-Holder et al., 2020b] and bounded weight vectors [Filippi et al., 2010, Faury et al., 2020].

Assumption 3.1 (Bounded features). *For all $\tau \in \mathcal{T}$, the feature embeddings are bounded: $\|\phi(\tau)\|_2 \leq B$ for some known constant $B < \infty$.*

Assumption 3.2 (Bounded weights). *There exists an unknown weight vector $w^* \in \{v \in \mathbb{R}^d : \|v\|_2 \leq W\}$ where the bound $W < \infty$ is known.*

Definition 1. *The degree of non-linearity of the sigmoid σ over the parameter space (denoting the first derivative of σ by σ') is given by*

$$\kappa := \sup_{\mathbf{x} \in B_B(d), \mathbf{w} \in B_S(d)} \frac{1}{\sigma'(\mathbf{w}^\top \mathbf{x})}.$$

We model the preference feedback through a Bradley-Terry model. Given two trajectories τ_1 and τ_2 , the binary preference outcome $o_{1,2} \sim \text{Ber}(P)$ is modeled as:

$$\mathbb{P}(\tau_1 \succ \tau_2) = \mathbb{P}(o_{1,2} = 1 | \tau_1, \tau_2) = \sigma(\langle \phi(\tau_1) - \phi(\tau_2), w^* \rangle),$$

where $\sigma(x) = (1 + e^{-x})^{-1}$ is the logistic function. This formulation corresponds to a latent utility model where the inner product $\langle \phi(\tau), w^* \rangle$ represents the utility of trajectory τ .

From this model, we derive a score function for trajectories $s(\tau) = \langle \phi(\tau), w^* \rangle$ and extend it to policies as $s^P(\pi) = \mathbb{E}_{\tau \sim \mathbb{P}_P^\pi} [s(\tau)]$, where P^* are the true transition dynamics. The preference between two policies π_1 and π_2 can now be written as follows: $\mathbb{P}(\pi_1 \succ \pi_2) = \sigma(\langle \phi^P(\pi_1) -$

137 $\phi^P(\pi_2), w^*)$). This represents an expected preference over the distribution of trajectories, and
 138 captures the average preference when comparing behaviors induced by different policies.

139 **Offline estimation quality.** For the offline phase, we measure the quality of estimation using
 140 distributional distance metrics in the space of trajectory distributions. Specifically, we will construct
 141 confidence sets in the form of Hellinger balls around our estimated density policy and dynamics.
 142 Notably, the Hellinger distance relates directly to the L^2 norm between square-root densities, enabling
 143 a geometric interpretation of our confidence sets as Euclidean balls in the space of density embeddings,
 144 with computational advantages over alternative divergences. The precise construction of these
 145 confidence sets and their properties will be detailed in the Section 4.

146 **Online regret.** We quantify our online learning phase’s performance through regret measurement. In
 147 each round $t \in [T]$ of online learning, the agent selects policies π_t^1 and π_t^2 , receives binary preference
 148 feedback $o_t \in \{0, 1\}$, and accumulates regret measured against the optimal policy. We specifically
 149 use the pseudo-regret with respect to the policy class Π as in Saha et al. [2023]:

$$R_T^{\text{psr}} := \max_{\pi \in \Pi} \sum_{t=1}^T \frac{[2\phi^{P^*}(\pi) - \phi^{P^*}(\pi_t^1) - \phi^{P^*}(\pi_t^2)]^\top w^*}{2} = \sum_{t=1}^T \frac{2s^{P^*}(\pi^*) - (s^{P^*}(\pi_t^1) + s^{P^*}(\pi_t^2))}{2},$$

150 where $\pi^* := \arg \max_{\pi \in \Pi} s(\pi)$. All our performance guarantees will be expressed in terms of the
 151 MDP parameters (state space size $|S|$, action space size $|\mathcal{A}|$, horizon length H), offline data quantity
 152 n , online interaction rounds T , and confidence level δ of the offline estimation - establishing a direct
 153 connection between offline data quality and online learning efficiency.

154 **Notation.** We denote $[H] = \{1, \dots, H\}$ for $H \in \mathbb{N}$. For probability distributions P, Q , $H^2(P, Q)$ is
 155 the squared Hellinger distance and $\text{TV}(P, Q)$ the total variation distance. We denote $\mathbb{B}_2^d(R) := \{x \in$
 156 $\mathbb{R}^d : \|x\|_2 \leq R\}$ for the Euclidean ball of radius R , and for any $x \in \mathbb{R}^d$, we define $x^{\otimes 2} := xx^\top$ as
 157 the outer product.

158 4 Bridging offline behavioral cloning and online preference-based feedback

159 Our framework leverages offline expert demonstrations to improve the efficiency of online preference
 160 learning. The key insight is that we can use maximum likelihood estimation (MLE) on the offline
 161 dataset $\mathbb{D}_n^H$ to construct confidence sets in the policy space that likely contain the expert policy. This
 162 approach has two important features: First, by using Hellinger distance, we obtain confidence sets that
 163 correspond to Euclidean norm balls in the space of square-root densities, providing a geometrically
 164 intuitive interpretation. The radius of this ball shrinks at a rate of $O(1/\sqrt{n})$ with offline sample size,
 165 establishing a quantifiable relationship between offline data quantity and online learning efficiency.
 166 Second, we develop a technique to make this confidence set computable using only observed data,
 167 despite the theoretical formulation involving unknowns.

168 When we restrict the online phase of BRIDGE to sample only policies from this confidence set, we
 169 significantly reduce the number of expert preference queries needed compared to algorithms without
 170 offline data access. This hybrid approach effectively trades offline demonstrations for reduced online
 171 expert interaction. Geometrically, our confidence set has a clear interpretation as a Hellinger ball in
 172 the space of trajectory distributions. However, when mapped to the policy space, it forms a more
 173 complex shape due to the nonlinearity of the inverse functional mapping from distribution metrics to
 174 policy parameters. Figure 1 explains the relation between these quantities. In the rest of this section,
 175 we formally explain how BRIDGE uses offline behavioral cloning data to warm-start an online
 176 preference-based learning process.

177 4.1 Offline behavioral cloning and uncertainty estimators

178 We apply maximum likelihood estimation (MLE) on the offline dataset $\mathbb{D}_n^H$ of expert trajectories to
 179 obtain separate estimators $\hat{\pi}, \hat{P}$ for the optimal policy and transition model respectively. We then
 180 derive a confidence set around the estimated optimal policy $\hat{\pi}$, which constrains the policy search
 181 space in the second, online preference-based learning part of our method. Relevant corollaries and
 182 their proofs are presented in Appendix B. We start by making a standard realizability assumption.

¹Saha et al. [2023] showed that the standard preference-based regret formulation is equivalent up to constant factors, when $B, W \leq 1$.

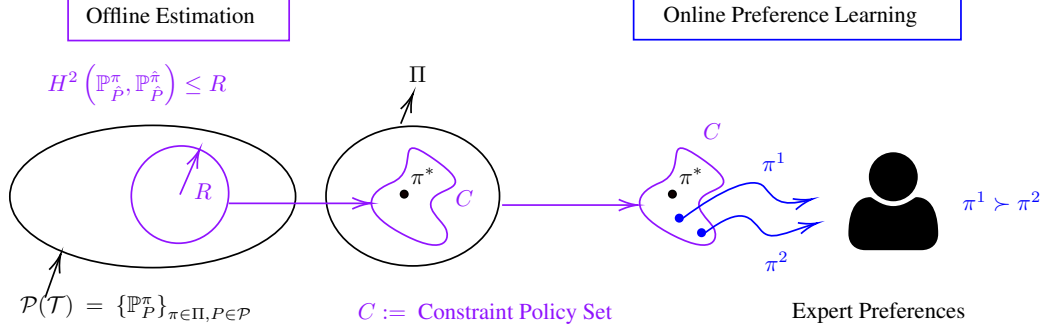


Figure 1: Framework overview: Offline estimation using dataset $\mathbb{D}_n^H$ constructs a confidence set as a Hellinger ball in trajectory distribution space (left), which translates to a constraint set in policy space containing π^* . This constraint set is then handed to the online preference learning phase (right), where policies are sampled from within this set and presented to the expert for preference feedback.

Assumption 4.1 (Realizability). *The optimal policy belongs to the policy class, $\pi^* \in \Pi$, and the true transition function belongs to the transition class, $P^* \in \mathcal{P}$.*

Policy estimation via log-loss Behavior Cloning. We define the log-loss behavioral cloning estimator as:

$$\hat{\pi} = \arg \max_{\pi \in \Pi} \sum_{i \in [n]} \sum_{h \in [H]} \log(\pi_h(a_h^i | s_h^i)). \quad (1)$$

We characterize the estimation error in terms of the Hellinger distance between trajectory distributions in Corollary 20 in Appendix B.3 using concentration results by Foster et al. [2024], cf. Appendix B.1

Transition model estimation via Maximum Likelihood Estimation (MLE). Similarly, we define the MLE transition estimator as:

$$\hat{P} = \arg \max_{P \in \mathcal{P}} \sum_{i \in [n]} \sum_{h \in [H]} \left(\log[P(s_{h+1}^i | s_h^i, a_h^i)] \right). \quad (2)$$

We describe the transition estimation quality in Corollary 24 in Appendix B.3.2, using maximum density likelihood concentration results by Foster et al. [2024]²

Policy confidence set construction. We start by defining the concentrability coefficient, usely defined in offline RL literature [Chen and Jiang 2019].

Definition 2 (Concentrability Coefficient).

$$C(\hat{\pi}, \pi^*) = \sup_{t \in [H]} \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}: d_{P^*}^{\pi^*,t}(s,a) > 0} \frac{d_{P^*}^{\hat{\pi},t}(s,a)}{d_{P^*}^{\pi^*,t}(s,a)}.$$

Intuitively, this coefficient measures how much the state-action visitation distribution of policy $\hat{\pi}$ can deviate from that of the expert policy π^* under the true dynamics P^* . To bound this quantity, we make the following standard assumption [Levine et al. 2020, Chen and Jiang 2019] about minimum state-action visitation:

Assumption 4.2 (Minimum Visitation Probability). *There exists a constant $\gamma_{min} > 0$ such that for all state-action-time tuples with non-zero probability under the optimal policy, that ensures that all relevant state-action pairs have a minimum positive probability under the expert policy.*

$$\min_{(s,a,t): d_{P^*}^{\pi^*,t}(s,a) > 0} d_{P^*}^{\pi^*,t}(s,a) \geq \gamma_{min}. \quad (3)$$

Given this assumption, we can derive a deterministic bound on the concentrability coefficient:

²Note that while we present results specifically for tabular, stochastic and stationary transitions, our framework readily adapts to other transition model classes by deriving appropriate covering number bounds using the general results in Appendix B

203 **Lemma 3** (Concentrability Coefficient Bound). *Consider a policy estimator $\hat{\pi}$ satisfying*

$$H^2(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*}) \leq R.$$

204 *Then, under Assumption 4.2 the concentration coefficient is bounded by*

$$C(\hat{\pi}, \pi^*) \leq 1 + \frac{2\sqrt{R}}{\gamma_{\min}}.$$

205 This bound is key to our approach: it allows us to replace the unknown concentrability coefficient with
 206 a deterministic upper bound that depends only on our policy estimation error and the minimum vis-
 207 itation probability. We can now construct a practical confidence set as shown in the following lemma:

208 **Lemma 4** (Offline Policy Confidence Set). *Assume the following events hold:*

$$\begin{aligned} E_1 &:= \left\{ H^2(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*}) \leq R_1(\delta_1) \right\}, \quad \text{such that } \text{Prob}(E_1) \geq 1 - \delta_1, \\ E_2 &:= \left\{ H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \leq R_2(\delta_2) \right\}, \quad \text{such that } \text{Prob}(E_2) \geq 1 - \delta_2. \end{aligned}$$

209 *Then, under Assumption 4.2 the policy set*

$$\Pi_{1-\delta}^{\text{offline}} := \left\{ \pi : \sqrt{H^2(\mathbb{P}_{\hat{P}}^{\pi}, \mathbb{P}_{\hat{P}}^{\hat{\pi}})} \leq \sqrt{R_1} + \sqrt{R_2} \cdot \left(1 + \sqrt{\left(1 + \frac{2\sqrt{R_1}}{\gamma_{\min}} \right) \cdot H} \right) \right\}$$

210 *is a confidence set of level $1 - \delta = 1 - (\delta_1 + \delta_2)$, i.e.,*

$$\text{Prob}(\pi^* \in \Pi_{1-\delta}^{\text{offline}}) \geq 1 - (\delta_1 + \delta_2).$$

211 The key insight is that the concentrability coefficient now appears as a deterministic term in the
 212 confidence set radius, allowing us to construct a practical confidence region using only quantities that
 213 can be computed from offline data, along with our domain knowledge about the minimum visitation
 214 probability. This confidence set will be central to our online learning phase, providing a principled
 215 way to constrain the policy search space. By combining our tabular setting results from Corollaries 20
 216 and 24 with the concentrability coefficient bound under Assumption 4.2 we can derive an explicit
 217 formula for the confidence set radius:

218 **Theorem 5** (Offline Confidence Set Radius). *Under the setting described above and Assumption 4.2*
 219 *with $\delta_1 = \delta_2 = \delta/2$ and defining*

$$\begin{aligned} \alpha &:= \sqrt{4 \cdot |\mathcal{S}| \cdot \log(|\mathcal{A}| \cdot 2/\delta)}, \\ \beta &:= \sqrt{4 \cdot |\mathcal{S}|^2 \cdot |\mathcal{A}| \cdot \log(nH \cdot 2/\delta)}. \end{aligned}$$

220 *The policy set*

$$\Pi_{1-\delta}^{\text{offline}} := \left\{ \pi : \sqrt{H^2(\mathbb{P}_{\hat{P}}^{\pi}, \mathbb{P}_{\hat{P}}^{\hat{\pi}})} \leq \text{Radius} \right\},$$

221 *is a confidence set of level $1 - \delta$ containing π^* with probability at least $1 - \delta$, where*

$$\text{Radius} = \frac{\alpha}{\sqrt{n}} + \frac{\beta}{\sqrt{n}} \cdot \left(1 + \sqrt{H \cdot \left(1 + \frac{2\alpha}{\gamma_{\min} \cdot \sqrt{n}} \right)} \right).$$

222 This result provides the fundamental connection between offline data and online learning efficiency:
 223 the confidence set radius scales as $O(1/\sqrt{n})$ with the offline sample size n . This inverse square root
 224 dependence means that as we collect more offline expert demonstrations, the confidence set shrinks,
 225 constraining the online policy search space more tightly. Since our online regret bounds will directly
 226 depend on the size of this confidence set, this establishes a quantifiable trade-off between offline data
 227 collection and online preference query efficiency, a key contribution of our work.

4.2 Online preference-based learning

In this section, we present how BRIDGE uses behavioral cloning estimation to guide online preference-based learning. Although we adapted the algorithm from Saha et al. [2023], our offline-to-online approach can be applied to other preference-based RL algorithms beyond BRIDGE.

Our online preference learning follows Saha et al. [2023], who adapted generalized linear models from parametric bandits [Filippi et al., 2010, Fauray et al., 2020] to the preference-based RL setting. As in Saha et al. [2023] we compute a regularized maximum likelihood estimator $\mathbf{w}_t^{MLE}$ to learn the preference weight vector w^* from pairwise comparisons. However, since $\mathbf{w}_t^{MLE}$ may not satisfy assumption 3.2, as in Saha et al. [2023] we define the *data matrix* $\mathbf{V}_t$, which approximates the Fisher Information Matrix (the negative expected Hessian of the log-likelihood):

$$\mathbf{V}_t = \kappa\lambda\mathbf{I}_d + \sum_{\ell=1}^{t-1} (\phi(\tau_\ell^1) - \phi(\tau_\ell^2))^{\otimes 2}, g_t(\mathbf{w}) = \sum_{l=1}^{t-1} \sigma(\langle \phi(\tau_l^1) - \phi(\tau_l^2), \mathbf{w} \rangle) (\phi(\tau_l^1) - \phi(\tau_l^2)) + \lambda\mathbf{w}.$$

Then, the projected parameter, a constrained version of $\mathbf{w}_t^{MLE}$, is given by

$$\mathbf{w}_t^{proj} = \arg \min_{\mathbf{w} \in \mathcal{W}} \|g(\mathbf{w}) - g(\mathbf{w}_t^{MLE})\|_{\mathbf{V}_t}. \quad (4)$$

This matrix $\mathbf{V}_t$ serves dual purposes: defining a confidence ellipsoid $C_t(\delta) = \{\mathbf{w} : \|\mathbf{w} - \mathbf{w}_t^{proj}\|_{\mathbf{V}_t} \leq 2\kappa\beta_t(\delta)\}$ containing w^* with high probability, shaped by the likelihood curvature, and guiding exploration by quantifying uncertainty through $\|\cdot\|_{\mathbf{V}_t^{-1}}$, prioritizing directions with sparse information.

This approach can be further strengthened by relating the empirical norm $\|\cdot\|_{\mathbf{V}_t}$ to an expected norm $\|\cdot\|_{\bar{\mathbf{V}}_t}$, where $\bar{\mathbf{V}}_t = \kappa\lambda\mathbf{I}_d + \sum_{\ell=1}^{t-1} (\phi^{\hat{P}_t}(\pi_\ell^1) - \phi^{\hat{P}_t}(\pi_\ell^2))^{\otimes 2}$. It can be shown that $\|\cdot\|_{\mathbf{V}_t}$ is approximately equivalent to $\|\cdot\|_{\bar{\mathbf{V}}_t}$ up to terms depending on the confidence bonus.

Our key contribution is connecting the offline estimation with the online learning process through two mechanisms: (i) leveraging the offline data to make a first estimation of the transition model, and (ii) constructing an initial dataset of policies that are plausible with the expert data. We refer the reader to (Appendix C) for a detailed explanation of the likelihood-based mechanisms inspired by online binary bandits that underpin this approach.

Transition Model Integration. We leverage our offline MLE transition estimate as the initialization for online learning. In the tabular setting, the MLE for transition probabilities from Eq. (2) is equivalent to a count-based estimator. As we collect online data, we update this estimator to combine both offline and online counts:

$$\hat{P}_t(s'|s, a) = \frac{N_{\text{off}}(s', s, a) + N_t(s', s, a)}{N_{\text{off}}(s, a) + N_t(s, a)},$$

where $N_{\text{off}}(s', s, a)$ counts transitions from state-action pair (s, a) to state s' in the offline dataset, $N_{\text{off}}(s, a)$ counts visits to (s, a) , and $N_t(s', s, a)$ and $N_t(s, a)$ are the corresponding counts from the first t rounds of online interaction. We define our adapted bonus incorporating both offline and online data as:

$$\hat{B}_t(\pi, \eta, \delta) = \mathbb{E}_{\tau \sim \mathbb{P}_{\hat{P}_t}^\pi} \left[\sum_{h \in [H]} \min \left(2\eta, 4\eta \sqrt{\frac{U_h}{N_{\text{off}}(s_h, a_h) + N_t(s_h, a_h)}} \right) \right],$$

where $U_h = H \log(|\mathcal{S}||\mathcal{A}|) + \log \left(\frac{6 \log(N_{\text{off}}(s_h, a_h) + N_t(s_h, a_h))}{\delta} \right)$. This adapts the bonus structure from Chatterji et al. [2021] to leverage our combined offline-online transition estimator.

Remark 6. This integration approach, while effective, has potential for further improvement. First, it does not fully leverage the independence structure of the offline dataset, which could lead to tighter concentration bounds. Second, potential distribution shifts between offline and online phases are not explicitly modeled. These refinements represent promising directions for future research, though our primary focus remains on policy fine-tuning rather than optimal transition modeling.

Feature Moment Bounds. We constrain policy selection to our offline-derived confidence set $\Pi_{1-\delta}^{\text{offline}}(\Pi)$. This enables us to bound the difference between expected feature representations of policies, which directly impacts the uncertainty quantification in our online learning phase.

Algorithm 1 BRIDGE: Bounded Regret with Imitation Data and Guided Exploration

```

1: Input:  $\mathbb{D}_n^H, T$ 
2: Estimate  $\hat{P}$  via MLE (Equation (2)) and compute confidence set  $\Pi_{1-\delta}^{\text{offline}}$  (Theorem 5)
3: Initialize  $\hat{P}_1 \leftarrow \hat{P}, \bar{\mathbf{V}}_1 \leftarrow \lambda \mathbf{I}_d$  ▷ Initialize model and data matrix
4: for  $t = 1, \dots, T$  do
5:   Compute  $w_t^{\text{proj}}$  via constrained MLE (Equation (4))
6:   Define policy set  $\Pi_t$  based on  $\Pi_{1-\delta}^{\text{offline}}$  and  $w_t^{\text{proj}}$  (Lemma 8)
7:    $(\pi_t^1, \pi_t^2) \leftarrow \arg \max_{\pi^1, \pi^2 \in \Pi_t} \{ \gamma_t \cdot \|\phi^{\hat{P}_t}(\pi^1) - \phi^{\hat{P}_t}(\pi^2)\|_{\bar{\mathbf{V}}_t^{-1}} + \hat{B}_t(\pi^1, 2WB, \delta') + \hat{B}_t(\pi^2, 2WB, \delta') \}$ 
8:   Sample trajectories  $\tau_t^1 \sim \mathbb{P}_{\hat{P}_t}^{\pi_t^1}, \tau_t^2 \sim \mathbb{P}_{\hat{P}_t}^{\pi_t^2}$  and obtain preference  $o_t = \mathbb{I}(\tau_t^1 \succ \tau_t^2)$ 
9:   Update matrix:  $\bar{\mathbf{V}}_{t+1} \leftarrow \bar{\mathbf{V}}_t + (\phi^{\hat{P}_t}(\pi_t^1) - \phi^{\hat{P}_t}(\pi_t^2))^{\otimes 2}$  and model  $\hat{P}_{t+1}$ 
10: end for
11: return Best policy from  $\Pi_T$  using final weight estimate  $w_T^{\text{proj}}$ 

```

Lemma 7 (Feature Moment Bounds). *Let X be a random variable on measurable space $(\mathcal{X}, \tilde{\mathcal{X}})$ and $f : \mathcal{X} \rightarrow \mathbb{R}^d$ be a bounded function such that $\|f(x)\|_2 \leq B < \infty$ for all $x \in \mathcal{X}$. For probability distributions P, Q that are absolutely continuous with respect to the Lebesgue measure, if $H^2(P, Q) \leq R$, then*

$$\|\mathbb{E}_{X \sim P}[f(X)] - \mathbb{E}_{X \sim Q}[f(X)]\|_2 \leq 2\sqrt{2} \cdot B \cdot \sqrt{R}.$$

This lemma provides a crucial connection between the Hellinger distance of trajectory distributions and the distance between their expected feature representations (embeddings) by setting $f = \phi$. In our preference-based learning framework, we apply this result to bound the elements of the expected feature covariance matrix $\bar{\mathbf{V}}_t$. The trace $\text{tr}(\bar{\mathbf{V}}_t) = \sum_{t' < t} \|\phi(\pi_{t'}^1) - \phi(\pi_{t'}^2)\|_2^2$ represents total variance, which Lemma 7 controls via our confidence set construction.

At the start of our online learning process ($t = 0$), for policies π_1^0 and π_2^0 that belong to our offline-derived confidence set $\Pi_{1-\delta}^{\text{offline}}$ from Theorem 5, the Hellinger distance between their induced trajectory distributions is bounded by the confidence set radius. Applying Lemma 7 with $f = \phi$, our trajectory embedding function, we obtain:

$$\|\phi^{\hat{P}_0}(\pi_1^0) - \phi^{\hat{P}_0}(\pi_2^0)\|_2 \leq 4\sqrt{2} \cdot B \cdot \mathcal{O}\left(\frac{1}{\sqrt{n}}\right).$$

This result has profound implications for our regret analysis. By constraining policies to our confidence set, we effectively control the variance of feature differences, allowing us to replace the naive bound $\|\phi^{\hat{P}_t}(\pi_1^t) - \phi^{\hat{P}_t}(\pi_2^t)\|_2 \leq 2B$ with our tighter bound that decreases with the offline sample size n . As online learning progresses, the transition model $\hat{P}_t$ improves through additional data collection. Importantly, this improvement in the transition model only strengthens our bound. The exact form of this improvement depends on the concentration properties of the online estimator and will be formalized in our final regret proof.

Based on these bounds, we can now define a policy confidence set that contains the optimal policy π^* with high probability while accounting for both estimation uncertainty and exploration bonuses:

Lemma 8 (Online Policy Confidence Set). *Let Π_t be the set of policies defined as*

$$\begin{aligned} \Pi_t := \{ \pi \in \Pi_{1-\delta}^{\text{offline}} \mid \forall \pi' \in \Pi_{1-\delta}^{\text{offline}} : \\ \langle \phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_t}(\pi'), w_t^{\text{proj}} \rangle + \gamma_t \cdot \|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_t}(\pi')\|_{\bar{\mathbf{V}}_t^{-1}} \\ + \hat{B}_t(\pi, 2WB, \delta') + \hat{B}_t(\pi', 2WB, \delta') \geq 0 \}, \end{aligned}$$

where $\delta' = \frac{\delta}{8\ell^3|\mathcal{A}||\mathcal{S}|}$. With probability at least $1 - \delta$, the optimal policy π^* remains in Π_t for all $t \in [T]$, where δ has been scaled appropriately via union bound to account for the separate probabilistic events in the offline confidence set, transition model estimation, and parameter estimation components. The confidence radius multiplier γ_t is provided in the appendix.

The complete algorithm is described in Algorithm 1.

296 4.3 Theoretical guarantees

297 We can now state the following regret bound for BRIDGE. The proof is shown in Appendix D.
 298 **Theorem 9** (Regret Bound with Offline-Enhanced Exploration). *Let n be the number of offline*
 299 *demonstrations with minimum visitation probability $\gamma_{\min} > 0$ for state-action pairs. With probability*
 300 *at least $1 - \delta$, the regret of BRIDGE is bounded by*

$$\begin{aligned}
 R_T \leq & \underbrace{2 \cdot \underbrace{\gamma_T}_{\text{Term 1}} \cdot \sqrt{T \cdot \log \left(1 + \underbrace{\frac{\tilde{O} \left(B^2 \cdot H \cdot |S|^2 \cdot \min \left\{ \frac{T}{n}, \ln(T) \right\} + \frac{T \cdot |S| \cdot B^2 \cdot \sqrt{|A| \cdot H}}{\sqrt{n \cdot \gamma_{\min}}}}_d \right)}_{\text{Term 2}} \right)}}_{\text{Term 2}} \\
 & + \underbrace{\tilde{O} \left(H|S| \sqrt{\frac{|A|TH}{n \cdot \gamma_{\min}}} + \frac{H^{5/2}WB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2WB \cdot \sqrt{T} \cdot \frac{|S|^{1/2}|A|^{1/4}}{n^{1/4}} \right)}_{\text{Term 3}},
 \end{aligned}$$

301 where

$$\gamma_T = \tilde{O} \left((\kappa + BW) \sqrt{d \log(T)} + H^2WB|S| \cdot \sqrt{\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} + \sqrt{HWPB} \right),$$

302 and we have set $\epsilon = \frac{1}{T}$ to optimize the bound.

303 From this regret bound we can observe that as $n \rightarrow \infty$ with fixed $\gamma_{\min} > 0$: (i) Term 1 ap-
 304 proaches $\tilde{O}((\kappa + BW) \sqrt{d \log(T)} + \sqrt{HWPB})$; (ii) Term 2, the logarithm, approaches $\log(1) = 0$;
 305 (iii) in Term 3, all components approach zero. The overall regret bound exhibits a $\sqrt{T}$ de-
 306 pendence as in Saha et al. [2023]. However, this results in a regret bound that can be made
 307 arbitrarily small with sufficiently high-quality offline data, changing the complexity of regret
 308 analysis without having access to an offline expert dataset. This result helps in closing the
 309 gap between empirical results in applying RL in real-world scenarios and theoretical works.
 310

311 4.4 Numerical simulations

312 We provide initial simulation results based on a practical imple-
 313 mentation of our algorithm. As baselines we used Foster et al.
 314 [2024]’s behavioral cloning (BC) and Saha et al. [2023]’s PbRL
 315 algorithms (for which no implementations are publicly avail-
 316 able). We used the StarMDP and Gridworld environments
 317 described in Pace et al. [2024]. Appendix F gives more details
 318 on our experiments. Figure 2 shows that with as little as 2
 319 offline, optimal trajectories, BRIDGE prunes the policy set and
 320 converges to the optimal policy faster than PbRL.

321 5 Conclusions

322 We present BRIDGE, a novel algorithm for fine-tuning BC
 323 policies using online PbRL. Our approach is motivated by the
 324 practical challenges of deploying RL in real-world settings,
 325 where reward specification is difficult and exploration is risky.
 326 By combining these two feedbacks, we construct confidence
 327 sets that constrain the policy space and guide safe, sample-
 328 efficient learning. We provide the first theoretical regret bound
 329 for this hybrid learning paradigm, showing that offline data
 330 reduces regret through a shrinking uncertainty radius. Our anal-
 331 ysis builds on recent advances in BC and PbRL, and crucially
 332 integrates them into a unified regret framework with provable
 333 benefits. Our work opens new directions for interactive learn-
 334 ing systems that can safely and efficiently improve with human
 335 input, even in the absence of explicit reward signals.

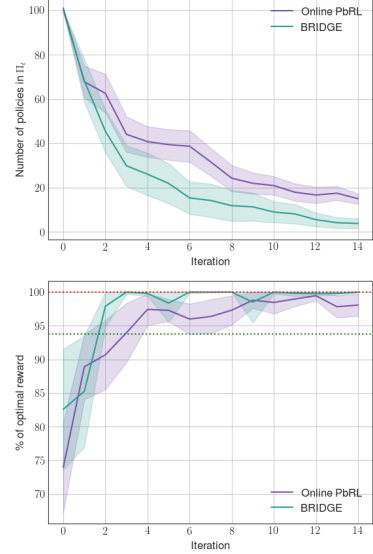


Figure 2: Performance comparison of our new algorithm with offline BC [Foster et al. 2024] and online PbRL [Saha et al. 2023]. The dotted red and green lines are the expected return of the optimal and BC policies respectively.

References

- P. J. Ball, L. Smith, I. Kostrikov, and S. Levine. Efficient online reinforcement learning with offline data. In *International Conference on Machine Learning*, pages 1577–1594. PMLR, 2023.
- D. Brown, R. Coleman, R. Srinivasan, and S. Niekum. Safe imitation learning via fast bayesian reward inference from preferences. In *International Conference on Machine Learning*, pages 1165–1177. PMLR, 2020.
- N. Chatterji, A. Pacchiano, P. Bartlett, and M. Jordan. On the theory of reinforcement learning with once-per-episode feedback. *Advances in Neural Information Processing Systems*, 34:3401–3412, 2021.
- J. Chen and N. Jiang. Information-theoretic considerations in batch reinforcement learning. In *International conference on machine learning*, pages 1042–1051. PMLR, 2019.
- P. F. Christiano, J. Leike, T. B. Brown, M. Martic, S. Legg, and D. Amodei. Deep reinforcement learning from human preferences. In *Proc. of Advances in Neural Information Processing Systems*, volume 2017-Decem, pages 4300–4308, 2017.
- G. Dulac-Arnold, D. Mankowitz, and T. Hester. Challenges of real-world reinforcement learning. *arXiv preprint arXiv:1904.12901*, 2019.
- L. Faury, M. Abeille, C. Calauzènes, and O. Fercoq. Improved optimistic algorithms for logistic bandits. In *International Conference on Machine Learning*, pages 3052–3060. PMLR, 2020.
- S. Filippi, O. Cappe, A. Garivier, and C. Szepesvári. Parametric bandits: The generalized linear case. *Advances in neural information processing systems*, 23, 2010.
- D. J. Foster, A. Block, and D. Misra. Is behavior cloning all you need? understanding horizon in imitation learning. *arXiv preprint arXiv:2407.15007*, 2024.
- T. Hester, M. Vecerik, O. Pietquin, M. Lanctot, T. Schaul, B. Piot, D. Horgan, J. Quan, A. Sendonaris, I. Osband, et al. Deep q-learning from demonstrations. In *Proceedings of the AAAI conference on artificial intelligence*, volume 32, 2018.
- I. Kostrikov, O. Nachum, V. Tomar, and S. Levine. Offline reinforcement learning with implicit q-learning. In *International Conference on Learning Representations (ICLR)*, 2022.
- T. Lattimore and C. Szepesvári. *Bandit algorithms*. Cambridge University Press, 2020.
- S. Lee, Y. Seo, K. Lee, P. Abbeel, and J. Shin. Offline-to-online reinforcement learning via balanced replay and pessimistic q-ensemble. In *Conference on Robot Learning*, pages 1702–1712. PMLR, 2022.
- J. Leike, D. Krueger, T. Everitt, M. Martic, V. Maini, and S. Legg. Scalable agent alignment via reward modeling: a research direction. *arXiv preprint arXiv:1811.07871*, 2018.
- S. Levine, A. Kumar, G. Tucker, and J. Fu. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *arXiv preprint arXiv:2005.01643*, 2020.
- A. Nair, B. McGrew, M. Andrychowicz, W. Zaremba, and P. Abbeel. Overcoming exploration in reinforcement learning with demonstrations. In *2018 IEEE international conference on robotics and automation (ICRA)*, pages 6292–6299. IEEE, 2018.
- A. Nair, G. Dalal, A. Gupta, and S. Levine. Awac: Accelerating online reinforcement learning with offline datasets. In *Conference on robot learning*, page 102–121. PMLR, 2020.
- L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray, et al. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744, 2022.

- 379 A. Pacchiano, J. Parker-Holder, Y. Tang, K. Choromanski, A. Choromanska, and M. Jordan. Learning
380 to score behaviors for guided policy optimization. In H. D. III and A. Singh, editors, *Proceedings*
381 *of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine*
382 *Learning Research*, pages 7445–7454. PMLR, 13–18 Jul 2020. URL [https://proceedings](https://proceedings.mlr.press/v119/pacchiano20a.html)
383 [mlr.press/v119/pacchiano20a.html](https://proceedings.mlr.press/v119/pacchiano20a.html).
- 384 A. Pace, B. Schölkopf, G. Rätsch, and G. Ramponi. Preference elicitation for offline reinforcement
385 learning. *arXiv preprint arXiv:2406.18450*, 2024.
- 386 S. Park, K. Frans, S. Levine, and A. Kumar. Is value learning really the main bottleneck in offline rl?
387 In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.
- 388 J. Parker-Holder, A. Pacchiano, K. Choromanski, and S. Roberts. Effective diversity in population
389 based reinforcement learning. In *Proceedings of the 34th International Conference on Neural*
390 *Information Processing Systems, NIPS ’20*, Red Hook, NY, USA, 2020a. Curran Associates Inc.
391 ISBN 9781713829546.
- 392 J. Parker-Holder, A. Pacchiano, K. M. Choromanski, and S. J. Roberts. Effective diversity in
393 population based reinforcement learning. *Advances in Neural Information Processing Systems*, 33:
394 18050–18062, 2020b.
- 395 D. A. Pomerleau. Alvin: An autonomous land vehicle in a neural network. *Advances in neural*
396 *information processing systems*, 1, 1988.
- 397 A. Rajeswaran, V. Kumar, A. Gupta, G. Vezzani, J. Schulman, E. Todorov, and S. Levine. Learning
398 complex dexterous manipulation with deep reinforcement learning and demonstrations. *arXiv*
399 *preprint arXiv:1709.10087*, 2017.
- 400 S. Ross and J. A. Bagnell. Agnostic system identification for model-based reinforcement learning.
401 In *Proceedings of the 29th International Conference on Machine Learning*, pages 1905–1912, 2012.
- 402 S. Ross, G. Gordon, and D. Bagnell. A reduction of imitation learning and structured prediction to
403 no-regret online learning. In *Proceedings of the fourteenth international conference on artificial*
404 *intelligence and statistics*, pages 627–635. JMLR Workshop and Conference Proceedings, 2011.
- 405 A. Saha, A. Pacchiano, and J. Lee. Dueling rl: Reinforcement learning with trajectory preferences. In
406 *International Conference on Artificial Intelligence and Statistics*, pages 6263–6289. PMLR, 2023.
- 407 I. Sason and S. Verdú. f -divergence inequalities. *IEEE Transactions on Information Theory*, 62(11):
408 5973–6006, 2016.
- 409 Y. Song, Y. Zhou, A. Sekhari, D. Bagnell, A. Krishnamurthy, and W. Sun. Hybrid RL: Using both
410 offline and online data can make RL efficient. In *The Eleventh International Conference on*
411 *Learning Representations*, 2023. URL <https://openreview.net/forum?id=yyBis80iUuU>.
- 412 R. S. Sutton and A. G. Barto. *Reinforcement learning: An introduction*. MIT press, 2018.
- 413 C. Tang, B. Abbatematteo, J. Hu, R. Chandra, R. Martín-Martín, and P. Stone. Deep reinforcement
414 learning for robotics: A survey of real-world successes. In *Proceedings of the AAAI Conference on*
415 *Artificial Intelligence*, volume 39, pages 28694–28698, 2025.
- 416 D. Tang, R. Jain, B. Hao, and Z. Wen. Efficient online learning with offline datasets for infinite
417 horizon mdps: A bayesian approach. *arXiv preprint arXiv:2310.11531*, 2023.
- 418 A. Tirinzoni, A. Touati, J. Farebrother, M. Guzek, A. Kanervisto, Y. Xu, A. Lazaric, and M. Pirota.
419 Zero-shot whole-body humanoid control via behavioral foundation models. *arXiv preprint*
420 *arXiv:2504.11054*, 2025.
- 421 S. A. van de Geer. Applications of empirical process theory. 2000. URL [https://api](https://api.semanticscholar.org/CorpusID:123051755)
422 [semanticscholar.org/CorpusID:123051755](https://api.semanticscholar.org/CorpusID:123051755).
- 423 M. Vecerik, T. Hester, J. Scholz, F. Wang, O. Pietquin, B. Piot, N. Heess, T. Rothörl, T. Lampe, and
424 M. Riedmiller. Leveraging demonstrations for deep reinforcement learning on robotics problems
425 with sparse rewards. *arXiv preprint arXiv:1707.08817*, 2017.
- 426

- 427 A. Wagenmaker and A. Pacchiano. Leveraging offline data in online reinforcement learning. In
428 *International Conference on Machine Learning*, pages 35300–35338. PMLR, 2023.
- 429 M. J. Wainwright. *High-dimensional statistics: A non-asymptotic viewpoint*, volume 48. Cambridge
430 university press, 2019.
- 431 T. Xie, N. Jiang, H. Wang, C. Xiong, and Y. Bai. Policy finetuning: Bridging sample-efficient
432 offline and online reinforcement learning. *Advances in neural information processing systems*, 34:
433 27395–27407, 2021.
- 434 T. Zhang. From ϵ -entropy to kl-entropy: Analysis of minimum information complexity density
435 estimation. *Annals of Statistics*, 34:2180–2210, 2006. URL [https://api.semanticscholar](https://api.semanticscholar.org/CorpusID:14152942)
436 [org/CorpusID:14152942](https://api.semanticscholar.org/CorpusID:14152942).

437

Appendix

438

439

440

Table of Contents

441

A Simplified Setup for Understanding Regret Analysis **14**

442

A.1 Setup for Known Dynamics 15

443

A.2 Regret Analysis: BRIDGE (Known Model) 16

444

A.3 Practical Regret Analysis with Fixed Offline Data 19

445

B Offline estimation **20**

446

B.1 Maximum Likelihood for Density Estimation 20

447

B.2 MLE Objective of Dataset of Independent Trajectories 20

448

B.3 Concentration Bounds 21

449

B.4 Performance Guarantees 29

450

C Online Estimation **30**

451

C.1 Elliptical Confidence Set 30

452

C.2 Norm Relation Between Data Matrices 31

453

C.3 Transition Estimation and Bonus Terms 32

454

C.4 Policy Set Π_t and Proof Lemma 8 34

455

C.5 Regret Bound 35

456

D Regret Analysis: Theorem 9 **37**

457

D.1 Term 1: Asymptotic bound 38

458

D.2 Term 2: Asymptotic bound 45

459

D.3 Term 3: Asymptotic bound 50

460

E Auxiliary Mathematical Results **61**

461

E.1 Bridging Offline Confidence Sets and Online Constraints 61

462

F Experiments **63**

463

F.1 Environments 63

464

F.2 Additional result on Gridworld 64

465

F.3 Embeddings 64

466

467

A Simplified Setup for Understanding Regret Analysis

In this section, we propose an analysis of the regret under a simplified setting: The underlying dynamic P^* is known. This idea is to help the reader understand how the construction of the confidence set over the policies from the offline learning estimation helps to reduce the number of policies to draw from in the online learning setting, without being overwhelmed by the estimation of the transitions. In this setting, it is then clear what part of our methods applies to the policies. The goal is to prepare the reader for the proof of our algorithm BRIDGE in Appendix D

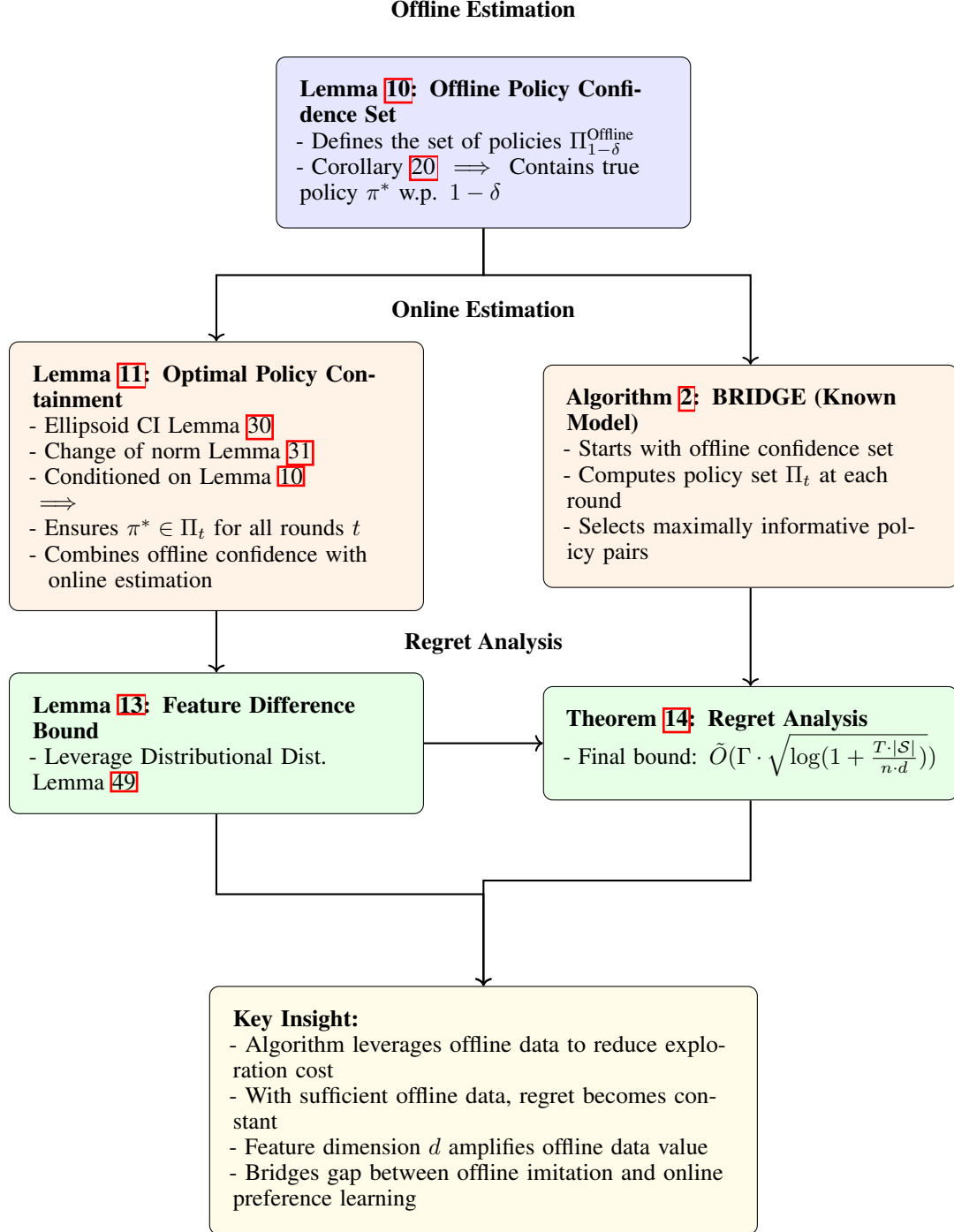


Figure 3: Proof Overview for BRIDGE Algorithm with Known Dynamics

476 A.1 Setup for Known Dynamics

477 A.1.1 Offline Estimation with Known Dynamics

478 Assume we get the offline data $\mathbb{D}_n^H = \{\tau_i\}_{i \in [n]}$. The underlying object describing the trajectories
 479 is a Finite MDP Reward Free setting as in the main paper. Assume that the set of possible policies
 480 is stationary and deterministic. Then under the fact that the underlying dynamic is known, the
 481 confidence set from Theorem 5 reduces to the following, by direct application of Corollary 20 i.e.,
 482 setting the radius around the MLE estimate π_{MLE} from Equation 7.

483 We formalize this into the following lemma:

484 **Lemma 10** (Offline Policy Confidence Set under Known Dynamics). *Let $\hat{\pi}$ be the log-loss BC*
 485 *estimator defined in Equation 7.*
 486 *The policy set*

$$\Pi_{1-\delta}^{Offline} := \left\{ \pi : H(\mathbb{P}_{P^*}^{\pi}, \mathbb{P}_{P^*}^{\hat{\pi}}) \leq \sqrt{\frac{6 \cdot |\mathcal{S}| \cdot \log(|\mathcal{A}| \cdot \delta^{-1})}{n}} \right\}$$

487 contains π^* with probability at least $1 - \delta$.

488 *Proof.* Note that by symmetry

$$H(\mathbb{P}_{P^*}^{\pi^*}, \mathbb{P}_{P^*}^{\hat{\pi}}) = H(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*})$$

489 Then the result follows from Corollary 20. □

490 A.1.2 Online Learning with Known Dynamics

491 Here we adapt our algorithm BRIDGE to the setting with known dynamics. This means we adapt
 492 the approach from Saha et al. (2023) under known dynamics to constrain the set of policies to choose
 493 from to our confidence set described in the previous section.

494 First, since the transitions are known, we define for this section:

$$\phi^{P^*}(\pi) := \phi(\pi) = \mathbb{E}_{\tau \sim \mathbb{P}_{P^*}^{\pi}} [\phi(\tau)]$$

495 We also define the expected data matrix $\bar{V}_t^{P^*}$ under the true transition dynamics P^* as follows (see
 496 Appendix C for an overview of results about data matrices):

$$\bar{V}_t^{P^*} = \kappa \lambda \mathbf{I}_d + \sum_{\ell=1}^{t-1} (\phi(\pi_\ell^1) - \phi(\pi_\ell^2)) (\phi(\pi_\ell^1) - \phi(\pi_\ell^2))^\top$$

497 Then we define the set of policies to draw from as:

$$\Pi_t := \left\{ \pi \in \Pi_{1-\delta}^{Offline} \mid \forall \pi' \in \Pi_{1-\delta}^{Offline} : \right. \\ \left. \langle \phi(\pi) - \phi(\pi'), \mathbf{w}_t^{\text{proj}} \rangle + \gamma_t \cdot \|\phi(\pi) - \phi(\pi')\|_{(\bar{V}_t^{P^*})^{-1}} \geq 0 \right\}$$

498 where $\gamma_t := 2\kappa\beta_t(\delta) + \alpha_{d,T}(\delta)$ and $\alpha_{d,T}(\delta)$ is defined as in Lemma 31.

499 **Lemma 11** (Optimal Policy Containment). *Conditioned on $E_{w^*} \cap E_{\bar{V}_T^{P^*}} \cap E_{offline}$ where:*

- 500 • E_{w^*} is the event defined in Lemma 30
- 501 • $E_{\bar{V}_T^{P^*}}$ is the event defined in Lemma 31
- 502 • $E_{offline} := \{\pi^* \in \Pi_{1-\delta}^{Offline}\}$

503 then

$$\pi^* \in \Pi_t \quad \forall t \in [T]$$

504 *Proof.* This follows directly from Lemma 2 in Saha et al. [2023]. We adapt the probability parameter
 505 δ to account for the additional condition that $\pi^* \in \Pi_{1-\delta}^{\text{Offline}}$, which holds with probability at least $1 - \delta$
 506 according to Lemma 28. $\square$

507 We now present the adapted version of BRIDGE for the known transition model:

Algorithm 2 BRIDGE: Bounded Regret with Imitation Data and Guided Exploration (Known Model)

1: **Input:** Offline dataset $\mathbb{D}_n^H$, time horizon T , true dynamics P^*
 2: Compute confidence set $\Pi_{1-\delta}^{\text{Offline}}$ using Lemma 10
 3: Initialize $\bar{\mathbf{V}}_1^{P^*} \leftarrow \kappa \lambda \mathbf{I}_d$ $\triangleright$ Initialize data matrix
 4: **for** $t = 1, \dots, T$ **do**
 5: Compute $\mathbf{w}_t^{\text{proj}}$ via constrained MLE (Equation (4))
 6: Define policy set Π_t based on $\Pi_{1-\delta}^{\text{Offline}}$ and $\mathbf{w}_t^{\text{proj}}$
 7: $(\pi_t^1, \pi_t^2) \leftarrow \arg \max_{\pi^1, \pi^2 \in \Pi_t} \{\|\phi(\pi^1) - \phi(\pi^2)\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}}\}$
 8: Sample trajectories $\tau_t^1 \sim \mathbb{P}_{P^*}^{\pi_t^1}$, $\tau_t^2 \sim \mathbb{P}_{P^*}^{\pi_t^2}$ and obtain preference $o_t = \mathbb{I}(\tau_t^1 \succ \tau_t^2)$
 9: Update matrix: $\bar{\mathbf{V}}_{t+1}^{P^*} \leftarrow \bar{\mathbf{V}}_t^{P^*} + (\phi(\pi_t^1) - \phi(\pi_t^2))(\phi(\pi_t^1) - \phi(\pi_t^2))^\top$
 10: **end for**
 11: **return** Best policy from Π_T using final weight estimate $\mathbf{w}_T^{\text{proj}}$

508 A.2 Regret Analysis: BRIDGE (Known Model)

509 We now present a regret analysis of the BRIDGE algorithm under known transition. We start by
 510 stating the following lemma:

511 **Lemma 12.** *The regret of BRIDGE under known dynamic is upper bounded as follow:*

$$R_T \leq 4 \cdot (\beta_T(\delta) + \alpha_{T,d}(\delta)) \cdot \sum_{t \in [T]} \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}}$$

512 *Proof.* For ease of notation we define $\delta^{*,1}\phi := \phi(\pi^*) - \phi(\pi^1)$. First we bound the instantenous
 513 regret

$$\begin{aligned} 2r_t &= \langle \delta^{*,1}\phi, w^* \rangle + \langle \delta^{*,2}\phi, w^* \rangle \\ &\leq \langle \delta^{*,1}\phi, \mathbf{w}_t^{\text{proj}} \rangle + \langle \delta^{*,2}\phi, \mathbf{w}_t^{\text{proj}} \rangle + \|w^* - \mathbf{w}_t^{\text{proj}}\|_{\bar{\mathbf{V}}_t^{P^*}} \left(\|\delta^{*,1}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} + \|\delta^{*,2}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} \right) \\ &\stackrel{\text{Lemma 31}}{\leq} \underbrace{\langle \delta^{*,1}\phi, \mathbf{w}_t^{\text{proj}} \rangle + \langle \delta^{*,2}\phi, \mathbf{w}_t^{\text{proj}} \rangle}_{\text{Lemma 31}} + (2\kappa\beta_t(\delta) + \alpha_{T,d}(\delta)) \cdot \left(\|\delta^{*,1}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} + \|\delta^{*,2}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} \right) \end{aligned}$$

514 Then notice that choosing π_t^1, π_t^2 as $\arg \max$ together with the fact that $\pi^* \in \Pi_t$ Lemma 11 yield

$$2r_t \leq \langle \delta^{*,1}\phi, \mathbf{w}_t^{\text{proj}} \rangle + \langle \delta^{*,2}\phi, \mathbf{w}_t^{\text{proj}} \rangle + (2\kappa\beta_t(\delta) + \alpha_{T,d}(\delta)) \cdot \left(\|\delta^{*,1}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} + \|\delta^{*,2}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} \right)$$

515 Next using the fact that $\pi_t^1, \pi_t^2, \pi^* \in \Pi_t$ we have the following constraints

$$\begin{aligned} \langle \delta^{*,i}\phi, \mathbf{w}_t^{\text{proj}} \rangle + \gamma_t \|\delta^{*,i}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} &\geq 0 \quad i \in \{1, 2\} \\ \Leftrightarrow \langle \delta^{*,i}\phi, \mathbf{w}_t^{\text{proj}} \rangle &\leq \gamma_t \|\delta^{*,i}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} \quad i \in \{1, 2\} \end{aligned}$$

516 yielding

$$2r_t \leq (2\kappa\beta_t(\delta) + \alpha_{T,d}(\delta)) \|\delta^{*,1}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}} \leq 4(2\kappa\beta_T(\delta) + \alpha_{T,d}(\delta)) \cdot \|\delta^{*,1}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}}$$

517 Hence

$$R_T = \sum_{t \in [T]} r_t \leq 2 \cdot \beta_T(\delta) + \alpha_{T,d}(\delta) \cdot \sum_{t \in [T]} \|\delta^{*,1}\phi\|_{(\bar{\mathbf{V}}_t^{P^*})^{-1}}$$

518 $\square$

519 The remaining step in our analysis is to bound the term:

$$\sum_{t \in [T]} \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_{(\bar{V}_t^{P^*})^{-1}} = \sum_{t \in [T]} \left\| \mathbb{E}_{\tau \sim \mathbb{P}_{P^*}^{\pi_t^1}} [\phi(\tau)] - \mathbb{E}_{\tau \sim \mathbb{P}_{P^*}^{\pi_t^2}} [\phi(\tau)] \right\|_{(\bar{V}_t^{P^*})^{-1}}$$

520 A simple approach would be to use Assumption 3.2 which states that the feature map ϕ is bounded
 521 in ℓ_2 -norm by B . However, our offline confidence set construction in Lemma 10 provides a more
 522 powerful result: policies in our set have distributions that are close not only in Hellinger distance but
 523 also in the resulting feature expectations.

524 This is precisely why we formulated our confidence set constraint using the square root of the squared
 525 Hellinger distance - it yields a bound on the L_2 norm of distribution differences. Through Lemma 49,
 526 we can translate bounds on Hellinger distance into bounds on the difference of feature expectations
 527 in the ℓ_2 -norm.

528 We formalize this connection in the following lemma:

529 **Lemma 13** (Feature Difference Bound Under Offline Constraints - Corrected). *For policies $\pi_t^1, \pi_t^2 \in$
 530 $\Pi_{1-\delta}^{\text{Offline}}$ selected by our algorithm at each round $t \in [T]$, the sum of feature differences measured in
 531 the data matrix norm is bounded as:*

$$\sum_{t \in [T]} \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_{(\bar{V}_t^{P^*})^{-1}} \leq \sqrt{2d \cdot \log \left(1 + \frac{192B^2T|\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n \cdot d \cdot \lambda} \right)}$$

532 where d is the feature dimension, B is the feature norm bound, $|\mathcal{S}|$ and $|\mathcal{A}|$ are the state and action
 533 space sizes, and n is the number of offline samples.

534 *Proof.* First note

$$\sum_{t \in [T]} \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_{(\bar{V}_t^{P^*})^{-1}} \leq \sqrt{T \cdot \sum_{t \in [T]} \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_{(\bar{V}_t^{P^*})^{-1}}^2}$$

535 then notice the inequality

$$u \leq 2 \log(1 + u) \quad u \geq 1 \implies \sum_{t \in [T]} \|\delta^{1,2} \phi\|_{(\bar{V}_t^{P^*})^{-1}}^2 \leq 2 \cdot \sum_{t \in [T]} \log(1 + \|\delta^{1,2} \phi\|_{(\bar{V}_t^{P^*})^{-1}}^2)$$

536 Using the definition of $\bar{V}_t^{P^*}$, we have

$$\begin{aligned} \bar{V}_{t+1}^{P^*} &= \lambda \cdot I_{d \times d} + \sum_{i \in [t]} (\phi(\pi_i^1) - \phi(\pi_i^2))(\phi(\pi_i^1) - \phi(\pi_i^2))^T \\ &= \bar{V}_t^{P^*} + (\phi(\pi_t^1) - \phi(\pi_t^2))(\phi(\pi_t^1) - \phi(\pi_t^2))^T \\ &= (\bar{V}_t^{P^*})^{1/2} \left(I + (\bar{V}_t^{P^*})^{-1/2} (\phi(\pi_t^1) - \phi(\pi_t^2))(\phi(\pi_t^1) - \phi(\pi_t^2))^T (\bar{V}_t^{P^*})^{-1/2} \right) (\bar{V}_t^{P^*})^{1/2} \end{aligned}$$

537 Using properties of determinant:

$$\begin{aligned} \det(\bar{V}_{t+1}^{P^*}) &= \det(\bar{V}_t^{P^*}) \cdot \det(I + (\bar{V}_t^{P^*})^{-1/2} (\phi(\pi_t^1) - \phi(\pi_t^2))(\phi(\pi_t^1) - \phi(\pi_t^2))^T (\bar{V}_t^{P^*})^{-1/2}) \\ &= \det(\bar{V}_t^{P^*}) \cdot (1 + \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_{(\bar{V}_t^{P^*})^{-1}}^2) \\ &= \det(V_0) \cdot \prod_{s \in [t]} (1 + \|\phi(\pi_s^1) - \phi(\pi_s^2)\|_{(\bar{V}_s^{P^*})^{-1}}^2) \\ &\Leftrightarrow \\ \log \left[\frac{\det(\bar{V}_{t+1}^{P^*})}{\det(V_0)} \right] &= \sum_{s \in [t]} \log(1 + \|\phi(\pi_s^1) - \phi(\pi_s^2)\|_{(\bar{V}_s^{P^*})^{-1}}^2) \end{aligned}$$

538 We have for the determinant:

$$\det(\bar{V}_{t+1}^{P*}) = \prod_{i \in [d]} \lambda_i \leq \left(\frac{1}{d} \cdot \text{Tr}\{\bar{V}_{t+1}^{P*}\} \right)^d$$

539 Using linearity of trace:

$$\begin{aligned} \text{Tr}\{\bar{V}_{t+1}^{P*}\} &= \text{Tr}\{\lambda I\} + \sum_{s \in [t]} \text{Tr}\{(\phi(\pi_s^1) - \phi(\pi_s^2))(\phi(\pi_s^1) - \phi(\pi_s^2))^T\} \\ &= d \cdot \lambda + \sum_{s \in [t]} \|\phi(\pi_s^1) - \phi(\pi_s^2)\|_2^2 \end{aligned}$$

540 Applying the corrected bound from Lemma 49:

$$\begin{aligned} \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_2^2 &\leq (2\sqrt{2} \cdot B \cdot \sqrt{H^2(\mathbb{P}_{P*}^{\pi_t^1}, \mathbb{P}_{P*}^{\pi_t^2})})^2 \\ &\leq 8B^2 \cdot \frac{24 \cdot |\mathcal{S}| \cdot \log(|\mathcal{A}| \cdot \delta^{-1})}{n} \\ &= \frac{192B^2 |\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n} \end{aligned}$$

541 Using this tighter bound in our trace calculation:

$$\begin{aligned} \text{Tr}\{\bar{V}_{t+1}^{P*}\} &\leq d \cdot \lambda + t \cdot \frac{192B^2 |\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n} \\ &= d\lambda \left(1 + \frac{192B^2 t |\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n \cdot d \cdot \lambda} \right) \end{aligned}$$

542 Hence:

$$\begin{aligned} \log \left[\frac{\det(\bar{V}_{t+1}^{P*})}{\det(V_0)} \right] &\leq d \cdot \log \left(\frac{\text{Tr}\{\bar{V}_{t+1}^{P*}\}}{d} \right) \\ &= d \cdot \log \left(\lambda \left(1 + \frac{192B^2 t |\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n \cdot d \cdot \lambda} \right) \right) \\ &= d \cdot \log(\lambda) + d \cdot \log \left(1 + \frac{192B^2 t |\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n \cdot d \cdot \lambda} \right) \end{aligned}$$

543 Since $\det(V_0) = \lambda^d$, the first logarithmic term cancels out:

$$\log \left[\frac{\det(\bar{V}_{t+1}^{P*})}{\det(V_0)} \right] = d \cdot \log \left(1 + \frac{192B^2 t |\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n \cdot d \cdot \lambda} \right)$$

544 Therefore:

$$\begin{aligned} \sum_{t \in [T]} \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_{(\bar{V}_t^{P*})^{-1}}^2 &\leq 2 \cdot \log \left[\frac{\det(\bar{V}_{T+1}^{P*})}{\det(V_0)} \right] \\ &\leq 2d \cdot \log \left(1 + \frac{192B^2 T |\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n \cdot d \cdot \lambda} \right) \end{aligned}$$

545 Taking the square root:

$$\sum_{t \in [T]} \|\phi(\pi_t^1) - \phi(\pi_t^2)\|_{(\bar{V}_t^{P*})^{-1}} \leq \sqrt{2d \cdot \log \left(1 + \frac{192B^2 T |\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n \cdot d \cdot \lambda} \right)}$$

546

□

547 **Theorem 14** (Regret Analysis for BRIDGE under Known Model). *Let $\delta \leq 1/e$ and $\lambda \geq \frac{B}{\kappa}$. Then,*
 548 *with probability at least $1 - \delta$, the expected regret of Algorithm 2 is bounded by:*

$$R_T \leq (2\kappa\beta_T(\delta) + \alpha_{d,T}(\delta)) \sqrt{2d \cdot \log \left(1 + \frac{192B^2T|\mathcal{S}| \log(|\mathcal{A}| \cdot \delta^{-1})}{n \cdot d \cdot \lambda} \right)}$$

549 *In asymptotic notation, this becomes:*

$$R_T = O \left(\left(W\sqrt{\kappa B} + WB \right) d \log(TB/\kappa\delta) \sqrt{\log \left(1 + \frac{T|\mathcal{S}|}{n \cdot d} \right)} \right)$$

550 *where the probability parameter δ accounts for the events*

$$\begin{aligned} E_{w^*} &\rightarrow \text{Lemma 30} \\ E_{\bar{V}_T^{P^*}} &\rightarrow \text{Lemma 31} \\ E_{\text{offline}} := \{\pi^* \in \Pi_{1-\delta}^{\text{Offline}}\} &\rightarrow \text{Lemma 10} \end{aligned}$$

551 **Remark 15.** *This result demonstrates a significant improvement over Saha et al. [2023]'s bound*
 552 *of $O \left(\left(W\sqrt{\kappa B} + WB \right) d \log(TB/\kappa\delta) \sqrt{T} \right)$. The key advantage lies in the term $\sqrt{\log(1 + \frac{T|\mathcal{S}|}{n})}$,*
 553 *which approaches zero as $n \rightarrow \infty$, potentially yielding constant regret.*

554 A.3 Practical Regret Analysis with Fixed Offline Data

555 For a fixed offline dataset of size n , our regret bound scales with horizon T as:

$$R_T = O \left(\Gamma \cdot \sqrt{\log \left(1 + \frac{CT|\mathcal{S}|}{n \cdot d} \right)} \right)$$

556 where $\Gamma = (W\sqrt{\kappa B} + WB)d \log(TB/\kappa\delta)$. This bound reveals three distinct regimes:

557 1. **Small T Regime** ($T|\mathcal{S}| \ll n \cdot d$): Using $\log(1+x) \approx x$ for small x :

$$R_T = O \left(\Gamma \cdot \sqrt{\frac{T|\mathcal{S}|}{n \cdot d}} \right) = O \left(\Gamma \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}}{\sqrt{n \cdot d}} \right)$$

558 2. **Transition Regime** ($T|\mathcal{S}| \approx n \cdot d$):

$$R_T = O(\Gamma) = O \left((W\sqrt{\kappa B} + WB)d \log(TB/\kappa\delta) \right)$$

559 3. **Large T Regime** ($T|\mathcal{S}| \gg n \cdot d$):

$$R_T = O \left(\Gamma \cdot \sqrt{\log(T)} \right)$$

560 These regimes highlight two key insights: (1) with sufficient offline data ($n = \Omega(\frac{T|\mathcal{S}|}{d})$), regret
 561 dramatically improves from $O(\sqrt{\log(T)})$ to $O(1)$ in the dependence on T ; and (2) feature dimension
 562 d amplifies the value of offline data, allowing the same regret reduction with $\sqrt{d}$ times less data. This
 563 explains why high-dimensional problems may benefit more significantly from offline data.

564 As n increases, regret transitions from logarithmic ($O(\log(T))$) to sublinear ($O(\sqrt{T/n})$) and even-
 565 tually approaches $O(1)$ when $n \gg \frac{T|\mathcal{S}|}{d}$. In the limiting case where $n \rightarrow \infty$, exploration becomes
 566 unnecessary, and regret is bounded only by statistical error in the offline estimation.

B Offline estimation

B.1 Maximum Likelihood for Density Estimation

In this section, we present Maximum Likelihood Estimation (MLE) for density estimation that forms the foundation of our concentration results. While these results are presented more extensively in Foster et al. [2024], we include them here for completeness and readability.

The analysis of MLE relies on standard concentration techniques following the well-established work of van de Geer [2000] and Zhang [2006], enhanced by new Freedman-type concentration inequalities developed in Foster et al. [2024] (Appendix B).

The key proof strategy connects MLE analysis to information-theoretic measures via Rényi divergence of order $1/2$. Specifically, the approach bounds expressions of the form $-n \cdot \log(\mathbb{E}_{z \sim g^*}[e^{\frac{1}{2} \log(g(z)/g^*(z))}])$, which equals $\frac{n}{2} \cdot D_{1/2}(g \| g^*)$. This term is bounded using Freedman-type inequalities for adapted sequences, which provide high-probability bounds of the form $\sum_{t=1}^{T'} -\log(\mathbb{E}_{t-1}[e^{-X_t}]) \leq \sum_{t=1}^{T'} X_t + \log(\delta^{-1})$. When combined with union bounds over ε -nets, this yields tight concentration results for the entire function class. The approach also leverages connections to Hellinger distance through the identity $H^2(g, g^*) = 1 - \int \sqrt{g(z)g^*(z)} dz$, providing geometrically interpretable guarantees.

To handle infinite classes, we introduce a tailored notion of covering number for log-loss:

Definition 16 (Log-Covering Number). *For a class $\mathcal{G} \subset \Delta(\mathcal{X})$, the class $\mathcal{G}' \subset \mathcal{X}$ is an ϵ -cover if for all $g \in \mathcal{G}$, there exists $g' \in \mathcal{G}'$ such that $\forall x \in \mathcal{X}$*

$$\log(g(x)/g'(x)) \leq \epsilon$$

The size of such cover is defined by $\mathcal{N}_{\log}(\mathcal{G}, \epsilon)$.

Consider the data $\mathbb{D}_n = \{x_i\}_{i \in [n]}$ consisting of i.i.d copies of $x \sim g^*$ where $g^* \in \Delta(\mathcal{X})$. We have a class $\mathcal{G} \subseteq \Delta(\mathcal{X})$ that may or may not contain g^* . The density MLE estimator is defined as

$$\hat{g} = \arg \max_{g \in \mathcal{G}} \sum_{i \in [n]} \log(g(x_i)) \quad (5)$$

Lemma 17 (Maximum Likelihood Estimator Bound). *The maximum likelihood estimator in Eq. Equation (5) has that with probability at least $1 - \delta$,*

$$H^2(\hat{g}, g^*) \leq \inf_{\varepsilon > 0} \left\{ \frac{6 \log(2\mathcal{N}_{\log}(\mathcal{G}, \varepsilon)/\delta^{-1})}{n} + 4\varepsilon \right\} + 2 \inf_{g \in \mathcal{G}} \log(1 + D_{\chi^2}(g^* \| g))$$

In particular, if $\mathcal{G}$ is finite, the maximum likelihood estimator satisfies

$$H^2(\hat{g}, g^*) \leq \frac{6 \log(2|\mathcal{G}|/\delta^{-1})}{n} + 2 \inf_{g \in \mathcal{G}} \log(1 + D_{\chi^2}(g^* \| g))$$

Note that the term $\inf_{g \in \mathcal{G}} \log(1 + D_{\chi^2}(g^ \| g))$ corresponds to misspecification error, and is zero if $g^* \in \mathcal{G}$.*

B.2 MLE Objective of Dataset of Independent Trajectories

Given a data set of reward free trajectories $\mathbb{D}_n^H = \{\tau_i\}_{i \in [H]}$ of n trajectories of length H where $\{\tau_i\} \sim_{i.i.d} \tau \sim \mathbb{P}_{P^*}^{\pi^*}$. The distribution $\mathbb{P}_{P^*}^{\pi^*}$ is assumed to be continuous w.r.t to Lebesgue measure. It is characterized by the policy density $\pi = \{\pi_i\}_{i \in [H]} \in \Pi$ and the stationary transition density $P = \mathcal{P}$ where $\Pi, \mathcal{P}$ characterize the policy and transition density spaces. The log-likelihood of the set with for a policy π and a transition P reads:

$$l_n(\pi, P) = \frac{1}{n} \sum_{i \in [n]} \log [P(s_1^i) \cdot \pi_1(a_1^i, s_1^i) \prod_{1 < j \leq H} P(s_j^i | s_{j-1}^i, a_{j-1}^i) \pi_j(a_j^i | s_j^i)]$$

The maximum likelihood objective over the density class $\{\mathbb{P}_P^\pi\}_{\pi \in \Pi, P \in \mathcal{P}}$ for the dataset $\mathbb{D}_n^H$

$$\arg \max_{\pi \in \Pi, P \in \mathcal{P}} \sum_{i \in [n]} \sum_{j \in [H]} \left(\log[\pi_i(a_j^i | s_j^i)] \right) + \sum_{i \in [n]} \sum_{j=0}^H \left(\log[P(s_{j+1}^i | s_j^i, a_j^i)] \right) \quad (6)$$

B.3 Concentration Bounds

In this section we provide concentration bounds for the MLE estimators of the policies and the transition model, as well as for our notion of concentrability coefficient. The important takeaway is that the control of the error, i.e., the decay of these concentration bounds depends only on values known to the user, which will allow us to compute confidence policy sets based on these bounds.

B.3.1 Policy estimation

Define the log-loss behavioral cloning estimator for dataset $\mathbb{D}_n^H$ as described in B.2 as

$$\hat{\pi} = \arg \max_{\pi \in \Pi} \sum_{i \in [n]} \sum_{h \in [H]} \log(\pi_h(a_h^i | s_h^i)) \quad (7)$$

which is from Equation (6) equivalent to performing maximum density estimation over the density class $\{\mathbb{P}_{P^*}^\pi\}_{\pi \in \Pi}$. Similar to definition 16 [Foster et al., 2024] define the following

Definition 18 (Policy Covering Number). *For a class $\Pi \subset \{\pi_h : \mathcal{S} \rightarrow \Delta(\mathcal{A})\}$, we say that $\Pi' \subset \{\pi_h : \mathcal{S} \rightarrow \Delta(\mathcal{A})\}$ is an ϵ -cover if for all $\pi \in \Pi$ there exists $\pi' \in \Pi'$ such that*

$$\log \left(\frac{\pi_h(a|s)}{\pi'_h(a|s)} \right) \leq \epsilon \quad \forall (s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]$$

We denote the size of the smallest such cover as $\mathcal{N}_{pol}(\Pi, \epsilon)$

We state the following theorem from [Foster et al., 2024, Appendix C]:

Theorem 19 (Generalization bound for logloss-BC). *The Logloss BC estimator Eq. 7 satisfies with probability $\geq 1 - \delta$*

$$H^2(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*}) \leq \inf_{\epsilon} \left\{ \frac{6 \log(2 \mathcal{N}_{pol}(\Pi, \epsilon/H) \delta^{-1})}{n} + \epsilon \right\}$$

in particular, if Π is finite

$$H^2(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*}) \leq \frac{6 \cdot \log(2 \cdot |\Pi| \cdot \delta^{-1})}{n}$$

Proof. See [Foster et al., 2024, Appendix C]. □

Corollary 20 (Deterministic Stationary Tabular Policies). *If $\Pi = \Pi_S^D$ i.e the set of deterministic tabular policies the log-loss BC estimator Eq. 7 has that with probability at least $1 - \delta$*

$$H^2(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*}) \leq \frac{6 \cdot |\mathcal{S}| \cdot \log(|\mathcal{A}| \cdot \delta^{-1})}{n}$$

Proof. We have $|\Pi_S^D| = |\mathcal{A}|^{|\mathcal{S}|}$ □

In the case we don't have deterministic but stochastic policies, we need to determine $\log(\mathcal{N}_{pol}(\Pi_S, \epsilon))$. This can be accomplished using a discretization argument, where we create a finite ϵ -net that approximates the continuous space of stochastic policies within the desired error tolerance.

B.3.2 Transition model estimation

Here we can give a similar argument as for the policy log loss BC estimator. We define the following estimator

$$\hat{P} = \arg \max_{P \in \mathcal{P}} \sum_{i \in [n]} \sum_{j=0}^H \left(\log[P(s_{j+1}^i | s_j^i, a_j^i)] \right) \quad (8)$$

which is from Eq. 6 equivalent to performing maximum density estimation over the density class $\{\mathbb{P}_P^{\pi^*}\}_{P \in \mathcal{P}}$. Similarly, we define the following notion of covering

629 **Definition 21** (Stationary Transition Log Covering Number). *For a class of stationary transition*
 630 *probability functions $\mathcal{P} \subset \{P : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})\}$ we define that $\mathcal{P}' \subset \{P : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})\}$ is an*
 631 *ϵ -cover if for all $P \in \mathcal{P}$ there exists $P' \in \mathcal{P}'$ such that*

$$\log \left(\frac{P(s'|s, a)}{P'(s'|s, a)} \right) \leq \epsilon \quad \forall (s', s) \in \mathcal{S}, a \in \mathcal{A}$$

632 *We denote the smallest such cover by $\mathcal{N}_{trans}(\mathcal{P}, \epsilon)$*

633 **Assumption B.1** (Realisability of Transition). *We assume the true transition density to be in the*
 634 *model class i.e $P^* \in \mathcal{P}$*

635 We can now give a similar guarantee as for the log loss policy estimate but for the transition estimate

636 **Theorem 22** (Generalisation Bound for MLE Transition Estimator). *The MLE transition estimator*
 637 *of Eq 8 satisfies with probability at least $1 - \delta$*

$$H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \leq \inf_{\epsilon} \left\{ \frac{6 \log(2\mathcal{N}_{trans}(\Pi, \epsilon/H)\delta^{-1})}{n} + \epsilon \right\}$$

638 *Proof.* Given a valid ϵ -cover of $\mathcal{P}$ from definition 21 we have

$$\log \left(\frac{\mathbb{P}_{\hat{P}}^{\pi}}{\mathbb{P}_{P'}^{\pi}} \right) = \sum_{h=1}^H \log \left(\frac{P(s_{h+1}|s_h, a_h)}{P'(s_{h+1}|s_h, a_h)} \right) \leq \epsilon \cdot H$$

639 this means that we get a valid $\epsilon \cdot H$ cover for the trajectory density class. The bound follow be a
 640 direct application of Lemma 17 $\square$

641 **Lemma 23** (Transition Covering Number: Stationary Tabular Stochastic Transitions). *For a class of*
 642 *stationary transition probability functions $\mathcal{P} \subset \{P : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$ where $|\mathcal{S} \times \mathcal{A}|$ is finite (tabular*
 643 *MDP), the ϵ -cover from Definition 21 satisfies:*

$$\log(\mathcal{N}_{trans}(\mathcal{P}, \epsilon)) \leq |\mathcal{S}| \cdot |\mathcal{A}| \cdot (|\mathcal{S}| - 1) \cdot \log \left(\frac{1}{\epsilon} + 1 \right)$$

644 *Proof.* The proof follow a standard geometric discretization argument for finite class of function (see
 645 Chapter 5 Wainwright (2019)). For a given $\epsilon > 0$ we construct a geometric grid:

$$\mathcal{G}_{\epsilon} = \{\delta, \delta \cdot \exp(\epsilon/2), \delta \exp(\epsilon), \delta \exp(3\epsilon/2), \dots, \delta \exp(k\epsilon/2)\}$$

646 where $\delta > 0$ is the minimum probability and k is chosen such that the grid represents a discretization
 647 of the continuous interval $[0, 1]$ i.e

$$\delta \exp(k\epsilon/2) \implies k \geq \frac{2 \log(1/\delta)}{\epsilon}$$

648 Thus the grid size is at most:

$$|\mathcal{G}_{\epsilon}| \leq \left\lceil \frac{2 \log(1/\delta)}{\epsilon} \right\rceil + 1$$

649 For each state action pairs (a, s) define $P(s_i|s, a) = p_i$ for $i \in |\mathcal{S}|$. Note that for the first $1, \dots, |\mathcal{S}| - 1$
 650 there exist $q_i := P'(s_i|s, a) \in \mathcal{G}_{\epsilon}$ that satisfies by construction

$$\exp(-\epsilon/2) \leq \frac{p_i}{q_i} \leq \exp(\epsilon/2)$$

651 For the last state $i = |\mathcal{S}|$, we need to determine $q_{|\mathcal{S}|}$ close enough to $p_{|\mathcal{S}|}$.

652 Let define $S_q := \sum_i^{|\mathcal{S}|-1} q_i$ and $S_p := \sum_i^{|\mathcal{S}|-1} p_i$ we have the constraint

$$\begin{aligned} p_{|\mathcal{S}|} &= 1 - S_p \\ q_{|\mathcal{S}|} &= 1 - S - q \end{aligned}$$

653 From the bound on the first $1 - |\mathcal{S}|$ elements we have

$$S_q \exp(-\epsilon/2) \leq S_p \leq S_q \exp(\epsilon/2)$$

654 For the ratio of the last probability

$$\frac{p_{|\mathcal{S}|}}{q_{|\mathcal{S}|}} = \frac{1 - S_p}{1 - S_q}$$

655 we have the following condition such that we have $\frac{p_{|\mathcal{S}|}}{q_{|\mathcal{S}|}} \geq \exp(-\epsilon)$

$$S_q \leq \frac{1 - \exp(-\epsilon)}{\exp(\epsilon/2) - \exp(-\epsilon)}$$

656 Similarly for the upper bound we have the condition

$$S_q \leq \frac{\exp(\epsilon) - 1}{\exp(\epsilon) - \exp(-\epsilon/2)}$$

657 Combining both constraints we have

$$S_q \leq \min \left\{ \frac{\exp(\epsilon) - 1}{\exp(\epsilon) - \exp(-\epsilon/2)}, \frac{\exp(\epsilon) - 1}{\exp(\epsilon) - \exp(-\epsilon/2)} \right\}$$

658 By Taylor approximation this boils down to

$$S_q \leq \frac{2}{3}$$

659 Hence we select only the combination of points that satisfies

$$\frac{2}{3} \leq S_q \leq 1 - \delta$$

660 It remains to count the number of point we have in our cover i.e the first $|\mathcal{S} - 1|$ that satisfies our
661 constrains

$$(\text{number of grid points})^{|\mathcal{S}-1|} \leq |\mathcal{G}_\epsilon|^{|\mathcal{S}-1|} \leq \left(\left\lceil \frac{2 \log(1/\delta)}{\epsilon} \right\rceil + 1 \right)^{|\mathcal{S}-1|}$$

662 Across all state action pair and taking the logarithm

$$\log(\mathcal{N}_{trans}(\mathcal{P}, \epsilon)) \leq |\mathcal{S}||\mathcal{A}|(|\mathcal{S} - 1|) \log \left(\left\lceil \frac{2 \log(1/\delta)}{\epsilon} \right\rceil + 1 \right)$$

663 choosing $\delta = O(\epsilon)$ yield the result. □

664 **Corollary 24** (Stochastic, stationary, tabular transition setting). *For finite $|\mathcal{S} \times \mathcal{A}|$ (tabular setting)*
665 *and assuming the transition density class to be stochastic and stationary we have with probability at*
666 *least $1 - \delta$*

$$H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \leq \mathcal{O} \left(\frac{|\mathcal{S}|^2 |\mathcal{A}| \log(nH\delta^{-1})}{n} \right)$$

667 where for the theoretical optimal constant $C_{theory} \approx 6$

668 *Proof.* From our lemma on the covering number of transition functions, we have:

$$\log \mathcal{N}_{trans}(\mathcal{P}, \epsilon/H) \leq |\mathcal{S}||\mathcal{A}|(|\mathcal{S} - 1|) \log \left(\frac{H}{\epsilon} + 1 \right)$$

669 For large $\frac{H}{\epsilon}$, we can approximate:

$$\log \left(\frac{H}{\epsilon} + 1 \right) \approx \log \left(\frac{H}{\epsilon} \right)$$

670 Substituting this into our bound:

$$\begin{aligned} H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) &\leq \inf_{\epsilon > 0} \left\{ \frac{6 \log(2) + 6|\mathcal{S}||\mathcal{A}|(|\mathcal{S} - 1|) \log \left(\frac{H}{\epsilon} \right) + 6 \log(\delta^{-1})}{n} + \epsilon \right\} \\ &= \inf_{\epsilon > 0} \left\{ \frac{6 \log(2) + 6D \log(H) - 6D \log(\epsilon) + 6 \log(\delta^{-1})}{n} + \epsilon \right\} \end{aligned}$$

671 where $D = |\mathcal{S}||\mathcal{A}|(|\mathcal{S}| - 1)$ for brevity.

672 To find the optimal ε , we differentiate with respect to ε and set to zero:

$$\begin{aligned} \frac{d}{d\varepsilon} \left[\frac{6 \log(2) + 6D \log(H) - 6D \log(\varepsilon) + 6 \log(\delta^{-1})}{n} + \varepsilon \right] &= -\frac{6D}{n\varepsilon} + 1 = 0 \\ \Rightarrow \varepsilon_{\text{opt}} &= \frac{6D}{n} = \frac{6|\mathcal{S}||\mathcal{A}|(|\mathcal{S}| - 1)}{n} \end{aligned}$$

673 Substituting this optimal value back:

$$\begin{aligned} H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) &\leq \frac{6 \log(2) + 6D \log(H) - 6D \log\left(\frac{6D}{n}\right) + 6 \log(\delta^{-1})}{n} + \frac{6D}{n} \\ &= \frac{6 \log(2) + 6D \log(H) - 6D \log(6D) + 6D \log(n) + 6 \log(\delta^{-1}) + 6D}{n} \\ &= \frac{6 \log(2) + 6 \log(\delta^{-1}) + 6D \log\left(\frac{nH}{6D}\right) + 6D}{n} \end{aligned}$$

674 For large state spaces where $|\mathcal{S}| - 1 \approx |\mathcal{S}|$, and defining $\tilde{D} = |\mathcal{S}|^2|\mathcal{A}|$, this becomes:

$$H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \leq \frac{6 \log(2) + 6 \log(\delta^{-1}) + 6\tilde{D} \log\left(\frac{nH}{6\tilde{D}}\right) + 6\tilde{D}}{n}$$

675 For large n and $\tilde{D}$, the dominant term is $\frac{6\tilde{D} \log(nH)}{n}$, and we can combine the logarithmic terms to
676 get:

$$H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) = O\left(\frac{|\mathcal{S}|^2|\mathcal{A}| \log(nH\delta^{-1})}{n}\right)$$

677 Note that the constant 6 appears in the full derivation. This completes the proof. $\square$

678 B.3.3 Concentrability Coefficient Upper Bound

679 **Definition 25** (Concentrability Coefficient). *We define the following quantity as the "concentrability
680 coefficient":*

$$C(\hat{\pi}, \pi^*) = \sup_{t \in [H]} \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}: d_{P^*}^{\pi^*,t}(s,a) > 0} \frac{d_{P^*}^{\pi,t}(s,a)}{d_{P^*}^{\pi^*,t}(s,a)}$$

681 *which measures the maximum ratio between the state-action distributions induced by policies $\hat{\pi}$ and
682 π^* under the true dynamics P^* .*

683 **Assumption B.2** (Minimum Visitation Probability). *There exists a constant $\gamma_{\min} > 0$ such that for
684 all state-action-time tuples with non-zero probability under the optimal policy:*

$$\min_{(s,a,t): d_{P^*}^{\pi^*,t}(s,a) > 0} d_{P^*}^{\pi^*,t}(s,a) \geq \gamma_{\min}$$

685 **Lemma 26** (Concentrability Coefficient Bound). *Consider a policy estimator $\hat{\pi}$ satisfying*

$$H^2(\mathbb{P}_{\hat{P}^*}, \mathbb{P}_{P^*}^{\pi^*}) \leq R$$

686 *Then, under Assumption [4.2](#) the concentration coefficient is bounded by:*

$$C(\hat{\pi}, \pi^*) \leq 1 + \frac{2\sqrt{R}}{\gamma_{\min}}$$

687 *Proof.* We will proceed by upper bounding the numerator using the condition on the Hellinger
688 distance followed by lower bounding with concentration the denominator.
689 For the upper bound note that

$$\sup_{a,s} |d_{P^*}^{\hat{\pi},t} - d_{P^*}^{\pi^*,t}| = 2 \cdot TV(d_{P^*}^{\hat{\pi},t}, d_{P^*}^{\pi^*,t})$$

690 Recalling that the state-action distribution $d_{P^*}^{\pi,t}(s, a)$ is the marginal distribution of the trajectory
691 distribution at time step t . Explicitly:

$$d_{P^*}^{\pi,t}(s, a) = \int_{\tau_{-t}} \mathbb{P}_{P^*}^{\pi}(\tau) d\tau_{-t} =: \mathbb{P}_{P^*}^{\pi}(s_t = s, a_t = a)$$

692 where τ_{-t} denotes all time steps in the trajectory except for time t , and $\mathbb{P}_{P^*}^{\pi}(\tau)$ is the probability of
693 trajectory τ under policy π and dynamics P^* .

694 Hence

$$\begin{aligned} TV(d_{P^*}^{\hat{\pi},t}, d_{P^*}^{\pi^*,t}) &= 2 \cdot TV(\mathbb{P}_{P^*}^{\hat{\pi}}(s_t = s, a_t = a), \mathbb{P}_{P^*}^{\pi^*}(s_t = s, a_t = a)) \\ &\leq 2 \cdot TV(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*}) \\ &\leq 2 \cdot \sqrt{H^2(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*})} \leq 2\sqrt{R} \end{aligned}$$

695 together with

696 By Assumption 1, we have a lower bound on the minimum state-action visitation probability:

$$\min_{(s,a,t): d_{P^*}^{\pi^*,t}(s,a) > 0} d_{P^*}^{\pi^*,t}(s, a) \geq \gamma_{min}$$

697 Finally, we combine the upper bound on the numerator and the lower bound from Assumption 1 to
698 get $\forall(a, s, t) \quad \text{s.t.} \quad d_{P^*}^{\pi^*,t}(a, s) > 0$:

$$\begin{aligned} C(\hat{\pi}, \pi^*) &= 1 + \frac{\sup_{a,s,t} |d_{P^*}^{\hat{\pi},t}(s, a) - d_{P^*}^{\pi^*,t}(s, a)|}{\inf_{a,s,t} d_{P^*}^{\pi^*,t}(s, a)} \\ &\leq 1 + \frac{2\sqrt{R}}{\inf_{a,s,t} d_{P^*}^{\pi^*,t}(s, a)} \\ &\leq 1 + \frac{2\sqrt{R}}{\gamma_{min}} \end{aligned}$$

699 This completes the proof, giving us a deterministic bound on the concentration coefficient that depends
700 on the Hellinger distance between trajectory distributions and the minimum visitation probability of
701 the optimal policy. $\square$

702 B.3.4 Confidence Set Construction

703 In this section we will derive a distributional confidence set on the trajectory space in the form of a
704 hellinger ball, accounting for the error of the MLE density estimates $\hat{\pi}$ and $\hat{P}$. We start by presenting
705 the following in between result

706 **Lemma 27** (Technical Results). *Assume finite state and action pair: $\mathcal{S} \times \mathcal{A}$. The following upper
707 bound is true $\forall \pi \in \Pi$ with π^* being the true policy and $P^*, \hat{P}$ being the true and estimated transition
708 models:*

$$H^2(\mathbb{P}_{\hat{P}}^{\pi}, \mathbb{P}_{P^*}^{\pi^*}) \leq H \cdot C(\pi, \pi^*) \cdot H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*})$$

709 where

$$C(\pi, \pi^*) = \sup_{t \in [H]} \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}: d_{P^*}^{\pi^*,t}(s,a) > 0} \frac{d_{P^*}^{\pi,t}(s, a)}{d_{P^*}^{\pi^*,t}(s, a)}$$

710 *Proof.* We derive the proof in three steps:

711

Step 1.

$$H^2(\mathbb{P}_{\hat{P}}^\pi, \mathbb{P}_{P^*}^\pi) \leq \sum_{t \in [H-1]} \mathbb{E}_{(s_t, a_t) \sim d_{P^*}^{\pi, t}} \left[H^2 \left(\hat{P}(\cdot | s_t, a_t), P^*(\cdot | s_t, a_t) \right) \right]$$

Step 2.

$$\mathbb{E}_{(s_t, a_t) \sim d_{P^*}^{\pi, t}} \left[H^2 \left(\hat{P}(\cdot | s_t, a_t), P^*(\cdot | s_t, a_t) \right) \right] \leq C(\pi, \pi^*) \cdot \mathbb{E}_{(s_t, a_t) \sim d_{P^*}^{\pi^*, t}} \left[H^2 \left(\hat{P}(\cdot | s_t, a_t), P^*(\cdot | s_t, a_t) \right) \right]$$

Step 3.

$$H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \geq \frac{1}{H} \cdot \mathbb{E}_{(s_t, a_t) \sim d_{P^*}^{\pi^*, t}} \left[H^2 \left(\hat{P}(\cdot | s_t, a_t), P^*(\cdot | s_t, a_t) \right) \right]$$

712 Proof Step 1:

713

$$\begin{aligned} H^2(\mathbb{P}_{\hat{P}}^\pi, \mathbb{P}_{P^*}^\pi) &= 1 - \int_{\mathcal{T}} 1 - \mu_0(s_0) \prod_{t=0}^{H-1} \pi(a_t | s_t) \sqrt{\hat{P}(s_{t+1} | a_t, s_t) P^*(s_{t+1} | s_t, a_t)} d\tau \\ &= 1 - \int_{\mathcal{T}} p_{P^*}^\pi(\tau) \cdot \frac{\mu_0(s_0) \prod_{t=0}^{H-1} \pi(a_t | s_t) \sqrt{\hat{P}(s_{t+1} | a_t, s_t) P^*(s_{t+1} | s_t, a_t)}}{\mu_0(s_0) \prod_{t=0}^{H-1} \pi(a_t | s_t) P^*(s_{t+1} | s_t, a_t)} d\tau \\ &= 1 - \int_{\mathcal{T}} p_{P^*}^\pi(\tau) \cdot \prod_{t=0}^{H-1} \sqrt{\frac{\hat{P}(s_{t+1} | s_t, a_t)}{P^*(s_{t+1} | s_t, a_t)}} d\tau \end{aligned}$$

714 Next define for ease of notation :

$$\begin{aligned} \alpha_t(s_{t+1}, a_t, s_t) &:= \sqrt{\frac{\hat{P}(s_{t+1} | s_t, a_t)}{P^*(s_{t+1} | s_t, a_t)}} \\ \gamma_t(s_t, a_t) &:= \int_{s'} \sqrt{\hat{P}(s_{t+1} | s_t, a_t) P^*(s_{t+1} | s_t, a_t)} ds' = \int_{s'} P^*(s' | s_t, a_t) \cdot \alpha_t(s', s_t, a_t) ds' \end{aligned}$$

715 Notice that γ_t is a BC coefficient i.e

$$1 - \gamma_t(s_t, a_t) = H^2(\hat{P}(\cdot | s_t, a_t), P^*(\cdot | s_t, a_t))$$

716 Using notation above:

$$H^2(\mathbb{P}_{\hat{P}}^\pi, \mathbb{P}_{P^*}^\pi) = 1 - \mathbb{E}_{\tau \sim \mathbb{P}_{P^*}^\pi} \left[\prod_{t=0}^{H-1} \alpha_t(s_{t+1}, s_t, a_t) \right]$$

717 Using conditional expectation (law of iterated expectation) we change the distribution in the expecta-
718 tion from $\mathbb{P}_{P^*}^\pi$ to the so called state-action distribution $d_{P^*}^{\pi, t}$. To show this argument we show it for
719 state action pair (a_0, s_0, s_1) . The rest follows by using the same idea:

$$\begin{aligned} \mathbb{E}_{\tau \sim \mathbb{P}_{P^*}^\pi} \left[\prod_{t=0}^{H-1} \alpha_t(s_{t+1}, s_t, a_t) \right] &= \mathbb{E}_{s_0, a_0} \left[\mathbb{E}_{s_1 | s_0, a_0} \left[\alpha_0(s_1, s_0, a_0) \right] \cdot \mathbb{E}_{a_1 \cup \tau_{[2:H-1]} | s_1} \left[\prod_{t=1}^{H-1} \alpha_t(s_{t+1}, s_t, a_t) \right] \right] \\ &= \int_{s_0, a_0} \mu_0(s_0) \cdot \pi_1(a_0 | s_0) \int_{s_1} P^*(s_1 | s_0, a_0) \cdot \alpha_0(s_1, s_0, a_0) \\ &\quad \cdot \mathbb{E}_{a_1 \cup \tau_{[2:H-1]} | s_1} \left[\prod_{t=1}^{H-1} \alpha_t(s_{t+1}, s_t, a_t) \right] \\ &= \int_{s_0, a_0} \mu_0(s_0) \cdot \pi_1(a_0 | s_0) \cdot \gamma_0(s_0, a_0) \cdot \left[\text{""} \right] \end{aligned}$$

720 Notice that

$$\mu_0(s_0) \cdot \pi_1(a_0|s_0) = \int p_{P^*}^\pi d\tau_{[1:H-1]} =: d_{P^*}^{\pi,0}(s_0, a_0) \quad \text{Marginal over } (s_0, a_0)$$

721 Hence using a recursive argument we have

$$\begin{aligned} H^2(\mathbb{P}_{\hat{P}}^\pi, \mathbb{P}_{P^*}^\pi) &= 1 - \mathbb{E}_{\tau \sim \mathbb{P}_{P^*}^\pi} \left[\prod_{t=0}^{H-1} \alpha_t(s_{t+1}, s_t, a_t) \right] \\ &= 1 - \prod_{t=0}^{H-1} \mathbb{E}_{d_{P^*}^{\pi,t}} \left[\gamma_t(s_t, a_t) \right] \end{aligned}$$

722 Using the fact that

$$1 - \prod_i x_i \leq \sum_i (1 - x_i) \quad \forall x_i \in [0, 1]$$

723 and by the fact that $\gamma_t \in [0, 1] \forall t$ we have

$$\begin{aligned} H^2(\mathbb{P}_{\hat{P}}^\pi, \mathbb{P}_{P^*}^\pi) &\leq \sum_{t=0}^{H-1} \mathbb{E}_{d_{P^*}^{\pi,t}} (1 - \gamma_t(s_t, a_t)) \\ &= \sum_{t=0}^{H-1} \mathbb{E}_{d_{P^*}^{\pi,t}} \left[H^2(\hat{P}(\cdot|s_t, a_t), P^*(\cdot|s_t, a_t)) \right] \end{aligned}$$

724 Proof Step 2:

725

$$\begin{aligned} \mathbb{E}_{(s_t, a_t) \sim d_{P^*}^{\pi,t}} [H^2(\hat{P}(\cdot|s_t, a_t), P^*(\cdot|s_t, a_t))] &= \sum_{s,a} d_{P^*}^{\pi,t}(s, a) \cdot H^2(\hat{P}(\cdot|s, a), P^*(\cdot|s, a)) \\ &= \sum_{s,a} \frac{d_{P^*}^{\pi,t}(s, a)}{d_{P^*}^{\pi^*,t}(s, a)} \cdot d_{P^*}^{\pi^*,t}(s, a) \cdot H^2(\hat{P}(\cdot|s, a), P^*(\cdot|s, a)) \\ &= \sum_{s,a} \frac{d_{P^*}^{\pi,t}(s, a)}{d_{P^*}^{\pi^*,t}(s, a)} \cdot d_{P^*}^{\pi^*,t}(s, a) \cdot H^2(\hat{P}(\cdot|s, a), P^*(\cdot|s, a)) \end{aligned}$$

726 By definition of the concentrability coefficient:

$$C(\pi, \pi^*) = \sup_{t \in [H]} \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}: d_{P^*}^{\pi^*,t}(s,a) > 0} \frac{d_{P^*}^{\pi,t}(s, a)}{d_{P^*}^{\pi^*,t}(s, a)}$$

727 Therefore:

$$\frac{d_{P^*}^{\pi,t}(s, a)}{d_{P^*}^{\pi^*,t}(s, a)} \leq C(\pi, \pi^*) \quad \forall t \quad \forall (s, a) \text{ where } d_{P^*}^{\pi^*,t} \geq 0$$

728 Proof Step 3:

729 Starting from our previous expression:

$$H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) = 1 - \prod_{t=0}^{H-1} \mathbb{E}_{d_{P^*}^{\pi^*,t}} [\gamma_t(s_t, a_t)]$$

730 Let's denote

$$x_i := 1 - \gamma_i(s_i, a_i) = H^2(\hat{P}(\cdot|s_i, a_i), P^*(\cdot|s_i, a_i))$$

731 Using the fact that $(1 - x) \leq e^{-x}$ for all x , we have:

$$\prod_{i=1}^H (1 - x_i) \leq \exp \left(- \sum_{i=1}^H x_i \right)$$

732 By the second-order Taylor expansion of the exponential function:

$$\exp\left(-\sum_{i=1}^H x_i\right) \leq 1 - \sum_{i=1}^H x_i + \frac{1}{2} \left(\sum_{i=1}^H x_i\right)^2$$

733 Since $x_i \leq 1$ for all i , we know that $\sum_{i=1}^H x_i \leq H$, which gives us:

$$\frac{1}{2} \left(\sum_{i=1}^H x_i\right)^2 \leq \frac{H}{2} \sum_{i=1}^H x_i$$

734 Therefore:

$$\begin{aligned} H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) &\geq \sum_{i=1}^H x_i - \frac{1}{2} \left(\sum_{i=1}^H x_i\right)^2 \\ &\geq \sum_{i=1}^H x_i - \frac{H}{2} \sum_{i=1}^H x_i \\ &= \left(1 - \frac{H}{2}\right) \sum_{i=1}^H x_i \end{aligned}$$

735 For the bound $H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \geq \frac{1}{H} \sum_{i=1}^H x_i$ to hold, we need:

$$\left(1 - \frac{H}{2}\right) \sum_{i=1}^H x_i \geq \frac{1}{H} \sum_{i=1}^H x_i$$

736 This is satisfied when:

$$\sum_{i=1}^H x_i \leq \frac{2(H-1)}{H}$$

737 This condition is typically met for good estimators where Hellinger distances are small. For large H ,
738 the bound approaches 2.

739 Under this condition, we can establish:

$$H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \geq \frac{1}{H} \sum_{t=0}^{H-1} \mathbb{E}_{(s_t, a_t) \sim d_{P^*}^{\pi^*, t}} [H^2(\hat{P}(\cdot|s_t, a_t), P^*(\cdot|s_t, a_t))]$$

740

□

741 **Lemma 28** (Policy Density Confidence Set). *Assume the following events hold:*

$$\begin{aligned} E_1 &:= \left\{ H^2(\mathbb{P}_{\hat{P}^*}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \leq R_1(\delta_1) \right\}, \quad \text{such that } P(E_1) \geq 1 - \delta_1, \\ E_2 &:= \left\{ H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \leq R_2(\delta_2) \right\}, \quad \text{such that } P(E_2) \geq 1 - \delta_2, \end{aligned}$$

742 where $\hat{\pi}$ and $\hat{P}$ are estimators of the policy and the transition dynamics, respectively.

743 Then, under Assumption 4.2 the policy set

$$C_{1-\delta} := \left\{ \pi : \sqrt{H^2(\mathbb{P}_{\hat{P}}^{\pi}, \mathbb{P}_{\hat{P}}^{\hat{\pi}})} \leq \sqrt{R_1} + \sqrt{R_2} \cdot \left(1 + \sqrt{\left(1 + \frac{2\sqrt{R_1}}{\gamma_{\min}}\right) \cdot H} \right) \right\}$$

744 is a confidence set of level $1 - \delta = 1 - (\delta_1 + \delta_2)$, i.e.,

$$P(\pi^* \in \Pi_{1-\delta}^{\text{offline}}) \geq 1 - (\delta_1 + \delta_2).$$

745 *Proof.* For ease of notation, let us define:

$$\sqrt{H^2(\mathbb{P}_{P_1}^{\pi_1}, \mathbb{P}_{P_2}^{\pi_2})} =: \|(\pi_1, P_1) - (\pi_2, P_2)\|$$

with $\|\cdot\| := \|\cdot\|_{L_2(\mu(\mathbb{R}))}$

746 We can then decompose by the triangle inequality:

$$\|(\pi^*, \hat{P}) - (\hat{\pi}, \hat{P})\| \leq \|(\pi^*, \hat{P}) - (\pi^*, P^*)\| + \|(\pi^*, P^*) - (\hat{\pi}, P^*)\| + \|(\hat{\pi}, P^*) - (\hat{\pi}, \hat{P})\|$$

747 From Lemma 27, we have:

$$\|(\hat{\pi}, P^*) - (\hat{\pi}, \hat{P})\| \leq \sqrt{C(\hat{\pi}, \pi^*) \cdot H} \cdot \|(\pi^*, \hat{P}) - (\pi^*, P^*)\|$$

748 From Assumption 4.2 and our concentrability coefficient bound, we have:

$$C(\hat{\pi}, \pi^*) \leq 1 + \frac{2\sqrt{R_1}}{\gamma_{\min}}$$

749 From event E_1 , we have:

$$\|(\pi^*, P^*) - (\hat{\pi}, P^*)\| \leq \sqrt{R_1}$$

750 From event E_2 , we have:

$$\|(\pi^*, \hat{P}) - (\pi^*, P^*)\| \leq \sqrt{R_2}$$

751 Then, assuming events $E_1 \cap E_2$ hold jointly, with probability at least $1 - (\delta_1 + \delta_2)$, we have:

$$\|(\pi^*, \hat{P}) - (\hat{\pi}, \hat{P})\| \leq \sqrt{R_1} + \sqrt{R_2} \cdot \left(1 + \sqrt{\left(1 + \frac{2\sqrt{R_1}}{\gamma_{\min}}\right) \cdot H}\right)$$

752 Hence, by construction, the set:

$$\Pi_{1-\delta}^{\text{offline}}(\Pi) := \left\{ \pi \in \Pi : \|(\pi, \hat{P}) - (\hat{\pi}, \hat{P})\| \leq \sqrt{R_1} + \sqrt{R_2} \cdot \left(1 + \sqrt{\left(1 + \frac{2\sqrt{R_1}}{\gamma_{\min}}\right) \cdot H}\right) \right\}$$

753 contains π^* with probability at least $1 - (\delta_1 + \delta_2)$. □

754 B.4 Performance Guarantees

755 We apply our method of constructing confidence sets based on distributional guarantees for maximum
 756 likelihood density estimation to the tabular reinforcement learning setting with state space $\mathcal{S}$ and
 757 action space $\mathcal{A}$. We consider deterministic stationary tabular policies ($\Pi = \Pi_S^D$) and stochastic
 758 stationary tabular transitions, though the method is versatile to other settings with appropriate
 759 adaptation of the corresponding covering numbers (cf. Definitions 18 and 21).

760 Let $\hat{\pi}$ be the log-loss BC estimator (Equation 11) of the true policy π^* , and $\hat{P}$ be the MLE estimator
 761 (Equation 2) of the true transition model P^* . The concentration bounds for these estimators are,
 762 with probability at least $1 - \delta_1$ and $1 - \delta_2$ respectively:

$$H^2(\mathbb{P}_{P^*}^{\hat{\pi}}, \mathbb{P}_{P^*}^{\pi^*}) \leq R_1 = \frac{4 \cdot |\mathcal{S}| \cdot \log(|\mathcal{A}| \cdot \delta_1^{-1})}{n} \quad (\text{Corollary 20})$$

$$H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{P^*}^{\pi^*}) \leq R_2 = \frac{4 \cdot |\mathcal{S}|^2 \cdot |\mathcal{A}| \cdot \log(nH\delta_2^{-1})}{n} \quad (\text{Corollary 24})$$

763 Additionally, we make the following assumption about the minimum visitation probability under the
 764 optimal policy:

765 **Assumption B.3** (Minimum Visitation Probability). *There exists a constant $\gamma_{\min} > 0$ such that for*
 766 *all state-action-time tuples with non-zero probability under the optimal policy:*

$$\min_{(s,a,t): d_{P^*}^{\pi^*,t}(s,a) > 0} d_{P^*}^{\pi^*,t}(s,a) \geq \gamma_{\min}$$

Under this assumption, the concentrability coefficient is bounded by:

$$C(\hat{\pi}, \pi^*) \leq 1 + \frac{2 \cdot \sqrt{R_1}}{\gamma_{\min}}$$

Theorem 29 (Policy Confidence Set). *Under the setting described above and Assumption B.3 with $\delta_1 = \delta_2 = \delta/2$ and defining*

$$\alpha := \sqrt{4 \cdot |\mathcal{S}| \cdot \log(|\mathcal{A}| \cdot 2/\delta)} \\ \beta := \sqrt{4 \cdot |\mathcal{S}|^2 \cdot |\mathcal{A}| \cdot \log(nH \cdot 2/\delta)}$$

The policy set

$$\Pi_{1-\delta}^{\text{offline}} := \left\{ \pi : \sqrt{H^2(\mathbb{P}_{\hat{P}}^{\pi}, \mathbb{P}_{\hat{P}}^{\hat{\pi}})} \leq \sqrt{R_1} + \sqrt{R_2} \cdot \left(1 + \sqrt{\left(1 + \frac{2\sqrt{R_1}}{\gamma_{\min}} \right) \cdot H} \right) \right\}$$

is a confidence set of level $1 - \delta$ containing π^* with probability at least $1 - \delta$. The radius of this confidence set is explicitly:

$$\text{Radius} = \frac{\alpha}{\sqrt{n}} + \frac{\beta}{\sqrt{n}} \cdot \left(1 + \sqrt{H \cdot \left(1 + \frac{2\alpha}{\gamma_{\min} \cdot \sqrt{n}} \right)} \right)$$

Proof. The proof follows directly from Lemma 28 by applying our bounds on $H^2(\mathbb{P}_{\hat{P}^*}^{\hat{\pi}}, \mathbb{P}_{\hat{P}^*}^{\pi^*})$ and $H^2(\mathbb{P}_{\hat{P}}^{\pi^*}, \mathbb{P}_{\hat{P}^*}^{\pi^*})$, along with our bound on the concentrability coefficient from Assumption B.3. Setting $\delta_1 = \delta_2 = \delta/2$ and substituting the appropriate values gives us the result. $\square$

C Online Estimation

The underlying setting is described in the Section Problem Setup

C.1 Elliptical Confidence Set

For completeness and to make our paper self-contained, we provide a brief overview of the online preference-based learning approach used in our method. The formulation presented in this section closely follows the work of Saha et al. [2023] and Faury et al. [2020], with adaptations to our specific setting. We include this background to help the reader understand the elliptical confidence set construction that forms a foundation for our theoretical analysis.

In the logistic model, a natural way of computing an estimator $\mathbf{w}_t$ of $\mathbf{w}^*$ given trajectory pairs $\{(\tau_\ell^1, \tau_\ell^2)\}_{\ell=1}^{t-1}$ and preference feedback values $\{o_\ell\}_{\ell=1}^{t-1}$ is via maximum likelihood estimation. At time t the regularized log-likelihood (or negative cross-entropy loss) of a parameter $\mathbf{w}$ can be written as:

$$\mathcal{L}_t^\lambda(\mathbf{w}) = \sum_{\ell=1}^{t-1} \left(o_\ell \log(\sigma(\langle \phi(\tau_\ell^1) - \phi(\tau_\ell^2), \mathbf{w} \rangle)) - \frac{\lambda}{2} \|\mathbf{w}\|^2 \right) \\ + (1 - o_\ell) \log(1 - \sigma(\langle \phi(\tau_\ell^1) - \phi(\tau_\ell^2), \mathbf{w} \rangle)),$$

where $\lambda > 0$ is a regularization parameter. The function $\mathcal{L}_t^\lambda$ is strictly concave for $\lambda > 0$. The maximum likelihood estimator $\hat{\mathbf{w}}_t^{\text{MLE}}$ can be written as $\hat{\mathbf{w}}_t^{\text{MLE}} = \arg \max_{\mathbf{w} \in \mathbb{R}^d} \mathcal{L}_t^\lambda(\mathbf{w})$. Unfortunately, $\hat{\mathbf{w}}_t^{\text{MLE}}$ may not satisfy the boundedness Assumption 1, so we instead make use of a projected version of $\hat{\mathbf{w}}_t^{\text{MLE}}$. Following Faury et al. [2020], and recalling Assumption 1, we define a data matrix and a transformation of $\hat{\mathbf{w}}_t^{\text{MLE}}$ given by

$$\mathbf{V}_t = \kappa \lambda \mathbf{I}_d + \sum_{\ell=1}^{t-1} (\phi(\tau_\ell^1) - \phi(\tau_\ell^2)) (\phi(\tau_\ell^1) - \phi(\tau_\ell^2))^\top \\ g_t(\mathbf{w}) = \sum_{\ell=1}^{t-1} \sigma(\langle \phi(\tau_\ell^1) - \phi(\tau_\ell^2), \mathbf{w} \rangle) (\phi(\tau_\ell^1) - \phi(\tau_\ell^2)) + \lambda \mathbf{w}$$

Then, the projected parameter, along with its confidence set, is given by

$$\mathbf{w}_t^{Proj} = \arg \min_{\mathbf{w} \text{ s.t. } \|\mathbf{w}\| \leq W} \|g_t(\mathbf{w}) - g_t(\hat{\mathbf{w}}_t^{\text{MLE}})\|_{\mathbf{V}_t^{-1}}$$

$$\mathcal{C}_t(\delta) = \{\mathbf{w} \text{ s.t. } \|\mathbf{w} - \mathbf{w}_t^P\|_{\mathbf{V}_t} \leq 2\kappa\beta_t(\delta)\}$$

where $\beta_t(\delta) = \sqrt{\lambda}W + \sqrt{\log(1/\delta) + 2d \log(1 + \frac{tB^2}{\kappa\lambda d})}$. We restate a bound by Fauray et al. [2020] that shows the probability of $\mathbf{w}_*$ being in $\mathcal{C}_t(\delta)$ for all $t \geq 1$ can be lower bounded.

Lemma 30 (Confidence Set Coverage). *Let $\delta \in (0, 1]$ and define the event that $\mathbf{w}_*$ is in the confidence interval $\mathcal{C}_t(\delta)$ for all $t \in \mathbb{N}$:*

$$E_{w^*} = \{\forall t \geq 1, \mathbf{w}_* \in \mathcal{C}_t(\delta)\}.$$

Then $\mathbb{P}(E_{w^*}) \geq 1 - \delta$.

Proof. This follows from Fauray et al. [2020] with a slight modification to account for our bounded feature assumption. $\square$

This elliptical confidence set construction, which has its roots in generalized linear bandits [Filippi et al., 2010, Fauray et al., 2020], forms a critical component of our online learning algorithm. By maintaining and updating these confidence sets as new preference data is collected, our algorithm can efficiently balance exploration and exploitation to identify the optimal policy. The confidence bounds ensure that with high probability, the true reward parameter lies within our constructed set throughout the learning process, which is essential for the regret guarantees we derive in the following sections.

C.2 Norm Relation Between Data Matrices

For completeness, we restate key results from Saha et al. [2023] concerning the relationships between various data matrices that arise in our analysis. These results are included to ensure the appendix is self-contained and to provide context for our subsequent analysis. The full proofs of these results can be found in the original paper.

Saha et al. [2023] establishes relationships between three key matrices:

- V_t - The empirical data matrix constructed from observed trajectories
- $\bar{V}_t^{P^*}$ - The expected data matrix under the true transition dynamics P^*
- $\bar{V}_t$ - The expected data matrix under the estimated transition dynamics $\hat{P}_t$

These matrices are defined as follows:

$$V_t = \kappa\lambda\mathbf{I}_d + \sum_{\ell=1}^{t-1} (\phi(\tau_\ell^1) - \phi(\tau_\ell^2))(\phi(\tau_\ell^1) - \phi(\tau_\ell^2))^\top$$

$$\bar{V}_t^{P^*} = \kappa\lambda\mathbf{I}_d + \sum_{\ell=1}^{t-1} (\phi(\pi_\ell^1) - \phi(\pi_\ell^2))(\phi(\pi_\ell^1) - \phi(\pi_\ell^2))^\top$$

$$\bar{V}_t = \kappa\lambda\mathbf{I}_d + \sum_{\ell=1}^{t-1} (\phi^{\hat{P}_t}(\pi_\ell^1) - \phi^{\hat{P}_t}(\pi_\ell^2))(\phi^{\hat{P}_t}(\pi_\ell^1) - \phi^{\hat{P}_t}(\pi_\ell^2))^\top$$

Where $\phi(\pi)$ represents the expected feature vector under policy π and the true transition dynamics P^* , while $\phi^{\hat{P}_t}(\pi)$ represents the expected feature vector under policy π and the estimated transition dynamics $\hat{P}_t$.

Saha et al. [2023] introduces a precision event that relates the empirical matrix V_T to the expected matrix $\bar{V}_T^{P^*}$:

$$E_{\bar{V}_T^{P^*}} = \{\bar{V}_T^{P^*} \preceq 2V_T + 84B^2d \log((1 + 2T)/\delta)\mathbf{I}_d\}$$

Under this event, they establish the following bound:

Lemma 31 (Adapted from Saha et al. [2023] Corollary 1). *Under Assumption 1, conditioned on event $E_{w^*} \cap E_{\bar{V}_T^{P^*}}$, for any $t \in [T]$*

$$\|\mathbf{w}^* - \mathbf{w}_t^L\|_{\bar{V}_t^{P^*}} \leq 4\kappa\beta_t(\delta) + \alpha_{d,T}(\delta),$$

where $\alpha_{d,T}(\delta) = 20BW\sqrt{d\log(T(1+2T)/\delta)}$. Furthermore, if $\delta \leq 1/e$, then $\mathbb{P}(E_{w^*} \cap E_{\bar{V}_T^{P^*}}) \geq 1 - \delta - \delta \log_2 T$.

Additionally, Saha et al. [2023] relates norms based on the matrix $\bar{V}_t^{P^*}$ with those based on $\bar{V}_t$:

Lemma 32 (Adapted from Saha et al. [2023] Lemma 3). *Let $\bar{\mathcal{E}}_0$ be the event that for all $t \in \mathbb{N}$,*

$$\|\mathbf{w}_t^{proj} - \mathbf{w}_*\|_{\bar{V}_t} \leq \sqrt{2}\|\mathbf{w}_t^{proj} - \mathbf{w}_*\|_{\bar{V}_t^{P^*}} + \sqrt{\sum_{\ell=1}^{t-1} 4 \left(\hat{B}_\ell \left(\pi, 2WB, \frac{\delta'}{8\ell^3|A||S|} \right) \right)^2} + \frac{1}{t}$$

where $\delta' = \frac{\delta}{(1+4W)^d}$ and $\epsilon = \frac{1}{2\kappa\lambda+4B^2t^\alpha}$. Then $\mathbb{P}(\bar{\mathcal{E}}_0) \geq 1 - \delta$.

Note that the bonus function $\hat{B}$ is defined in Lemma 33

These norm relations from Saha et al. [2023] are essential in our regret analysis, as they allow us to relate confidence bounds across different probability spaces and to bound the regret of our algorithm.

C.3 Transition Estimation and Bonus Terms

Note that the offline estimator of the transition probabilities based on the log-loss MLE in Equation (2), when the state-action space is discrete, is equivalent to the following count-based estimator (derivable using a simple Lagrange multiplier argument):

$$\hat{P}_{\text{offline}}(s'|s, a) = \frac{N_{\text{offline}}(s'|s, a)}{N_{\text{offline}}(s, a)}$$

where

$$N_{\text{offline}}(s, a) := \sum_{i \in [n]} \sum_{h \in [H]} \mathbb{I}\{s_h^i = s, a_h^i = a\}$$

$$N_{\text{offline}}(s'|s, a) := \sum_{i \in [n]} \sum_{h \in [H]} \mathbb{I}\{s_{h+1}^i = s', s_h^i = s, a_h^i = a\}$$

This equivalence allows us to initialize the online estimation process with the count estimator from the offline data (see line 3 in Algorithm 1), yielding the combined estimator for the transition model:

$$\hat{P}_t(s'|s, a) := \frac{N_{\text{offline}}(s'|s, a) + N_t(s'|s, a)}{N_{\text{offline}}(s, a) + N_t(s, a)} \quad (9)$$

From this estimator, we adapt two key lemmas from Chatterji et al. [2021] that will define our notion of bonus terms.

Lemma 33 (Moment Transition Difference Error). *Consider the transition count estimator $\hat{P}_t$ from Equation (9). Further assume the trajectory data follows a martingale structure adapted to the natural filtration of the problem. For any fixed policy $\pi \in \Pi$ and any scalar function $f : \mathcal{T} \rightarrow \mathbb{R}$ such that $|f(\tau)| < \eta$, with probability at least $1 - \delta$ for all $t \in \mathbb{N}$:*

$$\mathbb{E}_{\mathbb{P}_{P^*}}[f(\tau)] - \mathbb{E}_{\mathbb{P}_{\hat{P}_t}}[f(\tau)] \leq \mathbb{E}_{\mathbb{P}_{\hat{P}_t}} \left[\sum_{h \in [H]} \xi_{s_h, a_h}^t(\eta, \delta) \right] =: \hat{B}_t(\pi, \eta, \delta)$$

where

$$\xi_{s_h, a_h}^t(\eta, \delta) := \min \left(2\eta, 4\eta \sqrt{\frac{H \log(|\mathcal{S}| \cdot |\mathcal{A}|) + \log \left(\frac{6 \log(N_t(s_h, a_h) + N_{\text{offline}}(s_h, a_h))}{\delta} \right)}{N_t(s_h, a_h) + N_{\text{offline}}(s_h, a_h)}} \right)$$

The term $\hat{B}_t(\pi, \eta, \delta)$ serves as our "bonus" term.

849 *Proof.* Our combined estimator incorporates both online data (adapted to the natural filtration) and
 850 offline data (assumed i.i.d.). We can artificially treat the offline data as though it were adapted to the
 851 natural filtration as well, by considering it as "past" observations. This allows us to directly apply the
 852 proof methodology from Chatterji et al. [2021] (Lemma B.1) to our combined count estimator.

853 The key insight is that the martingale structure of the estimation error is preserved when combining
 854 offline and online counts, with the benefit of reduced variance due to the increased denominator
 855 $(N_t(s_h, a_h) + N_{\text{offline}}(s_h, a_h))$. This directly translates to tighter confidence bounds compared to
 856 using only online data. $\square$

857 We now present a stronger version of the lemma that holds uniformly for all policies π .

858 **Lemma 34** (Uniform Moment Transition Difference Error). *Consider the transition count estimator*
 859 $\hat{P}_t$ *from Equation (9). Further assume the trajectory data follows a martingale structure adapted to*
 860 *the natural filtration of the problem. For any scalar function $f : \mathcal{T} \rightarrow \mathbb{R}$ such that $|f(\tau)| < \eta$ and*
 861 *for any $\epsilon > 0$, with probability at least $1 - \delta$ for all $t \in \mathbb{N}$ and all $\pi \in \Pi$:*

$$\mathbb{E}_{\mathbb{P}_{\hat{P}_t}^\pi} [f(\tau)] - \mathbb{E}_{\mathbb{P}_{P^*}^\pi} [f(\tau)] \leq \underbrace{\mathbb{E}_{\mathbb{P}_{P^*}^\pi} \left[\sum_{h \in [H]} \bar{\xi}_{s_h, a_h}^t(\eta, \delta, \epsilon) \right]}_{=: B_t(\pi, \eta, \delta, \epsilon)} + \epsilon$$

862 where

$$\bar{\xi}_{s_h, a_h}^t(\eta, \delta, \epsilon) := \min \left(2\eta, 4\eta \sqrt{\frac{H \log(|\mathcal{S}| \cdot |\mathcal{A}|) + |\mathcal{S}| \log \left(\left\lceil \frac{4\eta H}{\epsilon} \right\rceil \right) + \log \left(\frac{6 \log(N_t(s_h, a_h) + N_{\text{offline}}(s_h, a_h))}{\delta} \right)}{N_t(s_h, a_h) + N_{\text{offline}}(s_h, a_h)}} \right)$$

863 *Proof.* The proof follows by applying similar techniques as in Lemma 33 but with additional care to
 864 ensure uniformity across all policies.

865 As before, we can artificially treat the offline data as adapted to the natural filtration. The uniform
 866 convergence over the policy class Π is achieved by applying a covering argument and the union
 867 bound, following the methodology in Chatterji et al. [2021] (Lemma B.2). The additional term
 868 $|\mathcal{S}| \log \left(\left\lceil \frac{4\eta H}{\epsilon} \right\rceil \right)$ appears due to this covering, which introduces an ϵ -discretization of the policy
 869 space.

870 The combined offline and online counts in the denominator $(N_t(s_h, a_h) + N_{\text{offline}}(s_h, a_h))$ provide
 871 tighter uniform confidence bounds compared to using online data alone. $\square$

872 To provide further intuition, we elaborate on the meaning and significance of the terms $\hat{B}_t$ and B_t
 873 introduced in the previous lemmas. In reinforcement learning literature, these would be referred to as
 874 the "empirical bonus" and "true bonus," respectively. Both terms quantify the concentration of our
 875 estimators around their true values.

876 The empirical bonus $\hat{B}_t(\pi, \eta, \delta)$ represents the expected sum of state-action-level uncertainty terms
 877 $\xi_{s_h, a_h}^t(\eta, \delta)$ under the *estimated* transition model $\hat{P}_t$. Importantly, this term can be directly computed
 878 from observed data.

879 In contrast, the true bonus $B_t(\pi, \eta, \delta, \epsilon)$ represents the expected sum of uncertainty terms
 880 $\bar{\xi}_{s_h, a_h}^t(\eta, \delta, \epsilon)$ under the *true* transition model P^* . This term cannot be directly computed as it
 881 depends on the unknown true model.

882 For our regret analysis, we need to relate these two quantities. The following lemma provides a
 883 crucial connection, showing that the empirical bonus $\hat{B}_t$ can be bounded in terms of the true bonus
 884 B_t uniformly across all policies π .

885 **Lemma 35** (Relationship Between Empirical and True Bonus Terms). *Let $\eta, \epsilon > 0$. For all policies*
 886 *$\pi \in \Pi$ simultaneously and for all $t \in \mathbb{N}$, with probability at least $1 - \delta$:*

$$\hat{B}_t(\pi, \eta, \delta) \leq 2B_t(\pi, 2H\eta, \delta, \epsilon) + \epsilon$$

887 *Proof.* Define the function $f : \mathcal{T} \rightarrow \mathbb{R}$ as:

$$f(\tau) := \sum_{h \in [H]} \xi_{s_h, a_h}^t(\eta, \delta)$$

888 By construction, $\hat{B}_t(\pi, \eta, \delta) = \mathbb{E}_{\mathbb{P}_{\hat{P}_t}^\pi} [f(\tau)]$. Since $\xi_{s_h, a_h}^t(\eta, \delta) \leq 2\eta$ for all state-action pairs, we
 889 have $|f(\tau)| \leq 2\eta H$.

890 Applying Lemma [34](#) with this $f(\tau)$ and the bound $2\eta H$:

$$\mathbb{E}_{\mathbb{P}_{\hat{P}_t}^\pi} [f(\tau)] - \mathbb{E}_{\mathbb{P}_{P^*}^\pi} [f(\tau)] \leq \mathbb{E}_{\mathbb{P}_{P^*}^\pi} \left[\sum_{h \in [H]} \bar{\xi}_{s_h, a_h}^t(2\eta H, \delta, \epsilon) \right] + \epsilon$$

891 By definition, the right-hand side equals $B_t(\pi, 2H\eta, \delta, \epsilon) + \epsilon$. Therefore:

$$\begin{aligned} \hat{B}_t(\pi, \eta, \delta) &= \mathbb{E}_{\mathbb{P}_{\hat{P}_t}^\pi} [f(\tau)] \\ &\leq \mathbb{E}_{\mathbb{P}_{P^*}^\pi} [f(\tau)] + B_t(\pi, 2H\eta, \delta, \epsilon) + \epsilon \end{aligned}$$

892 From Lemma [33](#) we know that:

$$\mathbb{E}_{\mathbb{P}_{P^*}^\pi} [f(\tau)] \leq \mathbb{E}_{\mathbb{P}_{\hat{P}_t}^\pi} [f(\tau)] + \hat{B}_t(\pi, \eta, \delta) = \hat{B}_t(\pi, \eta, \delta) + \hat{B}_t(\pi, \eta, \delta) = 2\hat{B}_t(\pi, \eta, \delta)$$

893 This gives us:

$$\begin{aligned} \hat{B}_t(\pi, \eta, \delta) &\leq 2\hat{B}_t(\pi, \eta, \delta) + B_t(\pi, 2H\eta, \delta, \epsilon) + \epsilon \\ \Rightarrow -\hat{B}_t(\pi, \eta, \delta) &\leq B_t(\pi, 2H\eta, \delta, \epsilon) + \epsilon \\ \Rightarrow \hat{B}_t(\pi, \eta, \delta) &\leq B_t(\pi, 2H\eta, \delta, \epsilon) + \epsilon \end{aligned}$$

894 Therefore, the lemma statement follows. $\square$

895 This lemma is instrumental for our regret analysis as it allows us to work with B_t instead of $\hat{B}_t$. The
 896 advantage is that B_t involves expectations with respect to the true transition model P^* , which makes
 897 it more amenable to theoretical analysis. By establishing this relationship, we effectively account for
 898 the transition estimation error and can focus on controlling the difference between empirical and true
 899 moments, which is a more tractable problem in our analytical framework.

900 **C.4 Policy Set Π_t and Proof Lemma [8](#)**

901 Recall that we define the policy set Π_t to draw from in line 7 of Algorithm [1](#) as

$$\begin{aligned} \Pi_t := \{ \pi \in \Pi_{1-\delta}^{\text{offline}} \mid \forall \pi' \in \Pi_{1-\delta}^{\text{offline}} : \\ \langle \phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_t}(\pi'), w_t^{\text{proj}} \rangle + \gamma_t \cdot \|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_t}(\pi')\|_{\bar{\mathbf{V}}_t^{-1}} \\ + \hat{B}_t(\pi, 2WB, \delta') + \hat{B}_t(\pi', 2WB, \delta') \geq 0 \} \end{aligned}$$

902 where $\Pi_{1-\delta}^{\text{offline}}$ is derived in Theorem [5](#). The radius γ_t is defined as

$$\gamma_t = \sqrt{2(4\kappa\beta_t(\delta) + \alpha_{d,T}(\delta))} + 2\sqrt{\sum_{\ell=1}^{t-1} \left(\hat{B}_\ell \left(\pi, 2WB, \frac{\delta'}{8\ell^3|A||S|} \right) \right)^2} + \frac{1}{t}$$

903 Then Lemma [8](#) state that with high probability, $\pi^* \in \Pi_t \quad \forall t \in [T]$

904 *Proof of Lemma [8](#)* We begin by conditioning on the following events:

- 905 • $E_{\text{offline}} = \{ \pi^* \in \Pi_{1-\delta}^{\text{offline}} \}$ from Theorem [5](#)

- 906 • E_{w^*} from Lemma 31 (confidence set for w^*)
- 907 • $E_{\overline{V}_T^{P^*}}$ from Lemma 31 (relation for data matrices)
- 908 • $\overline{\mathcal{E}}_0$ from Lemma 32 (estimated norm relation)
- 909 • $\mathcal{E}_3$ from Lemma 33 (bounds on the bonus terms $\hat{B}_t$)

910 By the union bound, these five events hold simultaneously with probability at least $1 - 5\delta$.

911 By the optimality condition, we have:

$$0 \leq \langle \phi(\pi^*) - \phi(\pi'), w^* \rangle = \langle \mathbb{E}_{\mathbb{P}_{P^*}} \phi(\tau) - \mathbb{E}_{\mathbb{P}_{P^*}} \phi(\tau), w^* \rangle$$

912 Then, by event $\mathcal{E}_3$ and defining $f(\tau) := \langle \phi(\tau), w^* \rangle$, we have from Assumption 3.2 that $|f(\tau)| \leq$
 913 $2WB$, which yields:

$$\langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi'), w^* \rangle + \hat{B}_t(\pi^*, 2WB, \delta/|\mathcal{A}|^{|S|}) + \hat{B}_t(\pi', 2WB, \delta/|\mathcal{A}|^{|S|})$$

914 where the probability parameter accounts for any $\pi' \in \Pi$, which covers the case of the offline
 915 confidence set being the whole policy space (i.e., not having enough offline data for learning).

916 Next, we bound the term:

$$\begin{aligned} \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi'), w^* \rangle &= \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi'), w_t^{\text{proj}} \rangle + \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi'), w^* - w_t^{\text{proj}} \rangle \\ &\leq \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi'), w_t^{\text{proj}} \rangle + \|\phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi')\|_{\overline{V}_t^{-1}} \cdot \|w_t^{\text{proj}} - w^*\|_{\overline{V}_t} \end{aligned}$$

917 We can now use event $\overline{\mathcal{E}}_0$:

$$\|w_t^{\text{proj}} - w^*\|_{\overline{V}_t} \leq \sqrt{2}\|w_t^{\text{proj}} - w^*\|_{\overline{V}_t^{P^*}} + 2\sqrt{\sum_{\ell=1}^{t-1} \left(\hat{B}_\ell \left(\pi, 2WB, \frac{\delta'}{8\ell^3|A|^{|S|}} \right) \right)^2} + \frac{1}{t}$$

918 Using events $E_{w^*} \cap E_{\overline{V}_T^{P^*}}$, we get:

$$\|w_t^{\text{proj}} - w^*\|_{\overline{V}_t} \leq \sqrt{2}(4\kappa\beta_t(\delta) + \alpha_{d,T}(\delta)) + 2\sqrt{\sum_{\ell=1}^{t-1} \left(\hat{B}_\ell \left(\pi, 2WB, \frac{\delta'}{8\ell^3|A|^{|S|}} \right) \right)^2} + \frac{1}{t} =: \gamma_t$$

919 Putting these results together yields that $\pi^* \in \Pi_t$ for all $t \in \mathbb{N}$ under the event E_{offline} .

920 The probability of this event is at least $1 - 5\delta$ by the union bound of all the events we conditioned on.
 921 By rescaling $\delta \mapsto \delta/5$, we obtain the desired result with probability at least $1 - \delta$. $\square$

922 C.5 Regret Bound

923 In this section, we provide a lemma as an intermediate step toward the full proof of the regret analysis
 924 of BRIDGE. This lemma separates the upper bound on the regret into three distinct terms, each of
 925 which we further analyze in Appendix D.

926 **Lemma 36** (Regret Analysis). *Under the following events:*

- 927 • $E_{\text{offline}} = \{\pi^* \in \Pi_{1-\delta}^{\text{offline}}\}$ from Theorem 5
- 928 • E_{w^*} from Lemma 31 (confidence set for w^*)
- 929 • $E_{\overline{V}_T^{P^*}}$ from Lemma 31 (relation for data matrices)
- 930 • $\overline{\mathcal{E}}_0$ from Lemma 32 (estimated norm relation)
- 931 • $\mathcal{E}_3$ from Lemma 33 (bounds on the bonus terms $\hat{B}_t$)

932 the regret of BRIDGE Algorithm 7 is upper bounded by:

$$R_T \leq 2 \cdot \underbrace{\gamma_T}_{\text{Term 1}} \cdot \underbrace{\sqrt{T \sum_{t \in [T]} \|\phi^{\hat{P}_t}(\pi_t^1) - \phi^{\hat{P}_t}(\pi_t^2)\|_{\bar{V}_t^{-1}}}}_{\text{Term 2}} + \underbrace{\sum_{i \in \{1,2\}} \sum_{t \in [T]} \hat{B}_t(\pi_t^i, 4WB, \delta)}_{\text{Term 3}}$$

933 where

$$\gamma_T = \sqrt{2}(4\kappa \cdot \beta_T(\delta) + \alpha_{d,T}(\delta)) + \frac{1}{T} + 4 \sqrt{\sum_{i \in \{1,2\}} \sum_{t \in [T]} B_T(\pi_t^i, 4HWB, \delta, \epsilon)^2 + 96T\epsilon HWB}$$

934 and

935 • $\alpha_{d,T}(\delta) = 20BW \sqrt{d \log(T(1+2T)/\delta)}$

936 • $\beta_T(\delta) = \sqrt{\lambda}W + \sqrt{\log(1/\delta) + 2d \log(1 + \frac{TB^2}{\kappa \lambda d})}$.

937 *Proof.* We start by writing

$$\begin{aligned} 2r_t &= \langle \phi(\pi^*) - \phi(\pi_t^1), w^* \rangle + \langle \phi(\pi^*) - \phi(\pi_t^2), w^* \rangle \\ &= \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1), w^* \rangle + \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2), w^* \rangle + 2\langle \phi^{\hat{P}_t}(\pi^*) - \phi(\pi^*), w^* \rangle \\ &\quad + \langle \phi^{\hat{P}_t}(\pi_t^1) - \phi(\pi_t^1), w^* \rangle + \langle \phi^{\hat{P}_t}(\pi_t^2) - \phi(\pi_t^2), w^* \rangle \end{aligned}$$

938 Then, by Lemma 33, we have with probability at least $1 - \delta$ for each of the following:

$$\begin{aligned} 2\langle \phi(\pi^*) - \phi^{\hat{P}_t}(\pi^*), w^* \rangle &\leq 2\hat{B}_t(\pi^*, 4WB, \delta) \\ \langle \phi^{\hat{P}_t}(\pi_t^1) - \phi(\pi_t^1), w^* \rangle &\leq \hat{B}_t(\pi_t^1, 4WB, \delta) \\ \langle \phi^{\hat{P}_t}(\pi_t^2) - \phi(\pi_t^2), w^* \rangle &\leq \hat{B}_t(\pi_t^2, 4WB, \delta) \end{aligned}$$

939 By the union bound, with high probability:

$$2r_t \leq \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1), w^* \rangle + \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2), w^* \rangle + \hat{B}_t(\pi_t^1, 4WB, \delta) + \hat{B}_t(\pi_t^2, 4WB, \delta) + 2\hat{B}_t(\pi^*, 4WB, \delta)$$

940 Next, we observe that:

$$\begin{aligned} &\langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1), w^* \rangle + \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2), w^* \rangle \\ &\leq \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1), w_t^{\text{proj}} \rangle + \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2), w_t^{\text{proj}} \rangle \\ &\quad + \|w^* - w_t^{\text{proj}}\|_{\bar{V}_t} \left(\|\phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1)\|_{\bar{V}_t^{-1}} + \|\phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2)\|_{\bar{V}_t^{-1}} \right) \end{aligned}$$

941 Conditioning on the joint event $\mathcal{E}_0 \cap E_{w^*} \cap E_{\bar{V}_t^{P^*}}$, we have with high probability:

$$\begin{aligned} &\langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1), w^* \rangle + \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2), w^* \rangle \\ &\leq \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1), w_t^{\text{proj}} \rangle + \langle \phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2), w_t^{\text{proj}} \rangle \\ &\quad + \gamma_t \cdot \left(\|\phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1)\|_{\bar{V}_t^{-1}} + \|\phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2)\|_{\bar{V}_t^{-1}} \right) \end{aligned}$$

942 Using Lemma 33 again, the following holds with high probability:

$$\begin{aligned} 2\langle \phi(\pi^*) - \phi^{\hat{P}_t}(\pi^*), w_t^{\text{proj}} \rangle &\leq 2\hat{B}_t(\pi^*, 4WB, \delta) \\ \langle \phi^{\hat{P}_t}(\pi_t^1) - \phi(\pi_t^1), w_t^{\text{proj}} \rangle &\leq \hat{B}_t(\pi_t^1, 4WB, \delta) \\ \langle \phi^{\hat{P}_t}(\pi_t^2) - \phi(\pi_t^2), w_t^{\text{proj}} \rangle &\leq \hat{B}_t(\pi_t^2, 4WB, \delta) \end{aligned}$$

943 Putting everything together yields:

$$2r_t \leq 2\gamma_t \left(\|\phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^1)\|_{\bar{V}_t^{-1}} + \|\phi^{\hat{P}_t}(\pi^*) - \phi^{\hat{P}_t}(\pi_t^2)\|_{\bar{V}_t^{-1}} \right) \\ + 2\hat{B}_t(\pi_t^1, 4WB, \delta) + 2\hat{B}_t(\pi_t^2, 4WB, \delta) + 4\hat{B}_t(\pi^*, 4WB, \delta)$$

944 Under the event $\pi^* \in \Pi_t$ from Lemma 8 and using the fact that $\pi_t^1, \pi_t^2 \in \Pi_t$, we have:

$$2r_t \leq \gamma_t \|\phi^{\hat{P}_t}(\pi_t^2) - \phi^{\hat{P}_t}(\pi_t^1)\|_{\bar{V}_t^{-1}} + 4\hat{B}_t(\pi_t^1, 4WB, \delta) + 4\hat{B}_t(\pi_t^2, 4WB, \delta)$$

945 Hence, the regret is:

$$R_T = \sum_{t \in [T]} 2r_t \\ \leq \sum_{t \in [T]} \left(\gamma_t \|\phi^{\hat{P}_t}(\pi_t^2) - \phi^{\hat{P}_t}(\pi_t^1)\|_{\bar{V}_t^{-1}} + 4\hat{B}_t(\pi_t^1, 4WB, \delta) + 4\hat{B}_t(\pi_t^2, 4WB, \delta) \right) \\ \leq \gamma_T \sqrt{T \sum_{t \in [T]} \|\phi^{\hat{P}_t}(\pi_t^2) - \phi^{\hat{P}_t}(\pi_t^1)\|_{\bar{V}_t^{-1}}^2} + \sum_{t \in [T]} \left(4\hat{B}_t(\pi_t^1, 4WB, \delta) + 4\hat{B}_t(\pi_t^2, 4WB, \delta) \right)$$

946 Note that by Lemma 33 with high probability:

$$\hat{B}_t(\pi_t^i, 2WB, \delta)^2 \leq 4B_t(\pi_t^i, 2HWB, \delta, \epsilon) + 24\epsilon HWB$$

947 Plugging this into γ_t yields:

$$\gamma_t \leq \sqrt{2(4\kappa\beta_t(\delta) + \alpha_{d,T}(\delta))} + \frac{1}{t} \\ + 4\sqrt{\sum_{\ell=1}^{t-1} B_t^2(\pi_t^1, 4HSB, \delta'_t, \epsilon) + B_t^2(\pi_t^2, 4HSB, \delta'_t) + 96(t-1)HWB} \quad \forall t$$

948 This completes the proof of the claimed result. $\square$

949 D Regret Analysis: Theorem 9

950 In this section, we present the complete regret analysis of our BRIDGE algorithm. We recommend
951 that readers first review Appendix A, where we analyze a simplified setting in which the dynamics
952 are assumed to be known. This simplified case captures the core idea of our approach: constraining
953 the set of policies considered during online preference learning using a confidence interval derived
954 from offline behavioral cloning estimation (see ??).

955 The key difference in the present analysis is that we now incorporate the estimation of the transition
956 model. Specifically, we first estimate the transition model offline and then use this estimate as the
957 starting point for online transition estimation. This approach reduces the error due to transition
958 uncertainty by a factor of $\mathcal{O}(1/\sqrt{n})$, which is the same rate of improvement we achieve for the policy
959 estimation through behavioral cloning. As we will show, this allows our algorithm to effectively
960 leverage offline demonstrations to reduce both sources of uncertainty, resulting in substantially
961 improved regret bounds.

962 **Theorem 37** (Regret Bound with Offline-Enhanced Exploration). *Let n be the number of offline*
963 *demonstrations with minimum visitation probability $\gamma_{\min} > 0$ for state-action pairs. With probability*
964 *at least $1 - \delta$, the regret of the algorithm is bounded by:*

$$\begin{aligned}
R_T \leq & 2 \cdot \underbrace{\gamma_T}_{\text{Term 1}} \cdot \underbrace{\sqrt{T \cdot \log \left(1 + \frac{\tilde{O} \left(B^2 \cdot H \cdot |S|^2 \cdot \min \left\{ \frac{T}{n}, \ln(T) \right\} + \frac{T \cdot |S| \cdot B^2 \cdot \sqrt{|A| \cdot H}}{\sqrt{n \cdot \gamma_{\min}}} \right)}{d} \right)}}_{\text{Term 2}} \\
& + \underbrace{\tilde{O} \left(H|S| \sqrt{\frac{|A|TH}{n \cdot \gamma_{\min}}} + \frac{H^{5/2}WB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2WB \cdot \sqrt{T} \cdot \frac{|S|^{1/2}|A|^{1/4}}{n^{1/4}} \right)}_{\text{Term 3}}
\end{aligned}$$

965 where

$$\gamma_T = \tilde{O} \left((\kappa + BW) \sqrt{d \log(T)} + H^2WB|S| \cdot \sqrt{\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} + \sqrt{HWB} \right)$$

966 and we have set $\epsilon = \frac{1}{T}$ to optimize the bound.

967 From Lemma 36, we analyze the three key terms in our regret bound: the confidence multiplier (Term
968 1), the logarithmic determinant ratio (Term 2), and the bonus function summation (Term 3). Each
969 term is examined in detail in the following subsections.

$$\text{Term 1} = \gamma_T = \sqrt{2(4\kappa \cdot \beta_T(\delta) + \alpha_{d,T}(\delta))} + \frac{1}{T} + 4 \sqrt{\sum_{i \in \{1,2\}} \sum_{t \in [T]} B_T(\pi_t^i, 4HWB, \delta, \epsilon)^2 + 96T\epsilon HWB}$$

$$\text{Term 2} = \sqrt{T \sum_{t \in [T]} \|\phi^{\hat{P}_t}(\pi_t^1) - \phi^{\hat{P}_t}(\pi_t^2)\|_{V_t^{-1}}^2}$$

$$\text{Term 3} = \sum_{i \in \{1,2\}} \sum_{t \in [T]} \hat{B}_t(\pi_t^i, 4WB, \delta)$$

970 D.1 Term 1: Asymptotic bound

971 We derive an asymptotic bound for Term 1 in Theorem 37 via Lemma 38. The auxiliary lemmata
972 used in the proof of Lemma 38 are found in Appendix D.1.1.

Lemma 38. *The asymptotic bound on γ_T can be expressed as:*

$$\gamma_T = \tilde{O} \left((\kappa + BW) \sqrt{d \log(T)} + H^2WB|S| \cdot \sqrt{\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} + \sqrt{T\epsilon HWB} \right)$$

973 *Proof.* We analyze each term in the expression for γ_T separately.

974 **Step 1: Analyze** $\sqrt{2(4\kappa \cdot \beta_T(\delta) + \alpha_{d,T}(\delta))}$

975 Given:

$$\begin{aligned}
\alpha_{d,T}(\delta) &= 20BW \sqrt{d \log(T(1+2T)/\delta)} \\
\beta_T(\delta) &= \sqrt{\lambda}W + \sqrt{\log(1/\delta) + 2d \log \left(1 + \frac{TB^2}{\kappa \lambda d} \right)}
\end{aligned}$$

976 For $\alpha_{d,T}(\delta)$, we have:

$$\begin{aligned}
\alpha_{d,T}(\delta) &= 20BW \sqrt{d \log(T(1+2T)/\delta)} \\
&= 20BW \sqrt{d \log(T) + d \log(1+2T) - d \log(\delta)} \\
&= 20BW \sqrt{d(\log(T) + \log(1+2T) - \log(\delta))} \\
&= O(BW \sqrt{d \log(T/\delta)})
\end{aligned}$$

977 For $\beta_T(\delta)$, we have:

$$\begin{aligned}
\beta_T(\delta) &= \sqrt{\lambda}W + \sqrt{\log(1/\delta) + 2d \log \left(1 + \frac{TB^2}{\kappa\lambda d} \right)} \\
&= \sqrt{\lambda}W + \sqrt{\log(1/\delta) + 2d \log \left(\frac{\kappa\lambda d + TB^2}{\kappa\lambda d} \right)} \\
&= \sqrt{\lambda}W + \sqrt{\log(1/\delta) + 2d \log \left(1 + \frac{TB^2}{\kappa\lambda d} \right)} \\
&\leq \sqrt{\lambda}W + \sqrt{\log(1/\delta) + 2d \log \left(\frac{2TB^2}{\kappa\lambda d} \right)} \quad (\text{for large enough } T) \\
&= \sqrt{\lambda}W + \sqrt{\log(1/\delta) + 2d \log(T) + 2d \log \left(\frac{2B^2}{\kappa\lambda d} \right)} \\
&= O(\sqrt{\lambda}W + \sqrt{d \log(T) + \log(1/\delta)})
\end{aligned}$$

978 Therefore, this term becomes:

$$\begin{aligned}
\sqrt{2}(4\kappa \cdot \beta_T(\delta) + \alpha_{d,T}(\delta)) &= O(\kappa \cdot (\sqrt{\lambda}W + \sqrt{d \log(T) + \log(1/\delta)}) + BW \sqrt{d \log(T/\delta)}) \\
&= O(\kappa \sqrt{\lambda}W + \kappa \sqrt{d \log(T) + \log(1/\delta)} + BW \sqrt{d \log(T/\delta)}) \\
&= O((\kappa + BW) \sqrt{d \log(T)} + \kappa \sqrt{\log(1/\delta)} + BW \sqrt{d \log(1/\delta)})
\end{aligned}$$

979 For a fixed confidence parameter δ , this simplifies to:

$$\sqrt{2}(4\kappa \cdot \beta_T(\delta) + \alpha_{d,T}(\delta)) = O((\kappa + BW) \sqrt{d \log(T)})$$

980 **Step 2: Analyze $\frac{1}{T}$**

981 This term is $O(\frac{1}{T})$ and becomes negligible for large T compared to other terms.

982 **Step 3: Analyze $4\sqrt{\sum_{i \in \{1,2\}} \sum_{t \in [T]} B_T(\pi_t^i, 4HWB, \delta, \epsilon)^2 + 96T\epsilon HWB}$**

983 Using the provided lemma on the sum of squared bonus terms, Lemma [41](#):

$$\begin{aligned}
\sum_{i \in \{1,2\}} \sum_{t \in [T]} B_T(\pi_t^i, 4HWB, \delta, \epsilon)^2 &\leq \tilde{O} \left((4HWB)^2 H^2 |S|^2 \cdot \min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\} \right) \\
&= \tilde{O} \left(16H^2 W^2 B^2 \cdot H^2 |S|^2 \cdot \min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\} \right) \\
&= \tilde{O} \left(16H^4 W^2 B^2 |S|^2 \cdot \min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\} \right)
\end{aligned}$$

984 For the second term inside the square root:

$$96T\epsilon HWB = O(T\epsilon HWB)$$

985 Therefore:

$$\begin{aligned}
& 4 \sqrt{\sum_{i \in \{1,2\}} \sum_{t \in [T]} B_T(\pi_t^i, 4HWB, \delta, \epsilon)^2 + 96T\epsilon HWB} \\
&= 4 \sqrt{\tilde{O} \left(16H^4 W^2 B^2 |S|^2 \cdot \min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\} \right) + O(T\epsilon HWB)} \\
&= \tilde{O} \left(4 \sqrt{16H^4 W^2 B^2 |S|^2 \cdot \min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} \right) + O(4\sqrt{T\epsilon HWB}) \\
&= \tilde{O} \left(16H^2 WB |S| \cdot \sqrt{\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} \right) + O(\sqrt{T\epsilon HWB}) \\
&= \tilde{O} \left(H^2 WB |S| \cdot \sqrt{\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} \right) + O(\sqrt{T\epsilon HWB})
\end{aligned}$$

986 **Step 4: Combine all terms**

987 Combining all terms from Steps 1-3, we get:

$$\begin{aligned}
\gamma_T &= O((\kappa + BW)\sqrt{d \log(T)}) + O\left(\frac{1}{T}\right) + \tilde{O} \left(H^2 WB |S| \cdot \sqrt{\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} \right) + O(\sqrt{T\epsilon HWB}) \\
&= O((\kappa + BW)\sqrt{d \log(T)}) + \tilde{O} \left(H^2 WB |S| \cdot \sqrt{\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} \right) + O(\sqrt{T\epsilon HWB}) + o(1)
\end{aligned}$$

988 Expressing this with $\tilde{O}$ notation to hide logarithmic factors:

$$\gamma_T = \tilde{O} \left((\kappa + BW)\sqrt{d \log(T)} + H^2 WB |S| \cdot \sqrt{\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\}} + \sqrt{T\epsilon HWB} \right)$$

989 □

990 **D.1.1 Term 1 asymptotic bound: auxiliary lemmata for Lemma 38**

991 **Lemma 39** (Offline-Enhanced Bonus Term Bound). *Let n be the number of offline demonstrations,*
992 *with a minimum visitation probability $\gamma_{\min} > 0$ for state-action pairs visited by the expert policy π^* .*
993 *Then, with probability at least $1 - 2\delta'$, the sum of squared bonus terms satisfies:*

$$\begin{aligned}
\sum_{t \in [T]} \sum_{h=1}^{H-1} \left(\xi_{s_t, h, a_t, h}^{(t)}(\epsilon, \eta, \delta) \right)^2 &\leq 32\eta^2 \left(H \log(|S||A|H) + |S| \log \left(\frac{4\eta H}{\epsilon} \right) + \log \left(\frac{6 \log(HT)}{\delta'} \right) \right) \\
&\quad \cdot |S_{\text{reach}}| \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right)
\end{aligned}$$

994 where $|S_{\text{reach}}|$ is the number of state-action pairs with non-zero visitation probability under the expert
995 policy.

996 **Proof. Step 1: Express Modified Bonus Terms with Offline Data.** We define our modified bonus
997 term to incorporate offline data:

$$\xi_{s,a}^{(t)}(\epsilon, \eta, \delta) = \min \left(2\eta, 4\eta \sqrt{\frac{U}{N_{\text{off}}(s, a) + N_t(s, a)}} \right)$$

998 where $U = H \log(|S||A|H) + |S| \log\left(\frac{4\eta H}{\epsilon}\right) + \log\left(\frac{6 \log(t)}{\delta}\right)$

999 **Step 2: Express the Sum of Squared Bonus Terms.** Following Pacchiano's structure but with our
 1000 modified bonus terms:

$$\sum_{t \in [T]} \sum_{h=1}^{H-1} \left(\xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, \eta, \delta) \right)^2 = \sum_{s \in S} \sum_{a \in A} \sum_{t=1}^{N_{T+1}(s,a)} \min \left(4\eta^2, 16\eta^2 \frac{H \log(|S||A|H) + |S| \log\left(\frac{4\eta H}{\epsilon}\right) + \log\left(\frac{6 \log(t)}{\delta}\right)}{N_{\text{off}}(s, a) + t} \right)$$

1001 **Step 3: Rearrange to Account for Offline Data.** The key insight: With offline data, we need to
 1002 adjust the indices of summation. For each state-action pair, we've already observed it $N_{\text{off}}(s, a)$ times
 1003 in the offline dataset. Therefore:

$$\sum_{t \in [T]} \sum_{h=1}^{H-1} \left(\xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, \eta, \delta) \right)^2 = \sum_{s \in S} \sum_{a \in A} \sum_{t'=N_{\text{off}}(s,a)+1}^{N_{\text{off}}(s,a)+N_{T+1}(s,a)} \min \left(4\eta^2, 16\eta^2 \frac{H \log(|S||A|H) + |S| \log\left(\frac{4\eta H}{\epsilon}\right) + \log\left(\frac{6 \log(t')}{\delta}\right)}{t'} \right)$$

1004 where t' represents the total count (offline + online).

1005 **Step 4: Simplify Using Common Term.** For clarity and following [Saha et al. \[2023\]](#) approach, let's
 1006 define:

$$V = H \log(|S||A|H) + |S| \log\left(\frac{4\eta H}{\epsilon}\right) + \log\left(\frac{6 \log(HT)}{\delta'}\right)$$

1007 For sufficiently large t' , the min is dominated by the second term:

$$\sum_{s \in S} \sum_{a \in A} \sum_{t'=N_{\text{off}}(s,a)+1}^{N_{\text{off}}(s,a)+N_{T+1}(s,a)} 16\eta^2 \frac{V}{t'} = 16\eta^2 \cdot V \cdot \sum_{s \in S} \sum_{a \in A} \sum_{t'=N_{\text{off}}(s,a)+1}^{N_{\text{off}}(s,a)+N_{T+1}(s,a)} \frac{1}{t'}$$

1008 **Step 5: Use the Harmonic Sum Property.** We know that $\sum_{i=a+1}^b \frac{1}{i} \leq \log\left(\frac{b}{a}\right)$. Therefore:

$$\sum_{t'=N_{\text{off}}(s,a)+1}^{N_{\text{off}}(s,a)+N_{T+1}(s,a)} \frac{1}{t'} \leq \log\left(\frac{N_{\text{off}}(s,a) + N_{T+1}(s,a)}{N_{\text{off}}(s,a)}\right) = \log\left(1 + \frac{N_{T+1}(s,a)}{N_{\text{off}}(s,a)}\right)$$

1009 **Step 6: Apply the Minimum Visitation Probability.** With our assumption that $d_{\mathcal{P}^*}^{\pi, t}(s, a) \geq \gamma_{\min}$
 1010 for all state-action pairs visited by the expert policy, we have:

$$N_{\text{off}}(s, a) \geq n \cdot H \cdot \gamma_{\min} \quad \forall (s, a) \in S_{\text{reach}}$$

1011 where S_{reach} is the set of state-action pairs with non-zero visitation probability under the expert policy.

1012 Therefore:

$$\log\left(1 + \frac{N_{T+1}(s,a)}{N_{\text{off}}(s,a)}\right) \leq \log\left(1 + \frac{N_{T+1}(s,a)}{n \cdot H \cdot \gamma_{\min}}\right) \quad \forall (s, a) \in S_{\text{reach}}$$

1013 **Step 7: Apply Jensen's Inequality.** We know $\sum_{s,a} N_{T+1}(s, a) = TH$ (total state-action visits in
 1014 online learning).

1015 By Jensen's inequality and the concavity of $\log(1+x)$:

$$\sum_{(s,a) \in S_{\text{reach}}} \log \left(1 + \frac{N_{T+1}(s, a)}{n \cdot H \cdot \gamma_{\min}} \right) \leq |S_{\text{reach}}| \cdot \log \left(1 + \frac{\sum_{(s,a) \in S_{\text{reach}}} N_{T+1}(s, a)}{|S_{\text{reach}}| \cdot n \cdot H \cdot \gamma_{\min}} \right)$$

1016 Since $\sum_{(s,a) \in S_{\text{reach}}} N_{T+1}(s, a) \leq TH$:

$$\sum_{(s,a) \in S_{\text{reach}}} \log \left(1 + \frac{N_{T+1}(s, a)}{n \cdot H \cdot \gamma_{\min}} \right) \leq |S_{\text{reach}}| \cdot \log \left(1 + \frac{TH}{|S_{\text{reach}}| \cdot n \cdot H \cdot \gamma_{\min}} \right)$$

1017 Simplifying:

$$\sum_{(s,a) \in S_{\text{reach}}} \log \left(1 + \frac{N_{T+1}(s, a)}{n \cdot H \cdot \gamma_{\min}} \right) \leq |S_{\text{reach}}| \cdot \log \left(1 + \frac{T}{|S_{\text{reach}}| \cdot n \cdot \gamma_{\min}} \right)$$

1018 For unreachable states, we can use Pacchiano's original bound, but these contribute negligibly to
 1019 regret as optimal policies don't visit them.

1020 **Step 8: Final Bound.** Substituting back:

$$\sum_{t \in [T]} \sum_{h=1}^{H-1} \left(\xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, \eta, \delta) \right)^2 \leq 16\eta^2 \cdot V \cdot |S_{\text{reach}}| \cdot \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right)$$

1021 Substituting V and accounting for approximation constants:

$$\begin{aligned} \sum_{t \in [T]} \sum_{h=1}^{H-1} \left(\xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, \eta, \delta) \right)^2 &\leq 32\eta^2 \left(H \log(|S||A|H) + |S| \log \left(\frac{4\eta H}{\epsilon} \right) + \log \left(\frac{6 \log(HT)}{\delta'} \right) \right) \\ &\quad \cdot |S_{\text{reach}}| \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right) \end{aligned}$$

1022 This completes our proof, showing explicitly how offline data (through n) and minimum visitation
 1023 probability $\gamma_{\min}$ reduce the bound on bonus terms, thereby reducing regret. $\square$

1024 **Lemma 40** (Offline-Enhanced Squared Bonus Term Bound). *Let $\eta, \epsilon > 0$ and $\delta, \delta' \in (0, 1)$. Let n be
 1025 the number of offline demonstrations with minimum visitation probability $\gamma_{\min} > 0$ for state-action
 1026 pairs visited by the expert policy. Define $\mathcal{E}_5(\delta')$ be the event that for all $t \in \mathbb{N}$ and $i \in \{1, 2\}$:*

$$\begin{aligned} \sum_{i \in \{1, 2\}} \sum_{\ell=1}^{t-1} \left(B_{\ell}(\pi_{\ell}^i, \eta, \delta/\ell^3, \epsilon) \right)^2 &\leq 12\eta^2 H^2 \left(1.4 \ln \ln (2 (\max(4\eta^2 H t, 1))) + \ln \frac{5.2}{\delta'} + 1 \right) \\ &\quad + 64\eta^2 H |S_{\text{reach}}| \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right) \\ &\quad \cdot \left(H \log(|S||A|H) + |S| \log \left(\left\lceil \frac{4\eta H}{\epsilon} \right\rceil \right) + \log \left(\frac{6 \log(HT)}{\delta} \right) \right) \end{aligned}$$

1027 Then $\mathbb{P}(\mathcal{E}_5(\delta')) \geq 1 - 2\delta'$.

1028 *Proof.* We follow [Saha et al. \[2023\]](#) proof structure, beginning with the martingale analysis and then
 1029 applying our offline-enhanced bounds.

1030 Observe that the bonus terms can be expressed as:

$$\begin{aligned} \left(B_\ell(\pi_\ell^1, \eta, \frac{\delta}{\ell^3}, \epsilon) \right)^2 + \left(B_\ell(\pi_\ell^2, \eta, \frac{\delta}{\ell^3}, \epsilon) \right)^2 &= \left(\mathbb{E}_{s_1^1 \sim \rho, \tau \sim \mathbb{P}_{\hat{P}_\ell}^{\pi_\ell^1}(\cdot | s_1^1)} \left[\sum_{h=1}^{H-1} \xi_{s_h^1, a_h^1}^{(\ell)}(\epsilon, \eta, \frac{\delta}{\ell^3}) \right] \right)^2 \\ &\quad + \left(\mathbb{E}_{s_1^2 \sim \rho, \tau \sim \mathbb{P}_{\hat{P}_\ell}^{\pi_\ell^2}(\cdot | s_1^2)} \left[\sum_{h=1}^{H-1} \xi_{s_h^2, a_h^2}^{(\ell)}(\epsilon, \eta, \frac{\delta}{\ell^3}) \right] \right)^2 \end{aligned}$$

1031 Using Jensen's inequality (as in the original proof):

$$\left(\mathbb{E}_{s_1^i \sim \rho, \tau \sim \mathbb{P}_{\hat{P}_\ell}^{\pi_\ell^i}(\cdot | s_1^i)} \left[\sum_{h=1}^{H-1} \xi_{s_h^i, a_h^i}^{(\ell)}(\epsilon, \eta, \frac{\delta}{\ell^3}) \right] \right)^2 \leq H \mathbb{E}_{s_1^i \sim \rho, \tau \sim \mathbb{P}_{\hat{P}_\ell}^{\pi_\ell^i}(\cdot | s_1^i)} \left[\sum_{h=1}^{H-1} \left(\xi_{s_h^i, a_h^i}^{(\ell)}(\epsilon, \eta, \frac{\delta}{\ell^3}) \right)^2 \right]$$

1032 Following the martingale analysis of Pacchiano, we define:

$$D_\ell^{(i)} = \mathbb{E}_{s_1^i \sim \rho, \tau \sim \mathbb{P}_{\hat{P}_\ell}^{\pi_\ell^i}(\cdot | s_1^i)} \left[\sum_{h=1}^{H-1} \left(\xi_{s_h^i, a_h^i}^{(\ell)}(\epsilon, \eta, \delta) \right)^2 \right] - \sum_{h=1}^{H-1} \left(\xi_{s_h^i, a_h^i}^{(\ell)}(\epsilon, \eta, \delta) \right)^2$$

1033 Since $\xi_{s,a}^{(\ell)}(\epsilon, \eta, \delta) \leq 2\eta$, we have $|D_\ell^{(i)}| \leq 8\eta^2 H$ and $\text{Var}_\ell^{(i)} \left(\sum_{h=1}^{H-1} \left(\xi_{s_h^i, a_h^i}^{(\ell)}(\epsilon, \eta, \delta) \right)^2 \right) \leq$
 1034 $16\eta^4 H^2$.

1035 Applying the Uniform Empirical Bernstein Bound (as in the original proof), we get:

$$\begin{aligned} \sum_{\ell=1}^{t-1} D_\ell^{(i)} &\leq \frac{1}{2} \mathbb{E}_{s_1^i \sim \rho, \tau \sim \mathbb{P}_{\hat{P}_\ell}^{\pi_\ell^i}(\cdot | s_1^i)} \left[\sum_{h=1}^{H-1} \left(\xi_{s_h^i, a_h^i}^{(\ell)}(\epsilon, \eta, \delta) \right)^2 \right] \\ &\quad + 6\eta^2 H \left(1.4 \ln \ln (2 (\max (4\eta^2 H t, 1))) + \ln \frac{5.2}{\delta'} \right) \end{aligned}$$

1036 Therefore, with high probability for $i \in \{1, 2\}$:

$$\begin{aligned} \mathbb{E}_{s_1^i \sim \rho, \tau \sim \mathbb{P}_{\hat{P}_\ell}^{\pi_\ell^i}(\cdot | s_1^i)} \left[\sum_{h=1}^{H-1} \left(\xi_{s_h^i, a_h^i}^{(\ell)}(\epsilon, \eta, \delta) \right)^2 \right] &\leq 2 \sum_{\ell=1}^{t-1} \sum_{h=1}^{H-1} \left(\xi_{s_h^i, a_h^i}^{(\ell)}(\epsilon, \eta, \delta) \right)^2 + 4\eta^2 H \\ &\quad + 6\eta^2 H \left(1.4 \ln \ln (2 (\max (4\eta^2 H t, 1))) + \ln \frac{5.2}{\delta'} \right) \end{aligned}$$

1037 Combining for both policies, with probability $1 - 2\delta'$:

$$\begin{aligned} \sum_{i \in \{1, 2\}} \sum_{\ell=1}^{t-1} (B_\ell(\pi_\ell^i, \eta, \delta/\ell^3))^2 &\leq 2H \sum_{i \in \{1, 2\}} \sum_{\ell=1}^{t-1} \sum_{h=1}^{H-1} \left(\xi_{s_h^i, a_h^i}^{(\ell, i)}(\epsilon, \eta, \delta) \right)^2 \\ &\quad + 12\eta^2 H^2 \left(1.4 \ln \ln (2 (\max (4\eta^2 H t, 1))) + \ln \frac{5.2}{\delta'} + 1 \right) \end{aligned}$$

1038 Now, using Lemma [39](#), we have:

$$\sum_{\ell=1}^{t-1} \sum_{h=1}^{H-1} \left(\xi_{s_h, a_h}^{(\ell, i)}(\epsilon, \eta, \delta) \right)^2 \leq 16\eta^2 V \cdot |S_{\text{reach}}| \cdot \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right)$$

1039 where $V = H \log(|S||A|H) + |S| \log \left(\left\lceil \frac{4\eta H}{\epsilon} \right\rceil \right) + \log \left(\frac{6 \log(HT)}{\delta} \right)$.

1040 Substituting this bound and combining terms:

$$\sum_{i \in \{1,2\}} \sum_{\ell=1}^{t-1} (B_\ell(\pi_\ell^i, \eta, \delta/\ell^3))^2 \leq 12\eta^2 H^2 \left(1.4 \ln \ln (2 (\max(4\eta^2 H t, 1))) + \ln \frac{5.2}{\delta'} + 1 \right) \\ + 64\eta^2 H V \cdot |S_{\text{reach}}| \cdot \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right)$$

1041 Expanding V :

$$\sum_{i \in \{1,2\}} \sum_{\ell=1}^{t-1} (B_\ell(\pi_\ell^i, \eta, \delta/\ell^3))^2 \leq 12\eta^2 H^2 \left(1.4 \ln \ln (2 (\max(4\eta^2 H t, 1))) + \ln \frac{5.2}{\delta'} + 1 \right) \\ + 64\eta^2 H \cdot |S_{\text{reach}}| \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right) \cdot \\ \left(H \log(|S||A|H) + |S| \log \left(\left\lceil \frac{4\eta H}{\epsilon} \right\rceil \right) + \log \left(\frac{6 \log(HT)}{\delta} \right) \right)$$

1042

□

1043 **Lemma 41** (Asymptotic Bound for Offline-Enhanced Squared Bonus Terms). *With n offline demon-*
1044 *strations and minimum visitation probability $\gamma_{\min}$, the sum of squared bonus terms is bounded*
1045 *as:*

$$\sum_{i \in \{1,2\}} \sum_{\ell=1}^{t-1} (B_\ell(\pi_\ell^i, \eta, \delta/\ell^3, \epsilon))^2 \leq \tilde{O} \left(\eta^2 H^2 |S|^2 \cdot \min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\} \right)$$

1046 where $\tilde{O}(\cdot)$ hides logarithmic factors in H , $|S|$, $|A|$, δ^{-1} , and ϵ^{-1} , as well as constant factors.

1047 *Proof.* We start from the detailed bound of Lemma 40:

$$\sum_{i \in \{1,2\}} \sum_{\ell=1}^{t-1} (B_\ell(\pi_\ell^i, \eta, \delta/\ell^3))^2 \leq 12\eta^2 H^2 \left(1.4 \ln \ln (2 (\max(4\eta^2 H t, 1))) + \ln \frac{5.2}{\delta'} + 1 \right) \\ + 64\eta^2 H \left(H \log(|S||A|H) + |S| \log \left(\left\lceil \frac{4\eta H}{\epsilon} \right\rceil \right) + \log \left(\frac{6 \log(HT)}{\delta} \right) \right) |S_{\text{reach}}| \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right)$$

1048 Analyzing each term:

1049 **Step 1: First term analysis.** The first term is:

$$12\eta^2 H^2 \left(1.4 \ln \ln (2 (\max(4\eta^2 H t, 1))) + \ln \frac{5.2}{\delta'} + 1 \right) = O(\eta^2 H^2 \log \log(T))$$

1050 Since $\log \log(T)$ grows extremely slowly, and we're using $\tilde{O}$ notation which hides logarithmic factors,
1051 this term is dominated by $\tilde{O}(\eta^2 H^2)$.

1052 **Step 2: Second term analysis.** For the second term, we have:

$$C \cdot \eta^2 H \cdot V \cdot |S_{\text{reach}}| \cdot \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right)$$

1053 where C is a constant and $V = \left(H \log(|S||A|H) + |S| \log \left(\left\lceil \frac{4\eta H}{\epsilon} \right\rceil \right) + \log \left(\frac{6 \log(HT)}{\delta} \right) \right)$.

1054 Within the factor V , the dominant term is $|S| \log \left(\left\lceil \frac{4\eta H}{\epsilon} \right\rceil \right)$ since it scales with $|S|$. Therefore,
1055 asymptotically:

$$V = \tilde{O}(|S|)$$

1056 Upper bounding $|S_{\text{reach}}| \leq |S|$ as requested, the second term becomes:

$$\tilde{O} \left(\eta^2 H^2 |S|^2 \cdot \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right) \right)$$

1057 **Step 3: Analysis of $\log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right)$.** We need to consider different regimes for this logarithmic
1058 term:

1059 **Case 1:** Small offline dataset ($n \cdot \gamma_{\min} \ll T$)

$$\begin{aligned} \log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right) &\approx \log \left(\frac{T}{n \cdot \gamma_{\min}} \right) \\ &= \log(T) - \log(n \cdot \gamma_{\min}) \\ &= O(\log(T)) \end{aligned}$$

1060 **Case 2:** Balanced regime ($n \cdot \gamma_{\min} \approx T$)

$$\log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right) \approx \log \left(1 + \frac{1}{\gamma_{\min}} \right) = O(1)$$

1061 **Case 3:** Large offline dataset ($n \cdot \gamma_{\min} \gg T$)

1062 Here we can use the approximation $\log(1+x) \approx x$ for small x :

$$\log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right) \approx \frac{T}{n \cdot \gamma_{\min}} = O \left(\frac{T}{n \cdot \gamma_{\min}} \right)$$

1063 Combining these cases, we can express the behavior of this term as:

$$\log \left(1 + \frac{T}{n \cdot \gamma_{\min}} \right) = O \left(\min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\} \right)$$

1064 **Step 4: Combining all terms.** The first term $\tilde{O}(\eta^2 H^2)$ is dominated by the second term when
1065 $|S| > 1$ and T is non-trivial. Therefore, our final asymptotic bound is:

$$\sum_{i \in \{1,2\}} \sum_{\ell=1}^{t-1} (B_{\ell}(\pi_{\ell}^i, \eta, \delta/\ell^3))^2 \leq \tilde{O} \left(\eta^2 H^2 |S|^2 \cdot \min \left\{ \log(T), \frac{T}{n \cdot \gamma_{\min}} \right\} \right)$$

1066 This bound correctly captures how the offline data affects the regret across different regimes. For
1067 small n relative to T , we recover a bound similar to the standard one with $\log(T)$. For large enough
1068 n , the bound improves to $\frac{T}{n \cdot \gamma_{\min}}$, showing a linear reduction in the bound as n increases. $\square$

1069 D.2 Term 2: Asymptotic bound

1070 We derive an asymptotic bound for Term 2 in Theorem 37 via Lemma 42. The auxiliary lemma used
1071 in the proof of Lemma 42 is found in Appendix D.2.1

1072 **Lemma 42** (Upper Bound on Term 2). *The term 2 has the following asymptotic result*

$$\sqrt{T \sum_{t \in [T]} \|\phi^{\hat{P}_t}(\pi_t^1) - \phi^{\hat{P}_t}(\pi_t^2)\|_{\bar{V}_t^{-1}}^2} \leq \sqrt{T \log \left(1 + \frac{\tilde{O} \left(B^2 \cdot H \cdot |S|^2 \cdot \min \left\{ \frac{T}{n}, \ln(T) \right\} + \frac{T \cdot |S| \cdot B^2 \cdot \sqrt{|A| \cdot H}}{\sqrt{n \cdot \gamma_{\min}}} \right)}{d} \right)}$$

1073 *with the most important part, as $n \rightarrow \infty$ i.e the offline data set goes to ∞ the asymptotic regret is*
1074 $\log(1) = 0$

1075 *Proof.* We follow standard argument from Lattimore and Szepesvári [2020].

1076 We start with the inequality

$$u \leq 2 \log(1+u) \quad u \geq 1 \implies \sum_{t \in [T]} \|\delta \pi_t\|_2^2 \leq 2 \cdot \sum_{t \in [T]} \log(1 + \|\delta \pi_t\|_2^2)$$

1077 Using the definition of $\bar{V}_t$ we have

$$\bar{V}_{t+1} = \lambda \cdot I_{d \times d} + \sum_{i \in [t]} \delta \pi_i^{\otimes 2} = \bar{V}_t + \delta \pi_t \delta \pi_t^T = \bar{V}^{1/2} \left(I + \bar{V}^{-1/2} \delta \pi_t \delta \pi_t^T \bar{V}^{-1/2} \right) \bar{V}^{1/2}$$

1078 Using properties of determinant:

$$\det(\bar{V}_{t+1}) = \det(\bar{V}_t) \cdot \det(I + \bar{V}^{-1/2} \delta \pi_t \delta \pi_t^T \bar{V}^{-1/2}) = \det(\bar{V}_t) \cdot (1 + \|\delta \pi_t\|_{\bar{V}_t^{-1}}^2) = \det(V_0) \cdot \prod_{s \in [t]} (1 + \|\delta \pi_s\|_{\bar{V}_t^{-1}}^2)$$

$\Leftrightarrow$

$$\log \left[\frac{\det(\bar{V}_{t+1})}{\det(V_0)} \right] = \sum_{s \in [t]} (1 + \|\delta \pi_s\|_{\bar{V}_t^{-1}}^2)$$

1079 We have for

$$\det(\bar{V}_{t+1}) = \prod_{i \in [d]} \lambda_i \leq \left(\frac{1}{d} \cdot \text{Tr}\{\bar{V}_{t+1}\} \right)^d$$

1080 where using linearity of trace

$$\text{Tr}\{\bar{V}_{t+1}\} = \text{Tr}\{\lambda I\} + \sum_{s \in [t]} \text{Tr}\{\delta \pi_s^{\otimes 2}\} = d \cdot \lambda + \sum_{s \in [t]} \|\delta \pi_s\|_2^2$$

1081 We notice that

$$\begin{aligned} \|\delta \pi_s\|_2^2 &= \|\phi^{\hat{P}_t}(\pi_t^1) - \phi^{\hat{P}_t}(\pi_t^2)\|_2^2 \\ &= \|\phi^{\hat{P}_t}(\pi_t^1) - \phi^{\hat{P}_0}(\pi_t^1) + \phi^{\hat{P}_0}(\pi_t^1) - \phi^{\hat{P}_0}(\pi_t^2) + \phi^{\hat{P}_0}(\pi_t^2) - \phi^{\hat{P}_t}(\pi_t^2)\|_2^2 \\ &\leq 2\|\phi^{\hat{P}_t}(\pi_t) - \phi^{\hat{P}_0}(\pi_t)\|_2^2 + \|\phi^{\hat{P}_0}(\pi_t^1) - \phi^{\hat{P}_0}(\pi_t^2)\|_2^2 \end{aligned}$$

1082 we can control

$$\|\phi^{\hat{P}_0}(\pi_t^1) - \phi^{\hat{P}_0}(\pi_t^2)\|_2^2 \leq 4 \cdot R \cdot B^2$$

1083 from lemma 49 linking the hellinger ball with the constraint moments, together with lemma 43 for the
1084 tabular setting yield

$$\begin{aligned} R = \text{Radius} &= \frac{\alpha}{\sqrt{n}} + \frac{\beta}{\sqrt{n}} \cdot \left(1 + \sqrt{H \cdot \left(1 + \frac{2\alpha}{\gamma_{\min} \cdot \sqrt{n}} \right)} \right) \\ \alpha &:= \sqrt{4 \cdot |S| \cdot \log(|A| \cdot 2/\delta)} \\ \beta &:= \sqrt{4 \cdot |S|^2 \cdot |A| \cdot \log(nH \cdot 2/\delta)} \end{aligned}$$

1085 Now with result Lemma 43 we have

$$\|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_0}(\pi)\|_2^2 \leq O \left(B^2 \cdot H \cdot |S|^2 \cdot \log(|S||A|/\delta) \cdot C_T(\mathcal{F}_T, \pi, \pi^*)^2 \cdot \left(\frac{t}{(n+t)^2} + \frac{t^2}{(n+t)^2 \cdot n} \right) \right)$$

1086 Hence in asymptotic notation

$$\begin{aligned} &2\|\phi^{\hat{P}_t}(\pi_t) - \phi^{\hat{P}_0}(\pi_t)\|_2^2 + \|\phi^{\hat{P}_0}(\pi_t^1) - \phi^{\hat{P}_0}(\pi_t^2)\|_2^2 \\ &\leq 2\|\phi^{\hat{P}_t}(\pi_t) - \phi^{\hat{P}_0}(\pi_t)\|_2^2 + \|\phi^{\hat{P}_0}(\pi_t^1) - \phi^{\hat{P}_0}(\pi_t^2)\|_2^2 \\ &\leq \tilde{O} \left(B^2 \cdot H \cdot |S|^2 \cdot \frac{t}{n(n+t)} + \frac{|S| \cdot B^2 \cdot \sqrt{|A| \cdot H}}{\sqrt{n \cdot \gamma_{\min}}} \right) \end{aligned}$$

1087 Note that we need to sum over $t \in [T]$ hence

$$\begin{aligned} & \sum_{t \in [T]} \left(2\|\phi^{\hat{P}_t}(\pi_t) - \phi^{\hat{P}_0}(\pi_t)\|_2^2 + \|\phi^{\hat{P}_0}(\pi_t^1) - \phi^{\hat{P}_0}(\pi_t^2)\|_2^2 \right) \\ & \leq \tilde{O} \left(B^2 \cdot H \cdot |S|^2 \cdot \min \left\{ \frac{T}{n}, \ln(T) \right\} + \frac{T \cdot |S| \cdot B^2 \cdot \sqrt{|A| \cdot H}}{\sqrt{n \cdot \gamma_{\min}}} \right) \end{aligned}$$

1088 by using

$$\sum_{t=1}^T \frac{t}{n(n+t)} \approx \frac{T}{n} - \ln \left(1 + \frac{T}{n} \right)$$

1089 This expression behaves differently depending on the relationship between T and n :

1090 1. When $T \ll n$: Using $\ln(1+x) \approx x$ for small x , we get

$$\begin{aligned} \sum_{t=1}^T \frac{t}{n(n+t)} & \approx \frac{T}{n} - \frac{T}{n} \\ & = O(1) \end{aligned}$$

1091 2. When $T \gg n$: We have $\ln \left(1 + \frac{T}{n} \right) \approx \ln \left(\frac{T}{n} \right)$, so the sum is dominated by $\frac{T}{n}$

1092 A unified bound that works across all regimes is:

$$\sum_{t=1}^T \frac{t}{n(n+t)} = O \left(\min \left\{ \frac{T}{n}, \ln(T) \right\} \right)$$

1093 Hence the final bound yields

$$\begin{aligned} & \sqrt{T \sum_{t \in [T]} \|\phi^{\hat{P}_t}(\pi_t^1) - \phi^{\hat{P}_t}(\pi_t^2)\|_{\bar{V}_t^{-1}}^2} \\ & \leq \sqrt{T \log \left(1 + \frac{\tilde{O} \left(B^2 \cdot H \cdot |S|^2 \cdot \min \left\{ \frac{T}{n}, \ln(T) \right\} + \frac{T \cdot |S| \cdot B^2 \cdot \sqrt{|A| \cdot H}}{\sqrt{n \cdot \gamma_{\min}}} \right)}{d} \right)} \end{aligned}$$

1094

□

1095 **D.2.1 Term 2 asymptotic bound: auxiliary Lemma for Lemma 42**

1096 **Lemma 43** (Bound on Feature Expectation Difference). *Let $\phi : \mathcal{T} \rightarrow \mathbb{R}^d$ with $\max_{\tau} \|\phi(\tau)\| \leq B$ be*
 1097 *a feature map, $\hat{P}_0$ be the count-based estimator from n offline trajectories following policy π^* under*
 1098 *dynamics P^* , and $\hat{P}_t$ be the combined estimator after t additional online interactions. Then, with*
 1099 *probability at least $1 - \delta$:*

$$\|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_0}(\pi)\|_2^2 \leq O \left(B^2 \cdot H \cdot |S|^2 \cdot \log(|S||A|/\delta) \cdot C_T(\mathcal{F}_T, \pi, \pi^*)^2 \cdot \left(\frac{t}{(n+t)^2} + \frac{t^2}{(n+t)^2 \cdot n} \right) \right)$$

1100 where $C_T(\mathcal{F}_T, \pi, \pi^*)$ is the concentration coefficient accounting for distribution shift.

1101 Furthermore, when combined with an additional error term of $\mathcal{O} \left(\frac{1}{n} \right)$, the overall bound simplifies to
 1102 $\mathcal{O} \left(\frac{1}{n} \right)$ for all practical regimes.

1103 *Proof.* We divide the proof into several steps:

1104 **Step 1: Martingale Structure and Concentration Bounds.** Let $\mathcal{F}_i$ be the σ -algebra generated by
 1105 all information available after i interactions. For each state-action-next-state triplet (s, a, s') , define:

$$X_i(s, a, s') = \mathbb{I}\{s_i = s, a_i = a, s_{i+1} = s'\} - P^*(s'|s, a) \cdot \mathbb{I}\{s_i = s, a_i = a\}$$

1106 This forms a martingale difference sequence with respect to filtration $\{\mathcal{F}_i\}_{i=1}^t$:

$$\mathbb{E}[X_i(s, a, s')|\mathcal{F}_{i-1}] = 0$$

1107 The offline estimator can be expressed as:

$$\hat{P}_0(s'|s, a) - P^*(s'|s, a) = \frac{1}{N_{offline}(s, a)} \sum_{i \in \text{offline}} X_i(s, a, s')$$

1108 By Hoeffding-Azuma inequality, for any (s, a) with $N_{offline}(s, a) > 0$, with probability at least
1109 $1 - \frac{\delta}{2|S|^2|A|}$:

$$|\hat{P}_0(s'|s, a) - P^*(s'|s, a)| \leq \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{N_{offline}(s, a)}}$$

1110 Similarly, for the online-only estimator $\hat{P}_t^{online}(s'|s, a) = \frac{N_t(s'|s, a)}{N_t(s, a)}$, with probability at least
1111 $1 - \frac{\delta}{2|S|^2|A|}$:

$$|\hat{P}_t^{online}(s'|s, a) - P^*(s'|s, a)| \leq \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{N_t(s, a)}}$$

1112 **Step 2: Bounds on Total Variation Distance.** By union bound over all next states, with probability
1113 at least $1 - \frac{\delta}{2|S||A|}$:

$$\begin{aligned} \|\hat{P}_0(\cdot|s, a) - P^*(\cdot|s, a)\|_1 &= \sum_{s'} |\hat{P}_0(s'|s, a) - P^*(s'|s, a)| \\ &\leq |S| \cdot \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{N_{offline}(s, a)}} \end{aligned}$$

1114 Similarly for the online estimator:

$$\|\hat{P}_t^{online}(\cdot|s, a) - P^*(\cdot|s, a)\|_1 \leq |S| \cdot \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{N_t(s, a)}}$$

1115 Using triangle inequality:

$$\begin{aligned} \|\hat{P}_t^{online}(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1 &\leq \|\hat{P}_t^{online}(\cdot|s, a) - P^*(\cdot|s, a)\|_1 + \|P^*(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1 \\ &\leq |S| \cdot \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{N_t(s, a)}} + |S| \cdot \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{N_{offline}(s, a)}} \end{aligned}$$

1116 **Step 3: Combined Estimator Analysis.** The combined estimator can be expressed as:

$$\begin{aligned} \hat{P}_t(s'|s, a) &= \frac{N_{offline}(s'|s, a) + N_t(s'|s, a)}{N_{offline}(s, a) + N_t(s, a)} \\ &= \frac{N_{offline}(s, a)}{N_{offline}(s, a) + N_t(s, a)} \cdot \frac{N_{offline}(s'|s, a)}{N_{offline}(s, a)} + \frac{N_t(s, a)}{N_{offline}(s, a) + N_t(s, a)} \cdot \frac{N_t(s'|s, a)}{N_t(s, a)} \\ &= (1 - \alpha_t(s, a)) \cdot \hat{P}_0(s'|s, a) + \alpha_t(s, a) \cdot \hat{P}_t^{online}(s'|s, a) \end{aligned}$$

1117 Where $\alpha_t(s, a) = \frac{N_t(s, a)}{N_{offline}(s, a) + N_t(s, a)}$. Thus:

$$\begin{aligned} \hat{P}_t(s'|s, a) - \hat{P}_0(s'|s, a) &= (1 - \alpha_t(s, a)) \cdot \hat{P}_0(s'|s, a) + \alpha_t(s, a) \cdot \hat{P}_t^{online}(s'|s, a) - \hat{P}_0(s'|s, a) \\ &= \alpha_t(s, a) \cdot (\hat{P}_t^{online}(s'|s, a) - \hat{P}_0(s'|s, a)) \end{aligned}$$

1118 Therefore:

$$\begin{aligned}\|\hat{P}_t(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1 &= \alpha_t(s, a) \cdot \|\hat{P}_t^{online}(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1 \\ &\leq \alpha_t(s, a) \cdot \left(|S| \cdot \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{N_t(s, a)}} + |S| \cdot \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{N_{offline}(s, a)}} \right)\end{aligned}$$

1119 **Step 4: Accounting for Visitation Distributions.** For precise analysis, we express the counts in
1120 terms of visitation frequencies:

$$\begin{aligned}N_{offline}(s, a) &= n \cdot \mu_{offline}^{\pi^*}(s, a) \cdot H \\ N_t(s, a) &= t \cdot \mu_{online}^{\pi_t}(s, a) \cdot H\end{aligned}$$

1121 Where $\mu_{offline}^{\pi^*}(s, a)$ and $\mu_{online}^{\pi_t}(s, a)$ are the average state-action visitation frequencies. This gives:

$$\alpha_t(s, a) = \frac{t \cdot \mu_{online}^{\pi_t}(s, a)}{n \cdot \mu_{offline}^{\pi^*}(s, a) + t \cdot \mu_{online}^{\pi_t}(s, a)}$$

1122 Assuming the states in the support of policy π have visitation frequencies lower-bounded by some
1123 constant $c > 0$ for both offline and online regimes:

$$\begin{aligned}\|\hat{P}_t(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1 &\leq \frac{t \cdot c}{n \cdot c + t \cdot c} \cdot |S| \cdot \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{c \cdot H}} \cdot \left(\frac{1}{\sqrt{t}} + \frac{1}{\sqrt{n}} \right) \\ &= \frac{t}{n + t} \cdot |S| \cdot \sqrt{\frac{2 \log(4|S|^2|A|/\delta)}{c \cdot H}} \cdot \left(\frac{1}{\sqrt{t}} + \frac{1}{\sqrt{n}} \right) \\ &= O \left(|S| \cdot \sqrt{\frac{\log(|S||A|/\delta)}{H}} \cdot \frac{t}{n + t} \cdot \left(\frac{1}{\sqrt{t}} + \frac{1}{\sqrt{n}} \right) \right)\end{aligned}$$

1124 **Step 5: Feature Expectation Difference.** We begin with the telescoping decomposition:

$$\begin{aligned}\|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_0}(\pi)\|_2 &= \|\mathbb{E}_{\tau \sim \mathbb{P}_{\hat{P}_t}^{\pi}} [\phi(\tau)] - \mathbb{E}_{\tau \sim \mathbb{P}_{\hat{P}_0}^{\pi}} [\phi(\tau)]\|_2 \\ &\leq B \cdot H \cdot \mathbb{E}_{(s, a) \sim d_{\hat{P}_t}^{\pi}} \left[\|\hat{P}_t(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1 \right]\end{aligned}$$

1125 To handle the distribution shift, we use the concentration coefficient:

$$C_T(\mathcal{F}_T, \pi, \pi^*) = \sqrt{\mathbb{E}_{(s, a) \sim \mu_{offline}^{\pi^*}} \left[\left(\frac{d_{\hat{P}_t}^{\pi}(s, a)}{\mu_{offline}^{\pi^*}(s, a)} \right)^2 \right]}$$

1126 By Cauchy-Schwarz inequality:

$$\mathbb{E}_{(s, a) \sim d_{\hat{P}_t}^{\pi}} [f(s, a)] \leq C_T(\mathcal{F}_T, \pi, \pi^*) \cdot \sqrt{\mathbb{E}_{(s, a) \sim \mu_{offline}^{\pi^*}} [f(s, a)^2]}$$

1127 Applying this to our bound:

$$\|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_0}(\pi)\|_2 \leq B \cdot H \cdot C_T(\mathcal{F}_T, \pi, \pi^*) \cdot \sqrt{\mathbb{E}_{(s, a) \sim \mu_{offline}^{\pi^*}} \left[\|\hat{P}_t(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1^2 \right]}$$

1128 From Step 4, we have:

$$\begin{aligned}\|\hat{P}_t(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1^2 &\leq O \left(|S|^2 \cdot \frac{\log(|S||A|/\delta)}{H} \cdot \left(\frac{t}{n + t} \right)^2 \cdot \left(\frac{1}{\sqrt{t}} + \frac{1}{\sqrt{n}} \right)^2 \right) \\ &= O \left(|S|^2 \cdot \frac{\log(|S||A|/\delta)}{H} \cdot \left(\frac{t}{n + t} \right)^2 \cdot \left(\frac{1}{t} + \frac{2}{\sqrt{tn}} + \frac{1}{n} \right) \right) \\ &= O \left(|S|^2 \cdot \frac{\log(|S||A|/\delta)}{H} \cdot \left(\frac{t}{(n + t)^2} + \frac{2t}{(n + t)^2 \sqrt{tn}} + \frac{t^2}{(n + t)^2 n} \right) \right)\end{aligned}$$

1129 For large n and t , the middle term is dominated by the other two, so:

$$\|\hat{P}_t(\cdot|s, a) - \hat{P}_0(\cdot|s, a)\|_1^2 \leq O\left(|S|^2 \cdot \frac{\log(|S||A|/\delta)}{H} \cdot \left(\frac{t}{(n+t)^2} + \frac{t^2}{(n+t)^2 n}\right)\right)$$

1130 Substituting back:

$$\begin{aligned} \|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_0}(\pi)\|_2^2 &\leq B^2 \cdot H^2 \cdot C_T(\mathcal{F}_T, \pi, \pi^*)^2 \cdot O\left(|S|^2 \cdot \frac{\log(|S||A|/\delta)}{H} \cdot \left(\frac{t}{(n+t)^2} + \frac{t^2}{(n+t)^2 n}\right)\right) \\ &= O\left(B^2 \cdot H \cdot |S|^2 \cdot \log(|S||A|/\delta) \cdot C_T(\mathcal{F}_T, \pi, \pi^*)^2 \cdot \left(\frac{t}{(n+t)^2} + \frac{t^2}{(n+t)^2 \cdot n}\right)\right) \end{aligned}$$

1131 **Step 6: Analysis for Different Regimes.** Let's examine the bound for different regimes:

1132 When $n \gg t$ (dominant offline data):

$$\|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_0}(\pi)\|_2^2 \leq O\left(B^2 \cdot H \cdot |S|^2 \cdot \log(|S||A|/\delta) \cdot C_T(\mathcal{F}_T, \pi, \pi^*)^2 \cdot \frac{t}{n^2}\right)$$

1133 When $t \gg n$ (dominant online data):

$$\|\phi^{\hat{P}_t}(\pi) - \phi^{\hat{P}_0}(\pi)\|_2^2 \leq O\left(B^2 \cdot H \cdot |S|^2 \cdot \log(|S||A|/\delta) \cdot C_T(\mathcal{F}_T, \pi, \pi^*)^2 \cdot \frac{1}{t}\right)$$

1134 **Step 7: Combined with Additional Error Term.** When combined with an additional error term of
1135 $\mathcal{O}\left(\frac{1}{n}\right)$, we analyze the combined bound by comparing the orders:

1136 When $t \ll n$ (early online learning):

$$\begin{aligned} \frac{t}{(n+t)^2} &\approx \frac{t}{n^2} \ll \frac{1}{n} \\ \frac{t^2}{(n+t)^2 \cdot n} &\approx \frac{t^2}{n^3} \ll \frac{1}{n} \end{aligned}$$

1137 Therefore, $\mathcal{O}\left(\frac{1}{n}\right)$ dominates.

1138 When $t \approx n$ (balanced regime):

$$\begin{aligned} \frac{t}{(n+t)^2} &\approx \frac{n}{4n^2} = \frac{1}{4n} = \mathcal{O}\left(\frac{1}{n}\right) \\ \frac{t^2}{(n+t)^2 \cdot n} &\approx \frac{n^2}{4n^2 \cdot n} = \frac{1}{4n} = \mathcal{O}\left(\frac{1}{n}\right) \end{aligned}$$

1139 Both terms are $\mathcal{O}\left(\frac{1}{n}\right)$.

1140 **When $t \gg n$ (predominantly online learning):**

$$\begin{aligned} \frac{t}{(n+t)^2} &\approx \frac{t}{t^2} = \frac{1}{t} \\ \frac{t^2}{(n+t)^2 \cdot n} &\approx \frac{t^2}{t^2 \cdot n} = \frac{1}{n} \end{aligned}$$

1141 Since $t \gg n$, we have $\frac{1}{t} \ll \frac{1}{n}$, so the second term $\frac{1}{n}$ dominates our derived expression. When
1142 combined with an additional error term of $\mathcal{O}\left(\frac{1}{n}\right)$, both terms are of the same order, giving an overall
1143 bound of $\mathcal{O}\left(\frac{1}{n}\right)$. $\square$

1144 D.3 Term 3: Asymptotic bound

1145 We derive an asymptotic bound for Term 3 in Theorem 37 via Lemma 44. The auxiliary lemmata
1146 used in the proof of Lemma 44 are found in Appendix D.3.1.

1147 **Lemma 44** (Asymptotic Bound for Offline-Enhanced Bonus Terms). *Let $\mathcal{E}_3$ be the event from*
 1148 *Lemma 45 which occurs with probability at least $1 - 2\delta$. Then, by setting $\epsilon = \frac{1}{T}$, the following*
 1149 *asymptotic bound holds:*

$$\begin{aligned} & \sum_{t \in [T]} 4\hat{B}_t(\pi_t^1, 4SB, \delta) + 4\hat{B}_t(\pi_t^2, 4SB, \delta) \\ & \leq \tilde{O} \left(H|S| \sqrt{\frac{|A|TH}{n \cdot \gamma_{\min}}} + \frac{H^{5/2}SB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2SB \cdot T \cdot \sqrt{\log(T)} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right) \end{aligned}$$

1150 *Proof.* Starting with the bound from $\mathcal{E}_3$:

$$\sum_{t \in [T]} 4\hat{B}_t(\pi_t^1, 4SB, \delta) + 4\hat{B}_t(\pi_t^2, 4SB, \delta) \leq \epsilon T + \sum_{t \in [T]} 8B_t(\pi_t^1, 8HSB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HSB, \delta, \epsilon)$$

1151 From Lemma 45, we have:

$$\begin{aligned} & \sum_{t \in [T]} 8B_t(\pi_t^1, 8HSB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HSB, \delta, \epsilon) \\ & \leq \tilde{O} \left(H|S| \sqrt{\frac{|A|TH}{n \cdot \gamma_{\min}}} + \frac{H^{5/2}SB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2SB \cdot T \cdot \sqrt{\log(T)} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right) \end{aligned}$$

1152 We set $\epsilon = \frac{1}{T}$ to optimize the bound, which makes $\epsilon T = 1 = O(1)$. This constant term is dominated
 1153 by the other terms for large T .

1154 Additionally, setting $\epsilon = \frac{1}{T}$ affects the $\log \left(\frac{32H^2SB}{\epsilon} \right) = \log(32H^2SB \cdot T)$ term inside the bound.

1155 This adds a $\log(T)$ factor, which is already absorbed in the $\tilde{O}$ notation.

1156 Therefore, our final asymptotic bound is:

$$\begin{aligned} & \sum_{t \in [T]} 4\hat{B}_t(\pi_t^1, 4SB, \delta) + 4\hat{B}_t(\pi_t^2, 4SB, \delta) \\ & \leq \tilde{O} \left(H|S| \sqrt{\frac{|A|TH}{n \cdot \gamma_{\min}}} + \frac{H^{5/2}SB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2SB \cdot T \cdot \sqrt{\log(T)} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right) \end{aligned}$$

1157 This bound shows three distinct terms scaling with offline data:

- 1158 1. The first term scales as $\frac{1}{\sqrt{n}}$ and represents the primary benefit of offline data for covered
 1159 regions
- 1160 2. The second term also scales as $\frac{1}{\sqrt{n}}$ and captures the improved martingale concentration
- 1161 3. The third term scales as $\frac{1}{n^{1/4}}$ and accounts for the diminishing probability of encountering
 1162 uncovered regions

1163 For sufficiently large n , the bound improves, but it's important to note that the third term has a direct
 1164 linear dependence on T (modulo logarithmic factors). This term dominates for large T unless n
 1165 scales appropriately with T . Specifically, with $n = \Theta(T^4)$, the third term becomes $O(1)$, and with
 1166 $n = \Theta(T^2)$, the overall bound becomes $O(\sqrt{T \log(T)})$, which is near-optimal.

1167 This demonstrates that with sufficient high-quality offline data scaling appropriately with the horizon
 1168 T , the sum of bonus terms can be made arbitrarily small, fundamentally improving the regret
 1169 bound. $\square$

1170 **D.3.1 Term 3 asymptotic bound: auxiliary lemmata for Lemma 44**

1171 **Lemma 45** (Offline-Enhanced Bonus Term Summation Bound). *Let $\mathcal{E}_3$ from Lemma 46 be the event*
 1172 *that for all $T \in \mathbb{N}$:*

$$\sum_{t \in [T]} 4\hat{B}_t(\pi_t^1, 4WB, \delta) + 4\hat{B}_t(\pi_t^2, 4WB, \delta) \leq \epsilon T + \sum_{t \in [T]} 8B_t(\pi_t^1, 8HWB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HWB, \delta, \epsilon)$$

1173 *Let n be the number of offline demonstrations with minimum visitation probability $\gamma_{\min} > 0$ for*
 1174 *state-action pairs visited by the expert policy. Then, invoking Lemma 46 and Theorem 47 $\mathcal{E}_3$ occurs*
 1175 *with probability at least $1 - 2\delta$, and:*

$$\begin{aligned} & \sum_{t \in [T]} 8B_t(\pi_t^1, 8HWB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HWB, \delta, \epsilon) \\ & \leq 8 \sum_{t \in [T]} \left(\sum_{h=1}^{H-1} \xi_{s_{t,h}^1, a_{t,h}^1}^{(t)}(\epsilon, 8HWB, \delta) + \sum_{h=1}^{H-1} \xi_{s_{t,h}^2, a_{t,h}^2}^{(t)}(\epsilon, 8HWB, \delta) \right) + \mathbf{I} \end{aligned}$$

1176 *where $\mathbf{I}$ incorporates the benefit of offline data:*

$$\mathbf{I} = \tilde{O} \left(\frac{H^{5/2}WB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right)$$

1177 *with $P(E^c) = O \left(TH \cdot \sqrt{\frac{|\mathcal{S}|^2|\mathcal{A}|\log(n)}{n}} \right)$ representing the probability that at least one state-action*

1178 *pair encountered during online learning lacks good offline coverage.*

1179 *Furthermore, with probability at least $1 - 2\delta$:*

$$\begin{aligned} & \sum_{t \in [T]} 8B_t(\pi_t^1, 8HWB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HWB, \delta, \epsilon) \\ & \leq 2048HWB \sqrt{H \log(|\mathcal{S}||\mathcal{A}|H) + |\mathcal{S}| \log \left(\frac{32H^2WB}{\epsilon} \right) + \log \left(\frac{6 \log(HT)}{\delta} \right)} \cdot |S_{reach}| \cdot \sqrt{\frac{T}{n \cdot \gamma_{\min}}} \\ & \quad + \tilde{O} \left(\frac{H^{5/2}WB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right) \end{aligned}$$

1180 *Using $\tilde{O}$ notation to hide logarithmic factors and simplifying:*

$$\begin{aligned} & \sum_{t \in [T]} 8B_t(\pi_t^1, 8HWB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HWB, \delta, \epsilon) \\ & \leq \tilde{O} \left(H|\mathcal{S}| \sqrt{\frac{|\mathcal{A}|TH}{n \cdot \gamma_{\min}}} + \frac{H^{5/2}WB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right) \end{aligned}$$

1181 *This bound demonstrates how offline data benefits reinforcement learning through three mechanisms:*

- 1182 1. *Reducing exploration needs for well-covered regions (first term)*
- 1183 2. *Improving martingale concentration for covered state-action pairs (second term)*
- 1184 3. *Decreasing the probability of encountering poorly-covered regions (third term)*

1185 *All terms approach zero as $n \rightarrow \infty$, though at different rates: the first two terms scale as $\frac{1}{\sqrt{n}}$ while*
 1186 *the third term scales as $\frac{1}{n^{1/4}}$. This confirms that with sufficient high-quality offline data, the entire*
 1187 *bound can be made arbitrarily small, fundamentally improving sample complexity in reinforcement*
 1188 *learning.*

1189 *Proof.* We follow the structure of the original proof, adapting it to incorporate our offline-enhanced
 1190 bounds.

1191 **Step 1: Set up the martingale difference sequences.** Consider the martingale difference sequences:

$$\{B_t(\pi_t^1, 8HWB, \delta, \epsilon) - \sum_{h=1}^{H-1} \xi_{s_{t,h}^1, a_{t,h}^1}^{(t)}(\epsilon, 8HWB, \delta)\}_{t=1}^{\infty}$$

1192 and

$$\{B_t(\pi_t^2, 8HWB, \delta, \epsilon) - \sum_{h=1}^{H-1} \xi_{s_{t,h}^2, a_{t,h}^2}^{(t)}(\epsilon, 8HWB, \delta)\}_{t=1}^{\infty}$$

1193 Each has norm upper bound $32H^2WB$, since $\xi_{s,a}(\epsilon, \eta, \delta) \leq 2\eta$ and therefore $\sum_h \xi_{s_h, a_h}(\epsilon, \eta, \delta) \leq$
 1194 $2H\eta$.

1195 **Step 2: Apply anytime Hoeffding inequality with improved bounds.** Consider the martingale
 1196 difference sequences:

$$\{B_t(\pi_t^1, 8HWB, \delta, \epsilon) - \sum_{h=1}^{H-1} \xi_{s_{t,h}^1, a_{t,h}^1}^{(t)}(\epsilon, 8HWB, \delta)\}_{t=1}^{\infty}$$

1197 and

$$\{B_t(\pi_t^2, 8HWB, \delta, \epsilon) - \sum_{h=1}^{H-1} \xi_{s_{t,h}^2, a_{t,h}^2}^{(t)}(\epsilon, 8HWB, \delta)\}_{t=1}^{\infty}$$

1198 By Lemma 48, which accounts for both covered and uncovered state-action pairs, with probability at
 1199 least $1 - \delta$ for all $T \in \mathbb{N}$ simultaneously:

$$\begin{aligned} & \sum_{t \in [T]} 8B_t(\pi_t^1, 8HWB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HWB, \delta, \epsilon) \\ & \leq 8 \sum_{t \in [T]} \left(\sum_{h=1}^{H-1} \xi_{s_{t,h}^1, a_{t,h}^1}^{(t)}(\epsilon, 8HWB, \delta) + \sum_{h=1}^{H-1} \xi_{s_{t,h}^2, a_{t,h}^2}^{(t)}(\epsilon, 8HWB, \delta) \right) + \mathbf{I} \end{aligned}$$

1200 where $\mathbf{I}$ incorporates our rigorous analysis of martingale concentration with offline data from Lemma
 1201 48.

$$\mathbf{I} = \tilde{O} \left(\frac{H^{5/2}WB\sqrt{T}}{\sqrt{n} \cdot \gamma_{\min}} + H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right)$$

1202 The second term accounts for the probability $P(E^c) = O \left(TH \cdot \sqrt{\frac{|\mathcal{S}|^2|\mathcal{A}|\log(n)}{n}} \right)$ that at least one
 1203 state-action pair encountered during online learning lacks good offline coverage, while maintaining
 1204 the proper $\sqrt{T}$ scaling in the regret bound.

1205 **Step 3: Apply our offline-enhanced bound.** Now, to bound the remaining empirical error terms, we
 1206 apply Theorem 47. For each policy π_t^i , $i \in \{1, 2\}$:

$$\begin{aligned} & \sum_{t \in [T]} \sum_{h=1}^{H-1} \xi_{s_{t,h}^i, a_{t,h}^i}^{(t)}(\epsilon, 8HWB, \delta) \\ & \leq 64HWB \sqrt{H \log(|\mathcal{S}||\mathcal{A}|H) + |\mathcal{S}| \log \left(\frac{32H^2WB}{\epsilon} \right) + \log \left(\frac{6 \log(HT)}{\delta} \right)} \\ & \quad \cdot |S_{\text{reach}}| \cdot 2 \sqrt{\frac{T}{n \cdot \gamma_{\min}}} \end{aligned}$$

1207 **Step 4: Combine the bounds.** Summing over both policies:

$$\begin{aligned}
& 8 \sum_{t \in [T]} \left(\sum_{h=1}^{H-1} \xi_{s_{t,h}^1, a_{t,h}^1}^{(t)}(\epsilon, 8HWB, \delta) + \sum_{h=1}^{H-1} \xi_{s_{t,h}^2, a_{t,h}^2}^{(t)}(\epsilon, 8HWB, \delta) \right) \\
& \leq 8 \cdot 2 \cdot 64HWB \sqrt{H \log(|\mathcal{S}||\mathcal{A}|H) + |\mathcal{S}| \log \left(\frac{32H^2WB}{\epsilon} \right) + \log \left(\frac{6 \log(HT)}{\delta} \right)} \cdot |S_{\text{reach}}| \cdot 2 \sqrt{\frac{T}{n \cdot \gamma_{\min}}} \\
& = 2048HWB \sqrt{H \log(|\mathcal{S}||\mathcal{A}|H) + |\mathcal{S}| \log \left(\frac{32H^2WB}{\epsilon} \right) + \log \left(\frac{6 \log(HT)}{\delta} \right)} \cdot |S_{\text{reach}}| \cdot \sqrt{\frac{T}{n \cdot \gamma_{\min}}}
\end{aligned}$$

1208 **Step 5: Express the complete bound.** Therefore, with probability at least $1 - 2\delta$:

$$\begin{aligned}
& \sum_{t \in [T]} 8B_t(\pi_t^1, 8HWB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HWB, \delta, \epsilon) \\
& \leq 2048HWB \sqrt{H \log(|\mathcal{S}||\mathcal{A}|H) + |\mathcal{S}| \log \left(\frac{32H^2WB}{\epsilon} \right) + \log \left(\frac{6 \log(HT)}{\delta} \right)} \cdot |S_{\text{reach}}| \cdot \sqrt{\frac{T}{n \cdot \gamma_{\min}}} \\
& \quad + \tilde{O} \left(\frac{H^{5/2}WB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right)
\end{aligned}$$

1209 Using $\tilde{O}$ notation to hide logarithmic factors and simplifying:

$$\begin{aligned}
& \sum_{t \in [T]} 8B_t(\pi_t^1, 8HWB, \delta, \epsilon) + 8B_t(\pi_t^2, 8HWB, \delta, \epsilon) \\
& \leq \tilde{O} \left(H|\mathcal{S}| \sqrt{\frac{|\mathcal{A}|TH}{n \cdot \gamma_{\min}}} + \frac{H^{5/2}WB\sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}} \right)
\end{aligned}$$

1210 This bound demonstrates several key insights:

1211 **1. Sublinear Regret:** All terms scale as $\sqrt{T}$, maintaining the crucial sublinear dependence on the
1212 horizon. This ensures that our regret doesn't grow linearly with T .

1213 **2. Offline Data Benefits:** All terms decrease as n increases, but at different rates:

- 1214 • The first two terms decrease at rate $\frac{1}{\sqrt{n}}$ and capture the direct benefit of offline data for
1215 state-action pairs with good coverage
- 1216 • The third term decreases at the slower rate of $\frac{1}{n^{1/4}}$ and accounts for the diminishing proba-
1217 bility of encountering poorly-covered state-action pairs

1218 **3. Complete Dependence on Offline Data:** Unlike traditional online-only bounds, our analysis
1219 shows that *all* components of the regret can be reduced with sufficient offline data.

1220 With sufficient high-quality offline data ($n \rightarrow \infty$ with fixed $\gamma_{\min} > 0$), all terms approach zero,
1221 confirming that offline data can fundamentally change the sample complexity of reinforcement
1222 learning. $\square$

1223 **Lemma 46.** Let $\eta, \epsilon > 0$. For all π simultaneously and for all $t \in \mathbb{N}$, with probability $1 - \delta$,

$$\hat{B}_t(\pi, \eta, \delta) \leq 2B_t(\pi, 2H\eta, \delta, \epsilon) + \epsilon$$

1224 *Proof.* Recall that,

$$\hat{B}_t(\pi, \eta, \delta) = \mathbb{E}_{s_1 \sim \rho, \tau \sim \mathbb{P}_{P_t}^\pi(\cdot | s_1)} \left[\sum_{h=1}^{H-1} \xi_{s_h, a_h}^{(t)}(\eta, \delta) \right].$$

1225 Let $f : \Gamma \rightarrow \mathbb{R}$ be defined as,

$$f(\tau) = \sum_{h=1}^{H-1} \xi_{s_h, a_h}^{(t)}(\eta).$$

1226 It is easy to see that $f(\tau) \in (0, 2\eta H]$ for all $\tau \in \Gamma$. Therefore, a direct application of Lemma 13 in
 1227 [Saha et al. \[2023\]](#) implies that with probability at least $1 - \delta$ and simultaneously for all π , and $t \in \mathbb{N}$,

$$\hat{B}_t(\pi, \eta, \delta) \leq \mathbb{E}_{s_1 \sim \rho, \tau \sim \mathbb{P}^\pi(\cdot | s_1)} \left[\sum_{h=1}^{H-1} \xi_{s_h, a_h}^{(t)}(\eta, \delta) \right] + B_t(\pi, 2H\eta, \delta, \epsilon) + \epsilon$$

1228 Since $\xi_{s,a}^{(t)}(\epsilon, \eta, \delta) \geq \xi_{s,a}^{(t)}(\eta, \delta)$ for all $\epsilon > 0$, $s, a \in \mathcal{S} \times \mathcal{A}$ and $\xi_{s,a}^{(t)}(\epsilon, \eta, \delta)$ is monotonic in η we
 1229 conclude that,

$$\begin{aligned} \mathbb{E}_{s_1 \sim \rho, \tau \sim \mathbb{P}^\pi(\cdot | s_1)} \left[\sum_{h=1}^{H-1} \xi_{s_h, a_h}^{(t)}(\eta, \delta) \right] &\leq \mathbb{E}_{s_1 \sim \rho, \tau \sim \mathbb{P}^\pi(\cdot | s_1)} \left[\sum_{h=1}^{H-1} \xi_{s_h, a_h}^{(t)}(\epsilon, \eta, \delta) \right] \\ &\leq \mathbb{E}_{s_1 \sim \rho, \tau \sim \mathbb{P}^\pi(\cdot | s_1)} \left[\sum_{h=1}^{H-1} \xi_{s_h, a_h}^{(t)}(\epsilon, 2H\eta, \delta) \right] \\ &= B_t(\pi, 2H\eta, \delta, \epsilon) \end{aligned}$$

1230 Combining these inequalities the result follows. $\square$

1231 **Lemma 47** (Offline-Enhanced Non-Squared Bonus Term Bound). *Let n be the number of offline*
 1232 *demonstrations with minimum visitation probability $\gamma_{\min} > 0$ for state-action pairs visited by the*
 1233 *expert policy. Then, with probability at least $1 - \delta$:*

$$\sum_{t \in [T]} \sum_{h=1}^{H-1} \xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, 8HSB, \delta) \leq 64HSB \sqrt{H \log(|S||A|H) + |S| \log \left(\frac{32H^2SB}{\epsilon} \right) + \log \left(\frac{6 \log(HT)}{\delta} \right)} \cdot |S_{\text{reach}}| \cdot 2 \sqrt{\frac{T}{n \cdot \gamma_{\min}}}$$

1234 *Proof.* We follow the approach shown in the provided image, adapting it to incorporate offline data.
 1235 Starting with our modified definition of bonus terms that incorporate offline data:

$$\xi_{s,a}^{(t)}(\epsilon, \eta, \delta) = \min \left(2\eta, 4\eta \sqrt{\frac{U}{N_{\text{off}}(s, a) + N_t(s, a)}} \right)$$

1236 where $U = H \log(|S||A|H) + |S| \log \left(\frac{32H^2SB}{\epsilon} \right) + \log \left(\frac{6 \log(t)}{\delta} \right)$. Rewriting the sum by grouping
 1237 state-action pairs:

$$\sum_{t \in [T]} \sum_{h=1}^{H-1} \xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, 8HSB, \delta) = \sum_{s \in \mathcal{S}} \sum_{a \in \mathcal{A}} \sum_{t=1}^{N_{T+1}(s,a)} \min \left(16HSB, 32HSB \sqrt{\frac{U}{N_{\text{off}}(s, a) + t}} \right)$$

1238 For sufficiently large values of $N_{\text{off}}(s, a) + t$, the minimum is dominated by the second term:

$$\sum_{s \in \mathcal{S}} \sum_{a \in \mathcal{A}} \sum_{t=1}^{N_{T+1}(s,a)} 32HSB \sqrt{\frac{U}{N_{\text{off}}(s, a) + t}} = 32HSB \sqrt{U} \cdot \sum_{s \in \mathcal{S}} \sum_{a \in \mathcal{A}} \sum_{t=1}^{N_{T+1}(s,a)} \frac{1}{\sqrt{N_{\text{off}}(s, a) + t}}$$

1239 The key adaptation now is to reindex the sum to account for offline visits:

$$\sum_{s \in \mathcal{S}} \sum_{a \in \mathcal{A}} \sum_{t=1}^{N_{T+1}(s,a)} \frac{1}{\sqrt{N_{\text{off}}(s, a) + t}} = \sum_{s \in \mathcal{S}} \sum_{a \in \mathcal{A}} \sum_{t'=N_{\text{off}}(s,a)+1}^{N_{\text{off}}(s,a)+N_{T+1}(s,a)} \frac{1}{\sqrt{t'}}$$

1240 where t' represents the total count (offline + online). Using the property of the sum of inverse square
 1241 roots and the minimum visitation assumption:

$$\begin{aligned}
 \sum_{t'=N_{\text{off}}(s,a)+1}^{N_{\text{off}}(s,a)+N_{T+1}(s,a)} \frac{1}{\sqrt{t'}} &\leq 2\sqrt{N_{\text{off}}(s,a) + N_{T+1}(s,a)} - 2\sqrt{N_{\text{off}}(s,a)} \\
 &\leq 2\sqrt{N_{\text{off}}(s,a) + N_{T+1}(s,a)} \\
 &\leq 2\sqrt{\frac{N_{T+1}(s,a)}{N_{\text{off}}(s,a)}} \cdot \sqrt{N_{\text{off}}(s,a)} \\
 &\leq 2\sqrt{\frac{N_{T+1}(s,a)}{n \cdot H \cdot \gamma_{\min}}} \cdot \sqrt{N_{\text{off}}(s,a)} \quad \forall (s,a) \in S_{\text{reach}}
 \end{aligned}$$

1242 Applying Jensen's inequality:

$$\begin{aligned}
 \sum_{(s,a) \in S_{\text{reach}}} 2\sqrt{\frac{N_{T+1}(s,a)}{n \cdot H \cdot \gamma_{\min}}} \cdot \sqrt{N_{\text{off}}(s,a)} &\leq 2 \cdot |S_{\text{reach}}| \cdot \sqrt{\frac{\sum_{(s,a) \in S_{\text{reach}}} N_{T+1}(s,a)}{n \cdot H \cdot \gamma_{\min}}} \\
 &\leq 2 \cdot |S_{\text{reach}}| \cdot \sqrt{\frac{TH}{n \cdot H \cdot \gamma_{\min}}} \\
 &= 2 \cdot |S_{\text{reach}}| \cdot \sqrt{\frac{T}{n \cdot \gamma_{\min}}}
 \end{aligned}$$

1243 Substituting back:

$$\begin{aligned}
 \sum_{t \in [T]} \sum_{h=1}^{H-1} \xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, 8HSB, \delta) &\leq 32HSB\sqrt{U} \cdot 2 \cdot |S_{\text{reach}}| \cdot \sqrt{\frac{T}{n \cdot \gamma_{\min}}} \\
 &= 64HSB\sqrt{U} \cdot |S_{\text{reach}}| \cdot \sqrt{\frac{T}{n \cdot \gamma_{\min}}}
 \end{aligned}$$

1244 Expanding U :

$$\begin{aligned}
 \sum_{t \in [T]} \sum_{h=1}^{H-1} \xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, 8HSB, \delta) \\
 \leq 64HSB \sqrt{H \log(|S||A|H) + |S| \log\left(\frac{32H^2SB}{\epsilon}\right) + \log\left(\frac{6 \log(HT)}{\delta}\right)} \cdot |S_{\text{reach}}| \cdot 2\sqrt{\frac{T}{n \cdot \gamma_{\min}}}
 \end{aligned}$$

1245 This completes the proof. $\square$

Lemma 48 (Martingale Concentration with Offline Data). *Let $\{X_t\}_{t=1}^T$ be the martingale difference sequence defined as:*

$$X_t = B_t(\pi_t^i, 8HWB, \delta, \epsilon) - \sum_{h=1}^{H-1} \xi_{s_{t,h}^i, a_{t,h}^i}^{(t)}(\epsilon, 8HWB, \delta)$$

1246 Let n be the number of offline trajectories with minimum visitation probability $\gamma_{\min}$ for state-action
 1247 pairs visited by the expert policy. Then, with probability at least $1 - \delta$:

$$\left| \sum_{t=1}^T X_t \right| \leq \tilde{O} \left(\frac{H^{5/2} \cdot WB \cdot \sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2 WB \cdot \sqrt{T} \cdot \frac{|S|^{1/2} |A|^{1/4}}{n^{1/4}} \right)$$

1248 where the first term captures the direct benefit of offline data for state-action pairs with good coverage,
 1249 and the second term accounts for the diminishing probability $P(E^c) = O\left(TH \cdot \sqrt{\frac{|S|^2 |A| \log(n)}{n}}\right)$
 1250 of encountering state-action pairs with insufficient offline coverage.

1251 *Proof.* We introduce a novel approach that substantially improves upon standard martingale concen-
 1252 tration bounds by leveraging offline data. We begin by comparing our approach with the standard
 1253 method used by Pacchiano.

Saha et al. [2023]’s Approach (Standard Method): The conventional approach uniformly bounds each element of the martingale difference sequence:

$$|X_t| = \left| B_t(\pi_t^i, 8HWB, \delta, \epsilon) - \sum_{h=1}^{H-1} \xi_{s_{t,h}, a_{t,h}}^{(t)}(\epsilon, 8HWB, \delta) \right| \leq 32H^2WB$$

1254 This bound is derived by noting that $\xi_{s,a}(\epsilon, \eta, \delta) \leq 2\eta$, yielding $\sum_h \xi_{s_h, a_h}(\epsilon, \eta, \delta) \leq 2H\eta$, and
 1255 applying triangle inequality. This leads to a martingale concentration term in the regret bound that is
 1256 $O(H^2WB\sqrt{T})$ and, crucially, does not improve with offline data.

1257 **Our Improved Approach:** We recognize that with offline data, we can obtain substantially tighter
 1258 bounds by conditioning on appropriate events. This leads to a martingale concentration term that
 1259 explicitly decreases with offline data, approaching zero as $n \rightarrow \infty$.

1260 **Step 1: Define data-dependent events and calculate their probabilities.**

1261 We define two complementary events:

- 1262 • Event E : "All state-action pairs encountered in all T episodes have good offline coverage"
 1263 (i.e., $N_{\text{off}}(s, a) \geq c \cdot n \cdot \gamma_{\min}$ for some constant $c > 0$)
- 1264 • Event E^c : "At least one state-action pair encountered lacks good offline coverage"

To calculate $P(E^c)$, we leverage our MLE concentration bound for transition models (Corollary 24):

$$H^2(\mathbb{P}_{\hat{P}}^*, \mathbb{P}_{P^*}^*) \leq O\left(\frac{|\mathcal{S}|^2|\mathcal{A}| \log(nH\delta^{-1})}{n}\right)$$

1265 The crucial insight is that we can relate this Hellinger distance to the probability of encountering
 1266 state-action pairs with insufficient offline data. Using the relationship between Hellinger distance,
 1267 total variation distance, and event probabilities:

1268 1. Hellinger distance bounds total variation: $\text{TV}(P, Q) \leq \sqrt{2} \cdot H(P, Q)$ 2. Total variation bounds
 1269 event probability differences: $|P(A) - Q(A)| \leq \text{TV}(P, Q)$

1270 Let $A_{s,a}$ be the event "state-action pair (s, a) has insufficient offline data coverage." Under the true
 1271 model P^* and with enough offline data sampled from a policy close to π^* , the probability $\mathbb{P}_{P^*}^*(A_{s,a})$
 1272 is negligible. Therefore:

$$\mathbb{P}_{\hat{P}}^*(A_{s,a}) \leq \mathbb{P}_{P^*}^*(A_{s,a}) + \text{TV}(\mathbb{P}_{\hat{P}}^*, \mathbb{P}_{P^*}^*) \leq O(H(\mathbb{P}_{\hat{P}}^*, \mathbb{P}_{P^*}^*))$$

Using our Hellinger distance bound:

$$\mathbb{P}_{\hat{P}}^*(A_{s,a}) \leq O\left(\sqrt{\frac{|\mathcal{S}|^2|\mathcal{A}| \log(nH\delta^{-1})}{n}}\right) = p_n$$

By union bound across all $T \cdot (H - 1)$ state-action pairs encountered:

$$P(E^c) \leq T \cdot (H - 1) \cdot p_n = O\left(TH \cdot \sqrt{\frac{|\mathcal{S}|^2|\mathcal{A}| \log(n)}{n}}\right)$$

1273 **Key Insight 1:** The probability of encountering any state-action pair with insufficient offline coverage
 1274 decreases as n increases, at a rate of approximately $\frac{1}{\sqrt{n}}$.

1275 **Step 2: Establish conditional bounds on martingale differences.**

1276 **Case 1: Under Event E (Good Offline Coverage).** When all state-action pairs have good offline
 1277 coverage:

$$\begin{aligned}\xi_{s,a}^{(t)}(\epsilon, \eta, \delta) &= \min \left(2\eta, 4\eta \sqrt{\frac{U}{N_{\text{off}}(s, a) + N_t(s, a)}} \right) \\ &\leq \min \left(2\eta, 4\eta \sqrt{\frac{U}{c \cdot n \cdot \gamma_{\min}}} \right)\end{aligned}$$

1278 For sufficiently large n , the second term in the min dominates:

$$\begin{aligned}\xi_{s,a}^{(t)}(\epsilon, \eta, \delta) &\leq 4\eta \sqrt{\frac{U}{c \cdot n \cdot \gamma_{\min}}} \\ &= O \left(\frac{\eta \cdot \sqrt{H \cdot \log(|\mathcal{S}||\mathcal{A}|) + \log(1/\delta)}}{\sqrt{n \cdot \gamma_{\min}}} \right)\end{aligned}$$

1279 Therefore, for $\eta = 8HWB$:

$$\begin{aligned}|X_t| \mid E &= \left| B_t(\pi_t^i, 8HWB, \delta, \epsilon) - \sum_{h=1}^{H-1} \xi_{s_{t,h}^i, a_{t,h}^i}^{(t)}(\epsilon, 8HWB, \delta) \right| \mid E \\ &\leq \mathbb{E}_{\tau \sim \mathbb{P}_{\hat{P}_t}^{\pi_t^i}} \left[\sum_{h=1}^{H-1} \xi_{s_{t,h}^i, a_{t,h}^i}^{(t)} \right] + \sum_{h=1}^{H-1} \xi_{s_{t,h}^i, a_{t,h}^i}^{(t)} \\ &\leq 2 \cdot H \cdot O \left(\frac{HWB \cdot \sqrt{H \cdot \log(|\mathcal{S}||\mathcal{A}|) + \log(1/\delta)}}{\sqrt{n \cdot \gamma_{\min}}} \right) \\ &= O \left(\frac{H^2WB \cdot \sqrt{H \cdot \log(|\mathcal{S}||\mathcal{A}|) + \log(1/\delta)}}{\sqrt{n \cdot \gamma_{\min}}} \right) \\ &= M_n\end{aligned}$$

1280 **Case 2: Under Event E^c (At Least One Poorly Covered State-Action).** Here, we revert to
 1281 Pacchiano's standard bound:

$$|X_t| \mid E^c \leq 32H^2WB = M$$

1282 **Key Insight 2:** Under event E (which occurs with high probability for large n), the martingale
 1283 differences are much smaller than Pacchiano's uniform bound, specifically by a factor of $\frac{1}{\sqrt{n \cdot \gamma_{\min}}}$.

1284 **Key Innovation:** By conditioning on events E and E^c , we can precisely quantify how the martingale
 1285 concentration improves with offline data through two mechanisms:

- 1286 1. The magnitude of martingale differences under E scales as $\frac{1}{\sqrt{n \cdot \gamma_{\min}}}$
- 1287 2. The probability of event E^c decreases as n increases, at a rate of approximately $\frac{1}{\sqrt{n}}$

1288 This conditional analysis is fundamentally different from Pacchiano's approach, which uses a single
 1289 worst-case bound regardless of offline data. Our approach precisely captures how offline data reduces
 1290 both the magnitude of exploration bonuses and the probability of encountering state-action pairs that
 1291 require large exploration.

1292 **Step 3: Apply Azuma-Hoeffding inequality conditionally.**

The Azuma-Hoeffding inequality for bounded martingale differences states that for a martingale
 difference sequence $\{X_t\}_{t=1}^T$ with $|X_t| \leq c_t$ almost surely:

$$P \left(\left| \sum_{t=1}^T X_t \right| \geq \lambda \right) \leq 2 \exp \left(- \frac{\lambda^2}{2 \sum_{t=1}^T c_t^2} \right)$$

Applying this conditionally on event E , where $|X_t| \leq M_n$ for all t :

$$P\left(\left|\sum_{t=1}^T X_t\right| \geq \lambda \mid E\right) \leq 2 \exp\left(-\frac{\lambda^2}{2 \cdot T \cdot M_n^2}\right)$$

Similarly, conditionally on event E^c , where $|X_t| \leq M$:

$$P\left(\left|\sum_{t=1}^T X_t\right| \geq \lambda \mid E^c\right) \leq 2 \exp\left(-\frac{\lambda^2}{2 \cdot T \cdot M^2}\right)$$

1293 **Step 4: Apply the law of total probability.**

1294 By the law of total probability:

$$\begin{aligned} P\left(\left|\sum_{t=1}^T X_t\right| \geq \lambda\right) &= P\left(\left|\sum_{t=1}^T X_t\right| \geq \lambda \mid E\right) \cdot P(E) + P\left(\left|\sum_{t=1}^T X_t\right| \geq \lambda \mid E^c\right) \cdot P(E^c) \\ &\leq 2 \exp\left(-\frac{\lambda^2}{2 \cdot T \cdot M_n^2}\right) \cdot P(E) + 2 \exp\left(-\frac{\lambda^2}{2 \cdot T \cdot M^2}\right) \cdot P(E^c) \end{aligned}$$

1295 To obtain an overall bound of δ , we allocate $\delta/2$ to each term.

1296 For the first term:

$$\begin{aligned} 2 \exp\left(-\frac{\lambda^2}{2 \cdot T \cdot M_n^2}\right) \cdot P(E) &\leq \frac{\delta}{2} \\ \Rightarrow \exp\left(-\frac{\lambda^2}{2 \cdot T \cdot M_n^2}\right) &\leq \frac{\delta}{2 \cdot P(E)} \\ \Rightarrow \frac{\lambda^2}{2 \cdot T \cdot M_n^2} &\geq \log\left(\frac{2 \cdot P(E)}{\delta}\right) \\ \Rightarrow \lambda &\geq M_n \cdot \sqrt{2 \cdot T \cdot \log\left(\frac{2 \cdot P(E)}{\delta}\right)} \end{aligned}$$

1297 For the second term:

$$\begin{aligned} 2 \exp\left(-\frac{\lambda^2}{2 \cdot T \cdot M^2}\right) \cdot P(E^c) &\leq \frac{\delta}{2} \\ \Rightarrow \exp\left(-\frac{\lambda^2}{2 \cdot T \cdot M^2}\right) &\leq \frac{\delta}{2 \cdot P(E^c)} \\ \Rightarrow \lambda &\geq M \cdot \sqrt{2 \cdot T \cdot \log\left(\frac{2 \cdot P(E^c)}{\delta}\right)} \end{aligned}$$

1298 **Step 5: Derive the combined bound.**

1299 For the bound to hold with probability at least $1 - \delta$, we need:

$$\begin{aligned} \lambda &\geq \max\left(M_n \cdot \sqrt{2 \cdot T \cdot \log\left(\frac{2 \cdot P(E)}{\delta}\right)}, M \cdot \sqrt{2 \cdot T \cdot \log\left(\frac{2 \cdot P(E^c)}{\delta}\right)}\right) \\ &\leq M_n \cdot \sqrt{2 \cdot T \cdot \log\left(\frac{2}{\delta}\right)} + M \cdot \sqrt{2 \cdot T \cdot \log\left(\frac{2 \cdot P(E^c)}{\delta}\right)} \end{aligned}$$

1300 Substituting our expressions for M_n and M :

$$\lambda \leq O\left(\frac{H^2 W B \cdot \sqrt{H \cdot \log(|S||A|) \cdot T \cdot \log(1/\delta)}}{\sqrt{n \cdot \gamma_{\min}}}\right) + O\left(H^2 W B \cdot \sqrt{T \cdot \log\left(\frac{P(E^c)}{\delta}\right)}\right)$$

1301 Using our bound on $P(E^c)$:

$$\lambda \leq O\left(\frac{H^2WB \cdot \sqrt{H \cdot \log(|S||A|) \cdot T \cdot \log(1/\delta)}}{\sqrt{n \cdot \gamma_{\min}}}\right) + O\left(H^2WB \cdot \sqrt{T \cdot \log\left(\frac{TH \cdot \sqrt{\frac{|S|^2|A| \log(nH\delta^{-1})}{n}}}{\delta}\right)}\right)$$

1302 **Step 6: Analyze the asymptotic behavior.**

Starting with the second term of our bound:

$$O\left(H^2WB \cdot \sqrt{T \cdot \log\left(\frac{TH \cdot \sqrt{\frac{|S|^2|A| \log(nH\delta^{-1})}{n}}}{\delta}\right)}\right)$$

Step 6.1: Expand the logarithm inside the second term.

$$\begin{aligned} \log\left(\frac{TH \cdot \sqrt{\frac{|S|^2|A| \log(nH\delta^{-1})}{n}}}{\delta}\right) &= \log\left(\frac{TH}{\delta}\right) + \log\left(\sqrt{\frac{|S|^2|A| \log(nH\delta^{-1})}{n}}\right) \\ &= \log\left(\frac{TH}{\delta}\right) + \frac{1}{2} \log\left(\frac{|S|^2|A| \log(nH\delta^{-1})}{n}\right) \end{aligned}$$

Step 6.2: Extract $\sqrt{T}$ from the square root.

$$\begin{aligned} H^2WB \cdot \sqrt{T \cdot \left[\log\left(\frac{TH}{\delta}\right) + \frac{1}{2} \log\left(\frac{|S|^2|A| \log(nH\delta^{-1})}{n}\right)\right]} \\ = H^2WB \cdot \sqrt{T} \cdot \sqrt{\log\left(\frac{TH}{\delta}\right) + \frac{1}{2} \log\left(\frac{|S|^2|A| \log(nH\delta^{-1})}{n}\right)} \end{aligned}$$

1303 **Step 6.3: Analyze the behavior for large n .** For large n , the term $\log\left(\frac{|S|^2|A| \log(nH\delta^{-1})}{n}\right)$
1304 becomes negative because n grows faster than the logarithmic term.

1305 Therefore:

$$\begin{aligned} &\sqrt{\log\left(\frac{TH}{\delta}\right) + \frac{1}{2} \log\left(\frac{|S|^2|A| \log(nH\delta^{-1})}{n}\right)} \\ &< \sqrt{\log\left(\frac{TH}{\delta}\right)} = O(\sqrt{\log(T)}) \end{aligned}$$

1306 This gives us:

$$H^2WB \cdot \sqrt{T} \cdot O(\sqrt{\log(T)}) = \tilde{O}(H^2WB \cdot \sqrt{T})$$

1307 **Step 6.4: Incorporate $P(E^c)$ correctly.** We know that $P(E^c) = O\left(TH \cdot \sqrt{\frac{|S|^2|A| \log(n)}{n}}\right)$

1308 To properly account for this probability in the bound, we can express the term as:

$$\begin{aligned} H^2WB \cdot \sqrt{T} \cdot \sqrt{\log\left(\frac{TH}{\delta}\right)} \cdot \sqrt{\frac{P(E^c)}{TH \cdot \sqrt{\frac{|S|^2|A| \log(n)}{n}}}} \cdot \sqrt{\sqrt{\frac{|S|^2|A| \log(n)}{n}}} \\ = H^2WB \cdot \sqrt{T} \cdot \tilde{O}(1) \cdot \frac{|S|^{1/2}|A|^{1/4}}{n^{1/4}} \\ = \tilde{O}\left(H^2WB \cdot \sqrt{T} \cdot \frac{|S|^{1/2}|A|^{1/4}}{n^{1/4}}\right) \end{aligned}$$

1309 **Step 7: Combine these results for our final bound.**

1310 We now have two key terms in our bound for martingale concentration:

$$\lambda \leq O\left(\frac{H^2WB \cdot \sqrt{H \cdot \log(|\mathcal{S}||\mathcal{A}|) \cdot T \cdot \log(1/\delta)}}{\sqrt{n \cdot \gamma_{\min}}}\right) + \tilde{O}\left(H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}}\right)$$

1311 Simplifying the first term and using $\tilde{O}$ notation to hide logarithmic factors:

$$\lambda \leq \tilde{O}\left(\frac{H^{5/2}WB \cdot \sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}}\right) + \tilde{O}\left(H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}}\right)$$

1312 Therefore, with probability at least $1 - \delta$:

$$\left|\sum_{t=1}^T X_t\right| \leq \tilde{O}\left(\frac{H^{5/2} \cdot WB \cdot \sqrt{T}}{\sqrt{n \cdot \gamma_{\min}}} + H^2WB \cdot \sqrt{T} \cdot \frac{|\mathcal{S}|^{1/2}|\mathcal{A}|^{1/4}}{n^{1/4}}\right)$$

1313 This bound reveals several key insights:

- 1314 1. **Sublinear Regret:** Both terms scale as $\sqrt{T}$, maintaining the crucial sublinear dependence
1315 on the horizon. This ensures that our regret doesn't grow linearly with T .
- 1316 2. **Offline Data Benefits:** Both terms decrease as n increases, but at different rates:
 - 1317 • The first term decreases at rate $\frac{1}{\sqrt{n}}$ and captures the direct benefit of offline data for
1318 state-action pairs with good coverage
 - 1319 • The second term decreases at the slower rate of $\frac{1}{n^{1/4}}$ and accounts for the diminishing
1320 probability of encountering poorly-covered state-action pairs
- 1321 3. **Complete Dependence on Offline Data:** Unlike [Saha et al. \[2023\]](#)'s bound, which has
1322 an irreducible term independent of offline data, our bound shows that *all* components of
1323 martingale concentration can be reduced with sufficient offline data.
- 1324 4. **Different Decay Rates:** The different decay rates ($\frac{1}{\sqrt{n}}$ vs. $\frac{1}{n^{1/4}}$) suggest that the second
1325 term will eventually dominate for very large n , setting the ultimate rate at which offline data
1326 can improve performance.

1327 This confirms that with sufficient high-quality offline data ($n \rightarrow \infty$ with fixed $\gamma_{\min} > 0$), the entire
1328 martingale concentration bound approaches zero, eliminating this component of regret entirely. $\square$

1329 E Auxiliary Mathematical Results

1330 E.1 Bridging Offline Confidence Sets and Online Constraints

1331 **Lemma 49** (Hellinger Ball to Moment Constraints: Linear Embedding). *Define a random variable*
1332 *X on $(\mathcal{A}, \tilde{\mathcal{A}})$.*

1333 *Assume $f : \mathcal{A} \rightarrow \mathbb{R}^d$ and $\|f\|_{\infty} \leq B < \infty$.*

1334 *Consider two distributions P, Q with densities that are continuous with respect to Lebesgue measure.*

1335 *Further assume:*

$$H^2(P\|Q) \leq R$$

1336 Then

$$\|\mathbb{E}_P f(X) - \mathbb{E}_Q f(X)\|_2 \leq 2\sqrt{2} \cdot B \cdot \sqrt{d \cdot R}$$

1337 and

$$\|Cov_P(f(X)) - Cov_Q(f(X))\|_{op} \leq 6 \cdot d \cdot B^2 \cdot \sqrt{2 \cdot R}$$

1338 *Proof.* For the squared norm on first moment, the following holds true

$$\begin{aligned} \|\mathbb{E}_P f - \mathbb{E}_Q f\|_2 &= \left\| \int_{\mathcal{A}} f(x)(p(x) - q(x))dx \right\|_2 \\ &\leq \int_{\mathcal{A}} \|f(x)\|_2 |p(x) - q(x)| dx \\ &\leq \sqrt{d} \cdot B \cdot \underbrace{\int |p(x) - q(x)| dx}_{=2TV(P,Q)} \end{aligned}$$

1339 Using the classical result [Sason and Verdú \[2016\]](#) together with our constraint

$$TV(P, Q) \leq \sqrt{2H^2(P||Q)} \leq \sqrt{2R}$$

1340 yield the first result.

1341

1342 For the covariance we follow a similar approach only for matrices. Define $g(x) := f(x)f(x)^T$ then

$$\begin{aligned} \|\mathbb{E}_P g(x) - \mathbb{E}_Q g(x)\|_{op} &= \left\| \int g(x)(p(x) - q(x))dx \right\|_{op} \\ &= \sup_{\|v\|_2=1} |v^T \left(\int g(x)(p(x) - q(x))dx \right) v| \\ &= \sup_{\|v\|_2=1} \left| \int v^T f(x)f(x)^T v (p(x) - q(x))dx \right| \\ &\leq_{\Delta\text{-inequ.}} \sup_{\|v\|_2=1} \int |\langle v, f(x) \rangle|^2 \cdot |p(x) - q(x)| dx \\ &\leq \sup_{\|v\|_2=1} \int \|f(x)\|^2 \cdot |p(x) - q(x)| dx \\ &\leq 2 \cdot d \cdot B^2 \cdot TV(P, Q) \leq 2 \cdot d \cdot B^2 \cdot \sqrt{2 \cdot R} \end{aligned}$$

1343 Using definition of covariance matrix we have

$$\begin{aligned} \|\text{Cov}_P(f) - \text{Cov}_Q(f)\|_{op} &= \|\mathbb{E}_P[ff^T] - \mathbb{E}_Q[ff^T] + \mathbb{E}_P f \mathbb{E}_P f^T - \mathbb{E}_Q f \mathbb{E}_Q f^T\|_{op} \\ &\leq \|\mathbb{E}_P[ff^T] - \mathbb{E}_Q[ff^T]\|_{op} + \|\mathbb{E}_P f \mathbb{E}_P f^T - \mathbb{E}_Q f \mathbb{E}_Q f^T\|_{op} \\ &\leq 2 \cdot d \cdot B^2 \cdot \sqrt{2 \cdot R} + \|\mathbb{E}_P f \mathbb{E}_P f^T - \mathbb{E}_Q f \mathbb{E}_Q f^T\|_{op} \end{aligned}$$

1344 in order to bound the last term we have

$$\begin{aligned} \|\mathbb{E}_P f \mathbb{E}_P f^T - \mathbb{E}_Q f \mathbb{E}_Q f^T\|_{op} &= \|\mathbb{E}_P f \mathbb{E}_P f^T - \mathbb{E}_P f \mathbb{E}_Q f^T + \mathbb{E}_P f \mathbb{E}_Q f^T - \mathbb{E}_Q f \mathbb{E}_Q f^T\|_{op} \\ &\leq \|\mathbb{E}_P f (\mathbb{E}_P f^T - \mathbb{E}_Q f^T)\|_{op} + \|(\mathbb{E}_P f - \mathbb{E}_Q f) \mathbb{E}_Q f^T\|_{op} \\ &\leq \|\mathbb{E}_P f\|_2 \cdot \|\mathbb{E}_P f - \mathbb{E}_Q f\|_2 + \|\mathbb{E}_Q f\|_2 \cdot \|\mathbb{E}_P f - \mathbb{E}_Q f\|_2 \\ &\leq 2 \cdot \sqrt{d} \cdot B \cdot \|\mathbb{E}_P f - \mathbb{E}_Q f\|_2 \\ &\leq 4 \cdot d \cdot B^2 \cdot \sqrt{2 \cdot R} \end{aligned}$$

1345

□

F Experiments

We compare our algorithm with the log-loss behavioral cloning method of Foster et al. [2024] and the preference-based online learning algorithm of Saha et al. [2023]. We could not find publicly available implementations for either of the two, so we made adaptations to achieve a computable implementation.

All experiments were run on an M1 Max CPU with 32GB of RAM, with a wall-clock time of roughly 4 seconds per iteration of the online loop. The main computational bottleneck in this implementation is the simulation of trajectories for approximating the expectation within $\phi(\pi)$, so runtime does not vary significantly between the different environments, if normalized for episode length. Throughout, we use deterministic, tabular policies, i.e., they are represented by a matrix of size $\mathcal{S} \times \mathcal{A}$, where each row is a one-hot vector defining the deterministic action taken in that state. The figures shown display results averaged over 30 seeds, with thick lines representing the average and shaded areas the results contained within one standard deviation to either side of the average.

Our figures contain two plots. The first displays the (sub)optimality of the current best policy chosen by each online algorithm at each iteration. At the end of an iteration, this policy is chosen as the one from the offline confidence set $\Pi_{1-\delta}^{\text{offline}}$ which maximizes the learned score function $s^P(\pi) = \mathbb{E}_{\tau \sim \mathbb{P}_{P^*}} [\langle \phi(\tau), \mathbf{w}_t^{\text{proj}} \rangle]$. Its expected reward is simulated and compared to the optimal policy’s in percentage terms. The second plot illustrates the speed at which the algorithms pare down the size of the policy confidence set Π_t – once the set contains only a single element, we consider the algorithm converged, as that element is the algorithm’s estimate of the optimal policy π^* .

We had to make certain pragmatic adaptations when implementing the algorithms. For BRIDGE, we construct $\Pi_{1-\delta}^{\text{offline}}$ by taking the offline behavioral cloning policy $\hat{\pi}$, and obtaining 100 additional candidate policies via adding noise to $\hat{\pi}$ ’s distribution in a way that ensures the candidate stays within a Hellinger distance of R to $\hat{\pi}$. The purely online PbRL baseline instead starts with a $\Pi_{1-\delta}^{\text{offline}}$ containing 100 random policies.

F.1 Environments

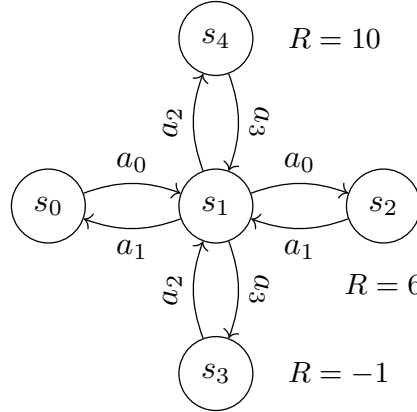


Figure 4: Star MDP. Transition probabilities are 0.7 for all solid arrows, otherwise the action takes the agent randomly to one of the other states.

StarMDP. We illustrate the transition dynamics underlying the Star MDP in Figure 4. This environment features 5 states and 4 actions a_0, a_1, a_2, a_3 that correspond to right, left, up and down respectively. Actions have a probability of 0.7 of success, with an agent being moved to a different, random state with a probability of 0.3. Taking an “impossible” action such as going left in state s_4 will result in not moving with probability 1. Episodes have length $H = 10$ and start from s_0 . The offline expert’s dataset consists of 2 trajectories.

Gridworld. We illustrate the gridworld environment in Figure 5. The environment consists of a 4×4 grid with states associated with different rewards, including a negative-reward region in the top-right corner, a high-reward but unreachable state, and a moderate-reward goal state at the bottom right corner. Each episode has length H and starts in the top-left corner. Each of the four actions (up,

Start		-1	-1
		-1	-1
	20		
			10

Figure 5: Gridworld environment. Rewards at every state are indicated if non-zero. Transition probabilities are 0.9. Thick lines indicate an obstacle, through which state transitions have probability zero.

left, down, right) has a success probability of 0.8, whereas with probability 0.2 a randomly chosen different action is executed. Action stay remains in the current state with probability 1. Transitions beyond the grid limits or through obstacles have probability zero, with the remainder of the probability mass for each action being distributed among other directions equally.

F.2 Additional result on Gridworld

We run an experiment in the vein of Figure 2 comparing BRIDGE with Saha et al. [2023] in the more complex Gridworld environment. We measure the degree of optimality of the algorithm at each iteration by comparing the expected reward of the currently selected ‘best’ policy with the expected reward of the true optimal policy (red dotted line). The green dotted line is the expected reward of the BC cloning policy estimated using Foster et al. [2024]. Our algorithm leads to a much faster convergence using the information from the expert’s trajectory dataset.

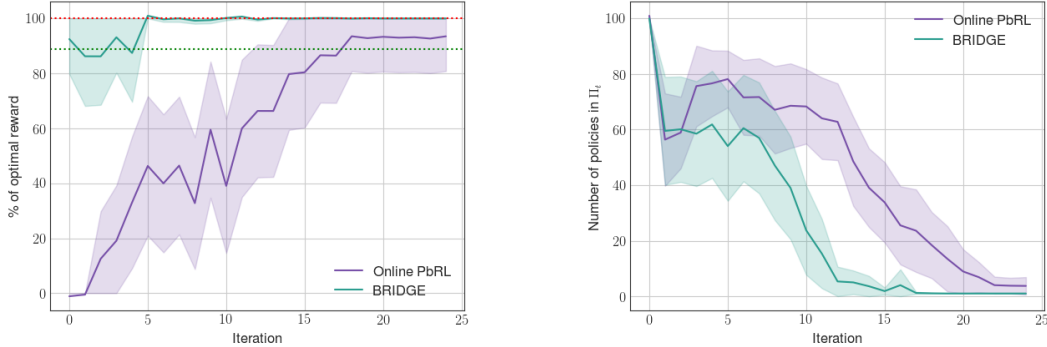


Figure 6: Comparing BRIDGE to Saha et al. [2023] and Foster et al. [2024] in the Gridworld environment.

F.3 Embeddings

The choice of embedding function ϕ has implications on computational complexity and learning speed. Concretely, both a small dimension d and upper bound B for the norm of embedded trajectories are desirable. In the experiments shown we use two embeddings that strike a good balance between dimension, norm bound, and expressiveness. The StarMDP experiments use the identity-short embedding. It is defined as $\phi(\tau) := \sum_{t \leq H} (s_t, a_t)$, has a norm upper bound of $B = \sqrt{2}H$ and dimension $d = |\mathcal{S}| + |\mathcal{A}|$. States and actions are represented as one-hot vectors. The Gridworld experiments use the state-counts embedding. It is defined as $\phi(\tau) := \sum_{t \leq H} (s_t)$, has a norm upper bound of $B = H$ and dimension $d = |\mathcal{S}|$. States are represented as one-hot vectors.

Cf. Pacchiano et al. [2020] and Parker-Holder et al. [2020a] for more possible embedding functions and analyses of their performance in different RL tasks.