

Improved Variance-Aware Confidence Sets for Linear Bandits and Linear Mixture MDP

Zihan Zhang*

Tsinghua University
zihan-zh17@mails.tsinghua.edu.cn

Jiaqi Yang*

Tsinghua University
yangjq17@gmail.com

Xiangyang Ji

Tsinghua University
xyji@tsinghua.edu.cn

Simon S. Du

University of Washington
ssdu@cs.washington.edu

Abstract

This paper presents new *variance-aware* confidence sets for linear bandits and linear mixture Markov Decision Processes (MDPs). With the new confidence sets, we obtain the follow regret bounds:

- For linear bandits, we obtain an $\tilde{O}(\text{poly}(d)\sqrt{1 + \sum_{k=1}^K \sigma_k^2})$ data-dependent regret bound, where d is the feature dimension, K is the number of rounds, and σ_k^2 is the *unknown* variance of the reward at the k -th round. This is the first regret bound that only scales with the variance and the dimension but *no explicit polynomial dependency on K* . When variances are small, this bound can be significantly smaller than the $\tilde{\Theta}(d\sqrt{K})$ worst-case regret bound.
- For linear mixture MDPs, we obtain an $\tilde{O}(\text{poly}(d, \log H)\sqrt{K})$ regret bound, where d is the number of base models, K is the number of episodes, and H is the planning horizon. This is the first regret bound that only scales *logarithmically* with H in the reinforcement learning with linear function approximation setting, thus *exponentially improving* existing results, and resolving an open problem in [Zhou et al., 2020a].

We develop three technical ideas that may be of independent interest: 1) applications of the peeling technique to both the input norm and the variance magnitude, 2) a recursion-based estimator for the variance, and 3) a new convex potential lemma that generalizes the seminal elliptical potential lemma.

1 Introduction

In sequential decision-making problems such as bandits and reinforcement learning (RL), the agent chooses an action based on the current state, with the goal to maximize the total reward. When the state-action space is large, function approximation is often used for generalization. One of the most fundamental and widely used methods is linear function approximation.

For (infinite-actioned) linear bandits, the minimax-optimal regret bound is $\tilde{\Theta}(d\sqrt{K})$ [Dani et al., 2008, Abbasi-Yadkori et al., 2011], where d is the feature dimension and K is the number of total rounds played by the agent.² However, oftentimes the worst-case analysis is overly pessimistic, and

*Equal contribution.

²We follow the reinforcement learning convention to use K to denote the total number of rounds / episodes.

it is possible to obtain data-dependent bound that is substantially smaller than $\tilde{O}(d\sqrt{K})$ in benign scenarios.

One direction to study is the variance magnitude. As a motivating example, in linear bandits, if there is no noise (variance is 0), one only needs to pay at most d regret to identify the best action because d samples are sufficient to recover the underlying linear coefficients (in general position). This constant-type regret bound is much smaller than the $\sqrt{K}$ -type regret bound in the worst case where the variance magnitude is a lower bounded constant. Therefore, a natural question is:

Can we design an algorithm that adapts to the variance magnitude, and its regret degrades gracefully from the benign noiseless constant-type bound to the worst-case $\sqrt{K}$ -type bound?

In RL, exploiting the variance information is also important. For tabular RL, one needs to utilize the variance information, e.g., Bernstein-type exploration bonus to achieve the minimax optimal regret [Azar et al., 2017, Zanette and Brunskill, 2019, Zhang et al., 2020c,a, Menard et al., 2021, Dann et al., 2019]. For example, the recently proposed MVP algorithm [Zhang et al., 2020a], enjoys an $\tilde{O}(\text{polylog}(H) \times (\sqrt{SAK} + S^2A))$ regret bound, where S is the number of states, A is the number of actions, H is the planning horizon, and K is the total number of episodes.³⁴ Notably, this regret bound only scales *logarithmically* with H . On the other hand, without using the variance information, e.g., using Hoeffding-type bonus instead of Bernstein-type bonus, algorithms would suffer a regret that scales *polynomially* with H [Azar et al., 2017].

Going beyond tabular RL, a recent line of work studied RL with linear function approximation with different assumptions [Yang and Wang, 2019, Modi et al., 2020, Jin et al., 2020, Ayoub et al., 2020, Zhou et al., 2020a, Modi et al., 2020]. Our paper studies the linear mixture Markov Decision Process (MDP) setting [Modi et al., 2020, Ayoub et al., 2020, Zhou et al., 2020a], where the transition probability can be represented by a linear function of some features or base models. This model-based assumption is motivated by problems in robotics and queuing systems. We refer readers to Ayoub et al. [2020] for more discussions.

For this linear mixture MDP setting, previous works can obtain regret bounds in the form $\tilde{O}(\text{poly}(d, H)\sqrt{K})$, where d is the number of base models. While these bounds do not scale with SA , they scale *polynomially* with H , because the algorithms in previous works do not use the variance information. In practice, H is often large, and even a polynomial dependency on H may not be acceptable. Therefore, a natural question is

Can we design an algorithm that exploits the variance information to obtain an $\tilde{O}(\text{poly}(d, \log H)\sqrt{K})$ regret bound for linear mixture MDP?

1.1 Our Contributions

In this paper, we develop new, *variance-aware* confidence sets for linear bandits and linear mixture MDP and answer the above two questions affirmatively.

Linear Bandits. For linear bandits, we obtain an $\tilde{O}(\text{poly}(d)\sqrt{1 + \sum_{k=1}^K \sigma_k^2})$ regret bound, where σ_k^2 is the *unknown* variance at the k -th round. To our knowledge, this is the first bound that solely depends on the variance and the feature dimension, and has no explicit polynomial dependency on K . When the variance is very small so that $\sigma_k^2 \ll 1$, this bound is substantially smaller than the worst-case $\tilde{O}(d\sqrt{K})$ bound. Furthermore, this regret bound naturally interpolates between the worst-case $\sqrt{K}$ -type bound and the noiseless-case constant-type bound.

Linear Mixture MDP. For linear mixture MDP, we obtain the desired $\tilde{O}(\text{poly}(d, \log H)\sqrt{K})$ regret bound. This is the first regret bound in RL with function approximation that 1) does not scale with the size of the state-action space, and 2) only scales *logarithmically* with the planning horizon

³ $\tilde{O}(\cdot)$ hides logarithmic factors. Sometimes we write out $\text{polylog } H$ explicitly to emphasize the logarithmic dependency on H .

⁴This bound holds for setting where the transition is homogeneous and the total reward is bounded by 1. We focus on this setting in this paper. See Section 2 and 3 for more discussions.

H . Therefore, we exponentially improve existing results on RL with linear function approximation in term of the H dependency, and resolve an open problem in [Zhou et al., 2020a]. More importantly, our result conveys the positive conceptual message for RL: it is possible to simultaneously overcome the two central challenges in RL, *large state-action space* and *long planning horizon*.

1.2 Main Difficulties and Technical Innovations

We first describe limitations of existing works why they cannot achieve the desired regret bounds described above.

Limitations of Existing Variance-Aware Confidence Sets Faury et al. [2020], Zhou et al. [2020a] applied Bernstein-style inequalities to construct a confidence sets of the least square estimator for linear bandits. However, their methods can not be applied directly to obtain the desired data-dependent regret bound. Abeille et al. [2021] also designed an variance-dependent confidence set for logistic bandits. However in their problem the rewards are Bernoulli and the variance is a function of the mean.

We give a simple example to illustrate their limitations. Consider the case where the variance is always $\sigma^2 \ll 1$. Let $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_{k-1}, y_{k-1})$ be the samples collected before the k -th round. Their confidence set at the k -th round is $\Theta_k = \{\theta \mid \|\theta - \hat{\theta}_k\|_{\Lambda_k^{-1}} \leq C(\sigma\sqrt{d} + 1 + \lambda^{1/2})\}$ (See In Equation (4.3) of Zhou et al. [2020a] and Theorem 1 of Faury et al. [2020]). where $\Lambda_{k-1} = \sum_{\tau=1}^{k-1} \mathbf{x}_\tau \mathbf{x}_\tau^\top + \lambda \mathbf{I}$ is the un-normalized covariance matrix, $\hat{\theta}_k = \Lambda_{k-1}^{-1} \sum_{\tau=1}^{k-1} y_\tau \mathbf{x}_\tau$ is the estimated linear coefficients by least squares, λ is a regularization parameter and C is a constant. Consider the case $d = 1$ and $\mathbf{x}_k = \sqrt{1/K}$ for $k = 1, \dots, K$. Their regret bound is roughly

$$\sum_{k=1}^K (\sigma\sqrt{d} + 1 + \lambda^{1/2}) \|\mathbf{x}_k\|_{\Lambda_k^{-1}} \geq (1 + \lambda^{1/2}) \sum_{i=1}^K \|\mathbf{x}_k\|_{\Lambda_k^{-1}} \geq (1 + \lambda^{1/2}) \sqrt{\frac{K}{1 + \lambda}} \geq \sqrt{K},$$

which is much larger than our bound, $O(\sqrt{K\sigma^2 + 1})$ when σ is very small. For more detailed discussion, please refer to Appendix B.

Below we describe our main techniques.

Elimination with Peeling. Instead of using least squares and upper-confidence-bound (UCB), we use an elimination approach. More precisely, for the underlying linear coefficients $\theta^* \in \mathbb{R}^d$, we build a confidence interval for $(\theta^*)^\top \mu$ for every μ in an ϵ -net of the d -dimensional unit ball, and we eliminate $\theta \in \mathbb{R}^d$ if $\theta^\top \mu$ fails to fall in the confidence interval of $(\theta^*)^\top \mu$ for some μ . To build the confidence intervals, we use 1) an empirical Bernstein inequality (cf. Theorem 4) and 2) the peeling technique to both the input norm and the variance magnitude. As will be clear in the proof (cf. Section D), this peeling step is crucial to obtain a tight regret bound for the example above. The new confidence region provides a tighter estimation for θ^* , which helps address the drawback in least squares.

Generalization of the Elliptical Potential Lemma. Since we use the peeling technique which comes with a clipping operation, we cannot use the seminal elliptic potential lemma Dani et al. [2008] any more. Instead, we propose a more general lemma below, which provides a bound of potential for a general class of convex functions though with a worse dependency on d than the bound in the elliptical potential lemma. We believe this lemma can be applied to other problems as well.

Lemma 1 (Generalized Quadratic Potential Lemma). *Let $f(x) \geq 0$ be a convex function over $\mathbb{R}$ such that $\frac{f(x)}{x^2} \leq \frac{f(y)}{y^2} \leq 1$ and $f(x) \geq f(y)$ if $x^2 \geq y^2 > 0$. Let $\mathbb{B}(1)$ denote the d -dimensional unit ball. Fix $\ell \in (0, 1]$. For any $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_t \in \mathbb{B}(1)$ and $\mu_1, \mu_2, \dots, \mu_t \in \mathbb{B}(1)$, we have that*

$$\sum_{i=1}^t \min \left\{ \frac{f(\mathbf{x}_i \mu_i)}{\sum_{j=1}^{i-1} f(\mathbf{x}_j \mu_j) + \ell^2}, 1 \right\} \leq O(d^4 \log(dt/\ell)).$$

Note that by choosing $f(x) = x^2$ and $\mu_i = \frac{\mathbf{x}_i \Lambda_i^{-1}}{\|\mathbf{x}_i \Lambda_i^{-1}\|}$ with $\Lambda_i = \sum_{j=1}^{i-1} \mathbf{x}_j \mathbf{x}_j^\top + \ell \mathbf{I}$, Lemma 1 reduces to the classical elliptic potential lemma [Dani et al., 2008]. Our proof consists of two major parts.

We first establish a symmetric version of Equation (??) using rearrangement inequality, and then bound the number of times the energy for some μ (i.e., $\sum_{j=1}^i f(x_j \mu) + l^2$) doubles. The full proof is deferred to Appendix C.

For linear mixture MDP, we propose another technique to further reduce the dependency on d .

Recursion-based Variance Estimation. In linear bandits, generally it is not possible to estimate the variance because the variance at each round can arbitrarily different. On the other hand, for linear mixture MDP, the variance is a quadratic function of the underlying coefficient θ^* . Furthermore, the higher moments are polynomial functions of θ^* . Utilizing this rich structure and leveraging the recursion idea in previous analyses on tabular RL [Lattimore and Hutter, 2012, Li et al., 2020, Zhang et al., 2020a], we explicitly estimate the variance and higher moments to further reduce the regret. See Section 5 for more explanations.

2 Related Work

Linear Bandits. There is a line of theoretical analyses of linear bandits problems [Auer et al., 2002, Dani et al., 2008, Chu et al., 2011, Abbasi-Yadkori et al., 2011, Li et al., 2019a,b]. For infinite-actioned linear bandits, the minimax regret bound is $\tilde{\Theta}(d\sqrt{K})$, and recent works tried to give fine-grained instance-dependent bounds [Katz-Samuels et al., 2020, Jedra and Proutiere, 2020]. For multi-armed bandits, Audibert et al. [2006] showed by exploiting the variance information, one can improve the regret bound. For linear bandits, only a few work studied how to use the variance information. Faury et al. [2020] studied logistic bandit problem with adaptivity to the variance of noise, where a Bernstein-style confidence set was proposed. However, they assume the variance is known and cannot attain the desired variance-dependent bound due to the example we gave above. Linear bandits can be also seen as a simplified version of RL with linear function approximation, where the planning horizon degenerates to $H = 1$.

RL with Linear Function Approximation. Recently, it is a central topic in the theoretical RL community to figure out the necessary and sufficient conditions that permit efficient learning in RL with large state-action space [Wen and Van Roy, 2013, Jiang et al., 2017, Yang and Wang, 2019, 2020, Du et al., 2019b, 2020a, 2019a, 2020b, Jiang et al., 2017, Feng et al., 2020, Sun et al., 2019, Dann et al., 2018, Krishnamurthy et al., 2016, Misra et al., 2019, Ayoub et al., 2020, Zanette et al., 2020, Wang et al., 2019, 2020c,b, Jin et al., 2020, Weisz et al., 2020, Modi et al., 2020, Shariff and Szepesvári, 2020, Jin et al., 2020, Cai et al., 2019, He et al., 2020, Zhou et al., 2020a]. However, to our knowledge, all existing regret upper bounds have a polynomial dependency on the planning horizon H , except works that assume the environment is deterministic [Wen and Van Roy, 2013, Du et al., 2020b].

This paper studies the linear mixture MDP setting [Ayoub et al., 2020, Zhou et al., 2020b,a, Modi et al., 2020], which assumes the underlying transition is a linear combination of some known base models. Ayoub et al. [2020] gave an algorithm, UCRL-VTR, with an $\tilde{O}(dH^2\sqrt{K})$ regret in the time-inhomogeneous model.⁵ Our algorithm improves the H -dependency from $\text{poly}(H)$ to $\text{polylog}(H)$, at the cost of a worse dependency on d .

Variance Information in Tabular MDP. The use of the variance information in tabular MDP was first proposed by Lattimore and Hutter [2012] in the discounted MDP setting, and was later adopted in the episodic MDP setting [Azar et al., 2017, Jin et al., 2018, Zanette and Brunskill, 2019, Dann et al., 2019, Zhang et al., 2020a,b]. This technique is crucial to tighten the dependency on H .

Concurrent Work by Zhou et al. [2020a]. While preparing this draft, we noticed a concurrent work by Zhou et al. [2020a], who also studied how to use the variance information for linear bandits

⁵The time-inhomogeneous model refers to the setting where the transition probability can vary at different levels, and the time-homogeneous model refers to the setting where the transition probability is the same at different levels. Roughly speaking, the model complexity of the time-inhomogeneous model is H times larger than that of the time-homogeneous model. In general, it is straightforward to tightly extend a result for the time-homogeneous model to the time-inhomogeneous model by extending the state-action space [Jin et al., 2018, Footnote 2], but not vice versa.

and linear mixture MDPs. We first compare their results with ours. For linear bandits, they proved an $\tilde{O}(\sqrt{dK} + d\sqrt{\sum_{i=1}^K \sigma_i^2})$ regret bound, while we prove an $\tilde{O}(d^{4.5}\sqrt{\sum_{i=1}^K \sigma_i^2} + d^5)$ regret bound. Our bound has a worse dependency on d , but in the regime where K is very large and the sum of the variances is small, our bound is stronger. Furthermore, they assumed *the variance is known while we do not need this assumption*. For linear mixture MDP, they proved an $\tilde{O}(\sqrt{d^2H} + dH^2\sqrt{K} + d^2H^2 + d^3H)$ bound for the time-inhomogeneous model, while we prove an $\tilde{O}(d^{4.5}\sqrt{K} + d^5) \times \text{polylog}(H)$ bound for the time-homogeneous model. Their bound has a better dependency on d than ours and is near-optimal in the regime $K = \Omega(\text{poly}(d, H))$ and $H = O(d)$. On the other hand, we have an exponentially better dependency on H in the time-homogeneous model. Indeed, obtaining a regret bound that is logarithmic in H (in the time-homogeneous model) was raised as an open question in their paper [Zhou et al., 2020a, Remark 5.5].

Next, we compare the algorithms and the analyses. The algorithms in the two papers are very different in nature: ours are based on elimination while theirs are based on least squares and UCB. We note that, for linear bandits, their current analysis cannot give a $\sqrt{K}$ -free bound because there is a term that scales *inversely* with the variance. This can be seen by plugging the first line of their (B.25) to their (B.23). For the same reason, they cannot give a horizon-free bound in the time-homogeneous linear mixture MDP. In sharp contrast, our analysis does not have the term depending on the inverse of the variance. On the other hand, their algorithms are computationally efficient (given certain computation oracles), but our algorithms are not because ours are elimination-based. See Section 6 for more discussions.

3 Preliminaries

Notations. We use $\mathbb{B}_p^d(r) = \{x \in \mathbb{R}^d : \|x\|_p \leq r\}$ to denote the d -dimensional ℓ_p -ball of radius r , so $\mathbb{B}(1) = \mathbb{B}_2^d(1)$. For any set $S \subseteq \mathbb{R}^d$, we use ∂S to denote its boundary. For $N \in \mathbb{N}$, we define $[N] = \{1, \dots, N\}$. One important operation used in our algorithms and analyses is clipping. Given $\ell > 0$ and $u \in \mathbb{R}$, we define

$$\text{clip}(u, \ell) = \min\{|u|, \ell\} \cdot \frac{u}{|u|}$$

for $u \neq 0$ and $\text{clip}(0, \ell) = 0$. For any two vectors $\mathbf{u}, \mathbf{v}$, to save notations, we use $\mathbf{u}\mathbf{v} = \mathbf{u}^\top \mathbf{v}$ to denote their inner product when no ambiguity.

Linear Bandits. We use K to denote the number of rounds in the linear bandits. At each round $k = 1, \dots, K$, the algorithm is first given the context set $\mathcal{A}_k \subseteq \mathbb{B}_2^d(1)$, then the algorithm chooses an action $\mathbf{x}_k \in \mathcal{A}_k$ and receives the noisy reward $r_k = \mathbf{x}_k \boldsymbol{\theta}^* + \varepsilon_k$, where $\boldsymbol{\theta}^* \in \mathbb{B}_2^d(1)$ is the unknown underlying linear coefficients and ε_k is the random noise. We define $\mathcal{F}_k = \sigma(\mathbf{x}_1, \varepsilon_1, \dots, \mathbf{x}_k, \varepsilon_k, \mathbf{x}_{k+1})$. We assume that $|r_k| \leq 1$ and that the noise ε_k satisfies $\mathbb{E}[\varepsilon_k | \mathcal{F}_k] = 0$ and $\mathbb{E}[\varepsilon_k^2 | \mathcal{F}_k] = \sigma_k^2$. The goal is to learn $\boldsymbol{\theta}^*$ and minimize the cumulative expected regret $\mathbb{E}[\mathfrak{R}^K]$, where

$$\mathfrak{R}^K = \sum_{k=1}^K [\max_{\mathbf{x} \in \mathcal{A}_k} \mathbf{x} \boldsymbol{\theta}^* - \mathbf{x}_k \boldsymbol{\theta}^*].$$

Remark 1. Here we assume the reward is uniformly bounded ($|r_k| \leq 1$) instead of 1-sub-Gaussian commonly used in the literature only for the ease of presentation, because in RL, it is standard to assume bounded reward. Note if the noise is 1-sub-Gaussian, our algorithm also applies with only an $O(\log K)$ overhead because a problem with 1-sub-Gaussian noise can be reduced to that with uniformly bounded noise by clipping the noise with a threshold $O(\log K)$.

Episodic MDP and Linear Mixture MDP. We use a tuple $(\mathcal{S}, \mathcal{A}, r, P, K, H)$ to define an episodic finite-horizon MDP. Here, $\mathcal{S}$ is its state space, $\mathcal{A}$ is its action space, $r : \mathcal{S} \times \mathcal{A} \rightarrow [0, 1]$ is its reward function, $P(s' | s, a)$ is the transition probability from the state-action pair (s, a) to the new state s' , K is the number of episodes, and H is the planning horizon of each episode. Without the loss of generality, we assume a fixed initial state s_1 . A sequence of functions $\pi = \{\pi_h : \mathcal{S} \rightarrow \Delta(\mathcal{A})\}_{h=1}^H$ is a policy, where $\Delta(\mathcal{A})$ denotes the set of all possible distributions over $\mathcal{A}$.

Algorithm 1 VOFUL: Variance-Aware Optimism in the Face of Uncertainty for Linear Bandits

- 1: **Initialize:** $\ell_i = 2^{2-i}$, $\iota = 16d \ln \frac{dK}{\delta}$, $L_2 = \lceil \log_2 K \rceil$, $\Lambda_2 = \{1, 2, \dots, L_2 + 1\}$, $\Theta_1 = \mathbb{B}_2^d(1)$,
Let $\mathcal{B}$ be an K^{-3} -net of $\mathbb{B}_2^d(2)$ with size not larger than $(\frac{4}{K})^{3d}$
- 2: **for** $k = 1, 2, \dots, K$ **do**
- 3: **Optimistic Action Selection:**
- 4: Observe context set $\mathcal{A}_k \subseteq \mathbb{B}_2^d(1)$
- 5: Compute $\mathbf{x}_k \leftarrow \arg \max_{\mathbf{x} \in \mathcal{A}_k} \max_{\boldsymbol{\theta} \in \Theta_k} \mathbf{x}\boldsymbol{\theta}$, choose action $\mathbf{x}_k$
- 6: Receive feedback y_k
- 7: **Construct Confidence Set:**
- 8: For each $\boldsymbol{\theta} \in \mathbb{B}_2^d(1)$, define $\epsilon_k(\boldsymbol{\theta}) = y_k - \mathbf{x}_k \boldsymbol{\theta}$, $\eta_k(\boldsymbol{\theta}) = (\epsilon_k(\boldsymbol{\theta}))^2$.
- 9: Define confidence set $\Theta_{k+1} = \bigcap_{j \in \Lambda_2} \Theta_{k+1}^j$, where

$$\Theta_{k+1}^j = \left\{ \boldsymbol{\theta} \in \mathbb{B}_2^d(1) : \left| \sum_{v=1}^k \text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}) \epsilon_v(\boldsymbol{\theta}) \right| \leq \sqrt{\sum_{v=1}^k \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}) \eta_v(\boldsymbol{\theta}) \iota + \ell_j \iota}, \forall \boldsymbol{\mu} \in \mathcal{B} \right\} \quad (1)$$

and $\text{clip}_j(\cdot) = \text{clip}(\cdot, \ell_j)$.

10: **end for**

At each episode $k = 1, \dots, K$, the algorithm outputs a policy π^k , which is then executed on the MDP by $a_h^k \sim \pi_h^k(s_h^k)$, $s_{h+1}^k \sim P(\cdot | s_h^k, a_h^k)$. We let $r_h^k = r(s_h^k, a_h^k)$ be the reward at time step h in episode k . Importantly, we assume the transition model $P(\cdot | \cdot, \cdot)$ is time-homogeneous, which is necessary to bypass the $\text{poly}(H)$ dependency. We assume that the reward function is known, which is standard in the theoretical RL literature to simplify the presentation [Modi et al., 2020, Ayoub et al., 2020]. We let π^* to denote the optimal policy which achieves the maximum reward in expectation.

We make the following regularity assumption on the rewards: the sum of reward, $\sum_{h=1}^H r_h$, in each episode is bounded by 1.

Assumption 2 (Non-uniform reward). $\sum_{h=1}^H r_h^k \leq 1$ almost surely for any policy π^k .

This assumption is much weaker than the common assumption where the reward at each time step is bounded by $1/H$ (uniform reward) because Assumption 2 allows one spiky reward as large as $\Omega(1)$. See more discussions about this reward scaling in Jiang and Agarwal [2018], Wang et al. [2020a], Zhang et al. [2020a].

For any policy π , we define its H -step V -function and Q -function as

$$V_h^\pi(s) = \max_{a \in \mathcal{A}} Q_h^\pi(s, a)$$

where $Q_h^\pi(s, a) = r(s, a) + \mathbb{E}_{s' \sim P(\cdot | s, a)} V_{h+1}^\pi(s')$ for $h = 1, \dots, H$

where we set $V_{H+1} = 0$. For simplicity, we also denote $V^\pi(s_1) = V_1^\pi(s_1)$ and $V^*(s_1) = V^{\pi^*}(s_1)$.

A linear mixture MDP is an episodic MDP with the extra assumption that its transition model is an unknown linear combination of a known set of models. Specifically, there is an unknown parameter $\boldsymbol{\theta}^* \in \mathbb{B}_1^d(1)$, such that $P = \sum_{i=1}^d \theta_i^* P_i$ where based models $P_1, \dots, P_d$ are given. The goal is to learn $\boldsymbol{\theta}^*$ and minimize the cumulative expected regret $\mathbb{E}[\mathfrak{R}^K]$, where

$$\mathfrak{R}^K = \sum_{k=1}^K [V^*(s_1) - V^k(s_1)].$$

4 Algorithm and Theory for Linear Bandits

In this section, we introduce our algorithm for linear bandits and analyze its regret. The pseudo-code is listed in Algorithm 1. The following theorem shows our algorithm achieves the desired variance-dependent regret bound. The full proof is deferred to Section D.

Theorem 3. *The expected regret of Algorithm 1 is bounded by $\mathbb{E}[\mathfrak{R}^K] \leq \tilde{O}(d^{4.5} \sqrt{\sum_{k=1}^K \sigma_k^2} + d^5)$.*

This theorem shows our algorithm's regret has no explicit polynomial dependency on the number of rounds K . In the worst-case where the variance is $\Omega(1)$, our bound becomes $\tilde{O}(d^{4.5}\sqrt{K} + d^5)$, which has a worse dependency on d compared with the minimax optimal algorithms [Dani et al., 2008, Abbasi-Yadkori et al., 2011]. However, in the benign case where the variance is $o(1)$, our bound can be much smaller. In particular, in the noiseless case, our bound is a constant-type regret bound, up to logarithmic factors. One future direction is to design an algorithm that is minimax optimal in the worst-case but also adapts to the variance magnitude like ours.

4.1 Main Algorithm

Now we describe our algorithm. Similar to many other linear bandit algorithms, the algorithm maintains confidence sets $\{\Theta_k\}_{k \geq 1}$ for the underlying parameter θ^* , and then choose the action greedily according to the confidence set.

To relax the known variance assumption, we use the following empirical Bernstein inequality that depends on the *empirical variance*, in contrast to the Bernstein inequality that depends on the *true variance*, which was used in existing works [Zhou et al., 2020b, Faury et al., 2020].

Theorem 4. *Let $\{\mathcal{F}_i\}_{i=0}^n$ be a filtration. Let $\{X_i\}_{i=1}^n$ be a sequence of real-valued random variables such that X_i is $\mathcal{F}_i$ -measurable. We assume that $\mathbb{E}[X_i | \mathcal{F}_{i-1}] = 0$ and that $|X_i| \leq b$ almost surely. For $\delta < e^{-1}$, we have*

$$\Pr \left[\left| \sum_{i=1}^n X_i \right| \leq 8 \sqrt{\sum_{i=1}^n X_i^2 \ln \frac{1}{\delta}} + 16b \ln \frac{1}{\delta} \right] \geq 1 - 6\delta \log_2 n. \quad (2)$$

Importantly, this inequality controls the deviation via the empirical variance, which is X_i^2 and can be computed once X_i is known. Note some previously proved inequalities require certain independence assumptions and thus cannot be directly applied to martingales [Maurer and Pontil, 2009, Peel et al., 2013], so they cannot be used for solving our linear bandits problem. The proof of the theorem is deferred to Appendix D.2.

More effort is devoted to designing a confidence set that fully exploits the variance information. Note Theorem 4 is for real-valued random variables, and it remains unclear how to generalize it to the linear regression setting, which is crucial for building confidence sets for linear bandits. Previous works built up their confidence sets based on analyzing the ordinary ridged least square estimator [Dani et al., 2008, Abbasi-Yadkori et al., 2011], or the weighted one [Zhou et al., 2020a].

We drop the least square estimators and instead, we take a testing-based approach, as done in Equation (1). To illustrate the idea, we first ignore the $\text{clip}_j(\cdot)$ operation and ℓ_j terms. We define the noise function $\epsilon_k(\theta)$ and the variance function $\eta_k(\theta)$ (Line 8 of Algorithm 1). Note that $\epsilon_k(\theta^*) = \varepsilon_k$ and $\eta_k(\theta^*) = \varepsilon_k^2$, so we have the following fact: if $\theta = \theta^*$, then Equation (2) would be true if we replace $X_k = w_k(\mu)\epsilon_k(\theta)$ and $X_k^2 = w_k^2(\mu)\eta_k(\theta)$ with high probability, where $\{w_k(\mu)\}$ is a proper sequence of weights depending on the test direction μ . Our approach uses the fact in the opposite direction: if weighted $w_k(\mu)\epsilon_k(\theta)$, $w_k^2(\mu)\eta_k(\theta)$ satisfies Equation (2) for all possible test directions μ in an K^{-3} -net of the d -dimensional unit ball, then we put θ into the confidence set.

Remark 2. *One can also view the algorithm as an elimination-based algorithm: if there exists some test direction μ such that Equation (2) fails for $X_k = w_k(\mu)\epsilon_k(\theta)$ and $X_k^2 = w_k^2(\mu)\eta_k(\theta)$, then we eliminate θ from the confidence set permanently.*

Given the test direction μ , following the least square estimation, $w_k(\mu)$ is set to be $x_k\mu$. However, with $w_k(\mu) = x_k\mu$, the right-hand-side of Equation (2) is at least $b \geq \max_{1 \leq k \leq n} |w_k(\mu)| = \max_{1 \leq k \leq n} |x_k\mu|$, which might be dominant compared with $\sum_{k=1}^n w_k^2(\mu)\eta_k(\theta)$ (See Appendix B for a toy example). To address this problem, we consider to peel $w_k(\mu)$ for various thresholds of difference level. More precisely, we construct confidence regions respectively with $w_k^j(\mu) = \text{clip}_j(x_k\mu)$, where $l_j = 2^{2-j}$ for $j = 1, 2, \dots, \lceil \log_2 K \rceil$. At last, we define the final confidence region as the intersections of all these confidence regions.

Remark 3. *Note that existing confidence sets in Equation (1) either do not exploit variance information [Dani et al., 2008, Abbasi-Yadkori et al., 2011], or require the variance to be known and do not fully exploit the variance information [Zhou et al., 2020a, Faury et al., 2020] as their regret bounds still have an $\tilde{O}(\sqrt{K})$ term.*

4.2 Proof Sketch of Theorem 3

Now we explain how our confidence set enables us to obtain a variance-dependent regret bound. We define $\theta_k = \arg \max_{\theta \in \Theta_k} \mathbf{x}_k(\theta - \theta^*)$ and $\mu_k = \theta_k - \theta^*$. Then our goal is to bound the regret $\sum_k \mathbf{x}_k \mu_k$. Our main idea is to consider $\{\mathbf{x}_k\}, \{\mu_k\}$ as two sequences of vectors. We decouple the complicated dependency between $\{\mathbf{x}_k\}$ and $\{\mu_k\}$ by a union bound over the net $\mathcal{B}$ (defined in Line 1 of Algorithm 1). To bound the regret, we implicitly divide all rounds $k \in [K]$ into norm layers based on $\log_2 |\mathbf{x}_k \mu_k|$ in the analysis.⁶ Within each layer, we apply Equation (1) to obtain the relations between μ_k and $\{\mathbf{x}_1, \dots, \mathbf{x}_{k-1}\}$, which would self-normalize the growth of the two sequences, ensuring that their in-layer total sum is properly bounded. Since we have logarithmically many layers, the total regret is then properly bounded. We highlight that our norm peeling technique ensures that the variance-dependent term dominates the other variance-independent term in Bernstein inequalities ($\sqrt{\sum X_i^2} \gtrsim b$ in Theorem 4), which resolves the variance-independent term in the final regret bound obtained by Zhou et al. [2020a].

We start the analysis by proving that the probability of failure events (i.e., the events where $\theta^* \notin \Theta_k$ for some $k \in [K]$) is properly bounded (see Lemma 18). Assuming the successful events happen, we have that $\theta^* \in \Theta_k$ for all $k \in [K]$. Then we obtain that

$$\mathfrak{R}^K := \sum_{k=1}^K \left(\max_{\mathbf{x} \in \mathcal{A}_k} \theta^* - \mathbf{x}_k \theta^* \right) \leq \sum_{k=1}^K \max_{\mathbf{x} \in \mathcal{A}_k, \theta \in \Theta_k} \mathbf{x}_k(\theta_k - \theta^*) \leq \sum_{k=1}^K \mathbf{x}_k(\theta_k - \theta^*) = \sum_{k=1}^K \mathbf{x}_k \mu_k.$$

Next we divide the time steps $[K]$ into $L_2 + 1$ disjoint subsets $\{\mathcal{K}_j\}_{j=1}^{L_2+1}$ according to the magnitude of $\mathbf{x}_k \mu_k$. More precisely, for $1 \leq j \leq L_2$ we assign k to $\mathcal{K}_j$ iff $\mathbf{x}_k \mu_k \in (l_j/2, l_j]$, and for $j = L_2 + 1$, we assign k to $\mathcal{K}_j$ iff $\mathbf{x}_k \mu_k \leq l_{L_2+1}/2$. Define

$$\Phi_k^j(\mu) = \sum_{v=1}^{k-1} \text{clip}_j(\mathbf{x}_v \mu) \mathbf{x}_v \mu + \ell_j^2, \quad \Psi_k^j(\mu) = \sum_{v=1}^{k-1} \text{clip}_j^2(\mathbf{x}_v \mu) \eta_v(\theta^*). \quad (3)$$

By the definition of Θ_k in (1), we have that (see Claim 20)

$$\sum_{k=1}^K \mathbf{x}_k \mu_k \leq 1 + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \mu_k \times \frac{3\sqrt{\Psi_k^j(\mu_k)\ell} + \sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \mu_k)(\mathbf{x}_v \mu_k)^2\ell} + 3\ell_j\ell}{\Phi_k^j(\mu_k)}. \quad (4)$$

Continuing the computation, we have that

$$\begin{aligned} & \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \mu_k \frac{\sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \mu_k)(\mathbf{x}_v \mu_k)^2\ell}}{\Phi_k^j(\mu_k)} \\ & \leq \frac{1}{2} \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \mu_k + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \mu_k \mathbb{I} \left\{ \frac{\sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \mu_k)(\mathbf{x}_v \mu_k)^2\ell}}{\Phi_k^j(\mu_k)} > \frac{1}{2} \right\} \\ & \leq \frac{1}{2} \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \mu_k + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \mu_k \frac{4l_j\ell}{\Phi_k^j(\mu_k)} \end{aligned} \quad (5)$$

$$\leq \frac{1}{2} \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \mu_k + O(d^4 |\Lambda_2| \ell \log^3(dK)), \quad (6)$$

⁶This cannot be done explicitly in the algorithm, since it would re-couple the two sequences.

Algorithm 2 VARLin: Variance-Aware RL with Linear Function Approximation

- 1: **Initialize:** $\ell_i = 2^{2-i}$, $\iota = 16d \ln \frac{dHK}{\delta}$, $L_0 = \lceil \log_2 KH \rceil$, $L_1 = L_2 = \lceil 5 \log_2(HK) + 3 \rceil$, $\Lambda_0 = \{0, 1, \dots, L_0\}$, $\Lambda_1 = \{1, \dots, L_1\}$, $\Lambda_2 = \{1, \dots, L_2\}$. $\mathcal{B}$ be an $(HK)^{-3}$ -net of $\mathbb{B}_1^d(2)$ with size no larger than $(\frac{4}{HK})^{3d}$. $\Theta_1 = \mathbb{B}_1^d(1)$.
- 2: **for** $k = 1, 2, \dots, K$ **do**
- 3: **Optimistic Planning:**
- 4: **for** $h = H, H-1, \dots, 1$ **do**
- 5: For each $(s, a) \in \mathcal{S} \times \mathcal{A}$, let $Q_h^k(s, a) = \min\{1, r(s, a) + \max_{\theta \in \Theta_k} \sum_{i=1}^d \theta_i P_{s,a}^i V_{h+1}^k\}$.
- 6: For each $s \in \mathcal{S}$, let $V_h^k(s) = \max_{a \in \mathcal{A}} Q_h^k(s, a)$.
- 7: **end for**
- 8: **for** $h = 1, 2, \dots, H$ **do**
- 9: Choose action $a_h^k \leftarrow \arg \max_{a \in \mathcal{A}} Q_h^k(s_h^k, a)$, observe the next state s_{h+1}^k .
- 10: **end for**
- 11: **Construct Confidence Set:**
- 12: For $m \in \Lambda_0, h \in [H]$, define the input $\mathbf{x}_{k,h}^m = [P_{s_h^k, a_h^k}^1 (V_{h+1}^k)^{2^m}, \dots, P_{s_h^k, a_h^k}^d (V_{h+1}^k)^{2^m}]^\top$.
- 13: For $m \in \Lambda_0, h \in [H]$, define the variance estimate $\eta_{k,h}^m = \max_{\theta \in \Theta_k} \{\theta \mathbf{x}_{k,h}^{m+1} - (\theta \mathbf{x}_{k,h}^m)^2\}$.
- 14: Denote $\epsilon_{v,u}^m(\theta) = \theta \mathbf{x}_{v,u}^m - (V_{u+1}^v(s_{u+1}^v))^{2^m}$ for $m \in \Lambda_0, u \in [H], v \in [k-1]$
- 15: Define $\mathcal{T}_{k+1}^{m,i} = \{(v, u) \in [k] \times [H] : \eta_{v,u}^m \in (\ell_{i+1}, \ell_i]\}$, $\mathcal{T}_{k+1}^{m, L_1+1} = \{(v, u) \in [k] \times [H] : \eta_{v,u}^m \leq \ell_{L_1+1}\}$.
- 16: Define the confidence ball $\Theta_{k+1} = \bigcap_{m,i,j} \Theta_{k+1}^{m,i,j}$, where

$$\Theta_{k+1}^{m,i,j} = \left\{ \theta \in \mathbb{B}_1^d(1) : \left| \sum_{(v,u) \in \mathcal{T}_{k+1}^{m,i}} \text{clip}_j(\mathbf{x}_{v,u}^m \mu) \epsilon_{v,u}^m(\theta) \right| \leq 4 \sqrt{\sum_{(v,u) \in \mathcal{T}_{k+1}^{m,i}} \text{clip}_j^2(\mathbf{x}_{v,u}^m \mu) \eta_{v,u}^m} \iota + 4\ell_j \iota, \forall \mu \in \mathcal{B} \right\} \quad (7)$$

and $\text{clip}_j(\cdot) = \text{clip}(\cdot, \ell_j)$

17: **end for**

where (5) is by the fact that $\frac{\sqrt{\sum_{v=1}^{k-1} 2 \text{clip}_j^2(\mathbf{x}_v \mu_k) (\mathbf{x}_v \mu_k)^{2\iota}}}{\Phi_k^j(\mu_k)} > \frac{1}{2}$ implies that $\frac{4\ell_j \iota}{\Phi_k^j(\mu_k)} > 1$, and (6) follows by Lemma 17. By (4) and (6), we have that

$$\begin{aligned} \sum_{k=1}^K \mathbf{x}_k \mu_k &\leq 12 \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \mu_k \times \frac{\sqrt{\Psi_k^j(\mu_k) \iota}}{\Phi_k^j(\mu_k)} + \tilde{O}(d^5) \\ &\leq \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \frac{12 \mathbf{x}_k \mu_k \ell_j}{\Phi_k^j(\mu_k)} \sqrt{\sum_{k=1}^K \eta_k(\theta^*) \iota} + \tilde{O}(d^5) \leq O(d^4 |\Lambda_2| \log^3(dK)) \sqrt{\left(\ln \frac{1}{\delta} + \sum_{k=1}^K \sigma_k^2 \right) \iota} + \tilde{O}(d^5), \end{aligned}$$

where the last inequality uses Lemma 17. Therefore, the regret bound is $\tilde{O}\left(d^{4.5} \sqrt{\sum_{k=1}^K \sigma_k^2} + d^5\right)$.

See Section D for the full proof.

5 Algorithm and Theory for Linear Mixture MDP

We introduce our algorithm and the regret bound for linear mixture MDP. Its pseudo-code is listed in Algorithm 2 and its regret bound is stated below. The proof is deferred to Section E.

Theorem 5. *The expected regret of Algorithm 2 is bounded by $\mathbb{E}[\mathfrak{R}^K] \leq \tilde{O}(d^{4.5} \sqrt{K} + d^9)$.*

Before describing our algorithm, we introduce some additional notations. In this section, we assume that, unless explicitly stated, the variables m, i, j, k, h iterate over the sets $\Lambda_0, \Lambda_1, \Lambda_2, [K], [H]$, respectively. See Line 1 of Algorithm 2 for the definitions of these sets. For example, at Line 16 of Algorithm 2, we have $\bigcap_{m,i,j} \Theta_{k+1}^{m,i,j} = \bigcap_{m \in \Lambda_0, i \in \Lambda_1, j \in \Lambda_2} \Theta_{k+1}^{m,i,j}$.

The starting point of our algorithm design is from [Zhang et al. \[2020a\]](#), in which the authors obtained a nearly horizon-free regret bound in tabular MDP. A natural idea is to combine their proof with our results for linear bandits and obtain a nearly horizon-free regret bound for linear mixture MDP.

Note that, however, there is one caveat for such direct combination: in Section 4, the confidence set Θ_k is updated at a per-round level, in that Θ_k is built using all rounds prior to k ; while for the RL setting, the confidence set Θ_k could only be updated at a per-episode level and use all time steps prior to episode k . Were it updated at a per-time-step level, severe dependency issues would prevent us from bounding the regret properly. Such discrepancy in update frequency results in a gap between the confidence set built using data prior to episode k , and that built using data prior to time step (k, h) . Fortunately, we are able to resolve this issue. In Lemma 22, we show that we can relate these two confidence intervals, except for $\tilde{O}(d)$ “bad” episodes. Therefore, we could adapt the analysis in [Zhang et al. \[2020a\]](#) only for the not “bad” episodes, and we bound the regret by 1 for the “bad” episodes. The resulting regret bound should be $\tilde{O}(d^{6.5}\sqrt{K})$.

To further reduce the horizon-free regret bound to $\tilde{O}(d^{4.5}\sqrt{K})$, we present another novel technique. We first note an important advantage of the linear mixture MDP setting over the linear bandit setting: in the latter setting, we cannot estimate the variance because there is no structure on the variance among different actions; while in the former setting, we could estimate an upper bound of the variance, because the variance is a quadratic function of θ^* . Therefore, we can use the peeling technique on the *variance magnitude* to reduce the regret (comparing Equation (30) and Equation (43) in appendix). We note that one can also apply this step to linear bandits if the variance can be estimated.

Along the way, we also need to bound the gap between estimated variance and true variance, which can be seen as the “regret of variance predictions.” Using the same idea, we can build a confidence set using the variance sequence (x^2) , and the regret of variance predictions can be bounded by the variance of variance, namely the 4-th moment. Still, a peeling step on the 4-th moment is required to bound the regret of variance predictions, we need to bound the gap between estimated 4-th moment and true 4-th moment, which requires predicting 8-th moment. We continue to use this idea: we estimate 2-th, 4-th, 8-th, $\dots$, $O(\log KH)$ -th moments. The index m is used for moments, and Λ_0 is the index set reserved for moments. We note that the proof in [\[Zhang et al., 2020a\]](#) also depends on the higher moments. The main difference is here we estimate these higher moments explicitly.

6 Discussions

By incorporating the variance information in the confidence set construction, we derive the first variance-dependent regret bound for linear bandits and the nearly horizon-free regret bound for linear mixture MDP. Below we discuss limitations of our work and some future directions.

One drawback of our result is that our dependency on d is large. The main reason is our bounds rely on the convex potential lemma (Lemma 17), which is $\tilde{O}(d^4)$. In analogous to the elliptical potential lemma in [\[Abbasi-Yadkori et al., 2011\]](#), we believe that this bound can be improved to $\tilde{O}(d)$. This improvement will directly reduce the dependencies on d in our bounds.

Another drawback is that our method is not computationally efficient. This is a common issue in elimination-based algorithms. We note that the issue of computational tractability is common in sequential decision-making problems [\[Zhang and Ji, 2019, Wang et al., 2020a, Bartlett and Tewari, 2012, Zanette et al., 2020, Krishnamurthy et al., 2016, Jiang et al., 2017, Sun et al., 2019, Jin et al., 2021, Du et al., 2021, Dong et al., 2020\]](#). We leave it as a future direction to design computationally efficient algorithms that enjoy variance-dependent bounds for settings considered in this paper.

Lastly, in this paper, we only study linear function approximation. It would be interesting to generalize the ideas in this paper to other settings with function approximation schemes [\[Yang and Wang, 2019, Jin et al., 2020, Zanette et al., 2020, Wang et al., 2020c, Russo and Van Roy, 2013, Jiang et al., 2017, Sun et al., 2019, Du et al., 2021, Jin et al., 2021\]](#).

Acknowledgement

Simon S. Du gratefully acknowledges funding from NSF Award’s IIS-2110170 and DMS-2134106.

References

- Yasin Abbasi-Yadkori, Dávid Pál, and Csaba Szepesvári. Improved algorithms for linear stochastic bandits. *Advances in neural information processing systems*, 24:2312–2320, 2011.
- Marc Abeille, Louis Faury, and Clément Calauzènes. Instance-wise minimax-optimal algorithms for logistic bandits. In *International Conference on Artificial Intelligence and Statistics*, pages 3691–3699. PMLR, 2021.
- Jean-Yves Audibert, Rémi Munos, and Csaba Szepesvari. Use of variance estimation in the multi-armed bandit problem. 2006.
- Peter Auer, Nicolò Cesa-Bianchi, Yoav Freund, and Robert E. Schapire. The nonstochastic multi-armed bandit problem. *SIAM Journal on Computing*, 32(1):48–77, 2002.
- Alex Ayoub, Zeyu Jia, Csaba Szepesvari, Mengdi Wang, and Lin F Yang. Model-based reinforcement learning with value-targeted regression. In *Proceedings of the 37th International Conference on Machine Learning*, 2020.
- Mohammad Gheshlaghi Azar, Ian Osband, and Rémi Munos. Minimax regret bounds for reinforcement learning. In *Proceedings of the 34th International Conference on Machine Learning*, pages 263–272, 2017.
- Kazuoki Azuma. Weighted sums of certain dependent random variables. *Tohoku Mathematical Journal, Second Series*, 19(3):357–367, 1967.
- Peter L Bartlett and Ambuj Tewari. Regal: A regularization based algorithm for reinforcement learning in weakly communicating mdps. *arXiv preprint arXiv:1205.2661*, 2012.
- Qi Cai, Zhuoran Yang, Chi Jin, and Zhaoran Wang. Provably efficient exploration in policy optimization. *arXiv preprint arXiv:1912.05830*, 2019.
- Wei Chu, Lihong Li, Lev Reyzin, and Robert Schapire. Contextual bandits with linear payoff functions. In *Proceedings of the Fourteenth International Conference on Artificial Intelligence and Statistics*, pages 208–214, 2011.
- Varsha Dani, Thomas P Hayes, and Sham M Kakade. Stochastic linear optimization under bandit feedback. In *Conference on Learning Theory*, 2008.
- Christoph Dann, Nan Jiang, Akshay Krishnamurthy, Alekh Agarwal, John Langford, and Robert E. Schapire. On oracle-efficient PAC-RL with rich observations. In *Advances in Neural Information Processing Systems*, 2018.
- Christoph Dann, Lihong Li, Wei Wei, and Emma Brunskill. Policy certificates: Towards accountable reinforcement learning. In *Proceedings of the 36th International Conference on Machine Learning*, pages 1507–1516, 2019.
- Kefan Dong, Jian Peng, Yining Wang, and Yuan Zhou. Root-n-regret for learning in Markov decision processes with function approximation and low Bellman rank. In *Conference on Learning Theory*, pages 1554–1557. PMLR, 2020.
- Simon S Du, Akshay Krishnamurthy, Nan Jiang, Alekh Agarwal, Miroslav Dudik, and John Langford. Provably efficient RL with rich observations via latent state decoding. In *International Conference on Machine Learning*, pages 1665–1674, 2019a.
- Simon S Du, Yuping Luo, Ruosong Wang, and Hanrui Zhang. Provably efficient Q-learning with function approximation via distribution shift error checking oracle. In *Advances in Neural Information Processing Systems*, pages 8058–8068, 2019b.
- Simon S Du, Sham M Kakade, Ruosong Wang, and Lin F Yang. Is a good representation sufficient for sample efficient reinforcement learning? In *International Conference on Learning Representations*, 2020a.

- Simon S Du, Jason D Lee, Gaurav Mahajan, and Ruosong Wang. Agnostic Q-learning with function approximation in deterministic systems: Tight bounds on approximation error and sample complexity. *Advances in Neural Information Processing Systems*, 2020b.
- Simon S Du, Sham M Kakade, Jason D Lee, Shachar Lovett, Gaurav Mahajan, Wen Sun, and Ruosong Wang. Bilinear classes: A structural framework for provable generalization in RL. *arXiv preprint arXiv:2103.10897*, 2021.
- Louis Faury, Marc Abeille, Clément Calauzènes, and Olivier Fercoq. Improved optimistic algorithms for logistic bandits. In *International Conference on Machine Learning*, pages 3052–3060. PMLR, 2020.
- Fei Feng, Ruosong Wang, Wotao Yin, Simon S Du, and Lin F Yang. Provably efficient exploration for RL with unsupervised learning. *arXiv preprint arXiv:2003.06898*, 2020.
- Jiafan He, Dongruo Zhou, and Quanquan Gu. Logarithmic regret for reinforcement learning with linear function approximation. *arXiv preprint arXiv:2011.11566*, 2020.
- Yassir Jedra and Alexandre Proutiere. Optimal best-arm identification in linear bandits. *Advances in Neural Information Processing Systems*, 33, 2020.
- Nan Jiang and Alekh Agarwal. Open problem: The dependence of sample complexity lower bounds on planning horizon. In *Conference On Learning Theory*, pages 3395–3398, 2018.
- Nan Jiang, Akshay Krishnamurthy, Alekh Agarwal, John Langford, and Robert E Schapire. Contextual decision processes with low Bellman rank are PAC-learnable. In *Proceedings of the 34th International Conference on Machine Learning*, pages 1704–1713, 2017.
- Chi Jin, Zeyuan Allen-Zhu, Sebastien Bubeck, and Michael I Jordan. Is Q-learning provably efficient? In *Advances in Neural Information Processing Systems*, pages 4863–4873, 2018.
- Chi Jin, Zhuoran Yang, Zhaoran Wang, and Michael I Jordan. Provably efficient reinforcement learning with linear function approximation. In *Conference on Learning Theory*, pages 2137–2143, 2020.
- Chi Jin, Qinghua Liu, and Sobhan Miryoosefi. Bellman Eluder dimension: New rich classes of RL problems, and sample-efficient algorithms. *arXiv preprint arXiv:2102.00815*, 2021.
- Julian Katz-Samuels, Lalit Jain, Kevin G Jamieson, et al. An empirical process approach to the union bound: Practical algorithms for combinatorial and linear bandits. *Advances in Neural Information Processing Systems*, 33, 2020.
- Akshay Krishnamurthy, Alekh Agarwal, and John Langford. PAC reinforcement learning with rich observations. In *Advances in Neural Information Processing Systems*, pages 1840–1848, 2016.
- Tor Lattimore and Marcus Hutter. Pac bounds for discounted mdps. In *International Conference on Algorithmic Learning Theory*, pages 320–334. Springer, 2012.
- Gen Li, Yuting Wei, Yuejie Chi, Yuantao Gu, and Yuxin Chen. Breaking the sample size barrier in model-based reinforcement learning with a generative model. In *Advances in Neural Information Processing Systems*, 2020.
- Yingkai Li, Yining Wang, and Yuan Zhou. Nearly minimax-optimal regret for linearly parameterized bandits. In *Conference on Learning Theory*, pages 2173–2174, 2019a.
- Yingkai Li, Yining Wang, and Yuan Zhou. Tight regret bounds for infinite-armed linear contextual bandits. *arXiv preprint arXiv:1905.01435*, 2019b.
- Andreas Maurer and Massimiliano Pontil. Empirical Bernstein bounds and sample variance penalization. In *Conference on Learning Theory*, 2009.
- Pierre Menard, Omar Darwiche Domingues, Xuedong Shang, and Michal Valko. Ucb momentum q-learning: Correcting the bias without forgetting. *arXiv preprint arXiv:2103.01312*, 2021.

- Dipendra Misra, Mikael Henaff, Akshay Krishnamurthy, and John Langford. Kinematic state abstraction and provably efficient rich-observation reinforcement learning. *arXiv preprint arXiv:1911.05815*, 2019.
- Aditya Modi, Nan Jiang, Ambuj Tewari, and Satinder Singh. Sample complexity of reinforcement learning using linearly combined model ensembles. In *International Conference on Artificial Intelligence and Statistics*, pages 2010–2020. PMLR, 2020.
- Thomas Peel, Sandrine Anthoine, and Liva Ralaivola. Empirical bernstein inequality for martingales: Application to online learning. 2013.
- Dan Russo and Benjamin Van Roy. Eluder dimension and the sample complexity of optimistic exploration. In *Advances in Neural Information Processing Systems*, pages 2256–2264, 2013.
- Roshan Shariff and Csaba Szepesvári. Efficient planning in large mdps with weak linear function approximation. *arXiv preprint arXiv:2007.06184*, 2020.
- Wen Sun, Nan Jiang, Akshay Krishnamurthy, Alekh Agarwal, and John Langford. Model-based RL in contextual decision processes: PAC bounds and exponential improvements over model-free approaches. In *Conference on Learning Theory*, pages 2898–2933, 2019.
- Ruosong Wang, Simon S Du, Lin F Yang, and Sham M Kakade. Is long horizon reinforcement learning more difficult than short horizon reinforcement learning? In *Advances in Neural Information Processing Systems*, 2020a.
- Ruosong Wang, Simon S Du, Lin F Yang, and Ruslan Salakhutdinov. On reward-free reinforcement learning with linear function approximation. In *Advances in Neural Information Processing Systems*, 2020b.
- Ruosong Wang, Ruslan Salakhutdinov, and Lin F Yang. Provably efficient reinforcement learning with general value function approximation. *Advances in Neural Information Processing Systems*, 2020c.
- Yining Wang, Ruosong Wang, Simon S Du, and Akshay Krishnamurthy. Optimism in reinforcement learning with generalized linear function approximation. *arXiv preprint arXiv:1912.04136*, 2019.
- Gellert Weisz, Philip Amortila, and Csaba Szepesvári. Exponential lower bounds for planning in mdps with linearly-realizable optimal action-value functions. *arXiv preprint arXiv:2010.01374*, 2020.
- Zheng Wen and Benjamin Van Roy. Efficient exploration and value function generalization in deterministic systems. In *Advances in Neural Information Processing Systems*, pages 3021–3029, 2013.
- Lin Yang and Mengdi Wang. Sample-optimal parametric Q-learning using linearly additive features. In *International Conference on Machine Learning*, pages 6995–7004, 2019.
- Lin F Yang and Mengdi Wang. Reinforcement learning in feature space: Matrix bandit, kernels, and regret bound. *International Conference on Machine Learning*, 2020.
- Andrea Zanette and Emma Brunskill. Tighter problem-dependent regret bounds in reinforcement learning without domain knowledge using value function bounds. In *International Conference on Machine Learning*, pages 7304–7312, 2019.
- Andrea Zanette, Alessandro Lazaric, Mykel Kochenderfer, and Emma Brunskill. Learning near optimal policies with low inherent bellman error. In *International Conference on Machine Learning*, 2020.
- Zihan Zhang and Xiangyang Ji. Regret minimization for reinforcement learning by evaluating the optimal bias function. In *Advances in Neural Information Processing Systems*, pages 2823–2832, 2019.
- Zihan Zhang, Xiangyang Ji, and Simon S Du. Is reinforcement learning more difficult than bandits? a near-optimal algorithm escaping the curse of horizon. *arXiv preprint arXiv:2009.13503*, 2020a.

Zihan Zhang, Yuan Zhou, and Xiangyang Ji. Almost optimal model-free reinforcement learning via reference-advantage decomposition. In *Advances in Neural Information Processing Systems*, 2020b.

Zihan Zhang, Yuan Zhou, and Xiangyang Ji. Model-free reinforcement learning: from clipped pseudo-regret to sample complexity. *arXiv preprint arXiv:2006.03864*, 2020c.

Dongruo Zhou, Quanquan Gu, and Csaba Szepesvari. Nearly minimax optimal reinforcement learning for linear mixture markov decision processes. *arXiv preprint arXiv:2012.08507*, 2020a.

Dongruo Zhou, Jiafan He, and Quanquan Gu. Provably efficient reinforcement learning for discounted mdps with feature mapping. *arXiv preprint arXiv:2006.13165*, 2020b.

Checklist

1. For all authors...
 - (a) Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope? [\[Yes\]](#)
 - (b) Did you describe the limitations of your work? [\[Yes\]](#) We discuss the limitations in Section 6.
 - (c) Did you discuss any potential negative societal impacts of your work? [\[N/A\]](#) This work is theoretical so the boarder impact does not apply.
 - (d) Have you read the ethics review guidelines and ensured that your paper conforms to them? [\[Yes\]](#)
2. If you are including theoretical results...
 - (a) Did you state the full set of assumptions of all theoretical results? [\[Yes\]](#) We present the main assumption in Section 3.
 - (b) Did you include complete proofs of all theoretical results? [\[Yes\]](#) We present the proofs in Appendix.
3. If you ran experiments...
 - (a) Did you include the code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL)? [\[N/A\]](#) We have no experiments.
 - (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they were chosen)? [\[N/A\]](#)
 - (c) Did you report error bars (e.g., with respect to the random seed after running experiments multiple times)? [\[N/A\]](#)
 - (d) Did you include the total amount of compute and the type of resources used (e.g., type of GPUs, internal cluster, or cloud provider)? [\[N/A\]](#)
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...
 - (a) If your work uses existing assets, did you cite the creators? [\[N/A\]](#) We do not use existing models.
 - (b) Did you mention the license of the assets? [\[N/A\]](#)
 - (c) Did you include any new assets either in the supplemental material or as a URL? [\[N/A\]](#)
 - (d) Did you discuss whether and how consent was obtained from people whose data you’re using/curating? [\[N/A\]](#)
 - (e) Did you discuss whether the data you are using/curating contains personally identifiable information or offensive content? [\[N/A\]](#)
5. If you used crowdsourcing or conducted research with human subjects...
 - (a) Did you include the full text of instructions given to participants and screenshots, if applicable? [\[N/A\]](#) This work is irrelevant with human subjects.
 - (b) Did you describe any potential participant risks, with links to Institutional Review Board (IRB) approvals, if applicable? [\[N/A\]](#)
 - (c) Did you include the estimated hourly wage paid to participants and the total amount spent on participant compensation? [\[N/A\]](#)

A Technical Lemmas

Lemma 6 ([Azuma, 1967]). *Let $(M_n)_{n \geq 0}$ be a martingale such that $M_0 = 0$ and $|M_n - M_{n-1}| \leq b$ almost surely for every $n \geq 1$. Then we have*

$$\Pr\left[|M_n| \geq b\sqrt{2n \log(2/\delta)}\right] \leq \delta.$$

Lemma 7 ([Zhang et al., 2020c], Lemma 9). *Let $\{\mathcal{F}_i\}_{i \geq 0}$ be a filtration. Let $\{X_i\}_{i \geq 1}$ be a real-valued stochastic process adapted to $\{\mathcal{F}_i\}_{i \geq 0}$ such that $0 \leq X_i \leq 1$ almost surely and that X_i is $\mathcal{F}_i$ -measurable. For every $\delta \in (0, 1)$, $c \geq 1$, we have*

$$\Pr\left[\exists n \geq 1 : \sum_{i=1}^n \mathbb{E}[X_i | \mathcal{F}_{i-1}] \geq 4c \ln \frac{4}{\delta}, \sum_{i=1}^n X_i \leq c \ln \frac{4}{\delta}\right] \leq \delta.$$

Lemma 8. *Let $\{\mathcal{F}_i\}_{i \geq 0}$ be a filtration. Let $\{X_i\}_{i \geq 1}$ be a real-valued stochastic process adapted to $\{\mathcal{F}_i\}_{i \geq 0}$ such that $0 \leq X_i \leq 1$ almost surely and that X_i is $\mathcal{F}_i$ -measurable. For every $\delta \in (0, 1)$, $c \geq 1$, we have*

$$\Pr\left[\exists n \geq 1 : \sum_{i=1}^n X_i \geq 4c \ln \frac{4}{\delta}, \sum_{i=1}^n \mathbb{E}[X_i | \mathcal{F}_{i-1}] \leq c \ln \frac{4}{\delta}\right] \leq \delta.$$

Proof. We follow the proof of Lemma 9 in [Zhang et al., 2020c]. Let $\lambda > 0$ be a parameter, $\mu_i = \mathbb{E}[X_i | \mathcal{F}_{i-1}]$. Define $Y_n = \exp(\lambda \sum_{i=1}^n X_i - (e^\lambda - 1) \sum_{i=1}^n \mu_i)$ for $n \geq 0$. Note that $\mathbb{E}[e^{\lambda X}] \leq \mu e^\lambda + (1 - \mu) \leq e^{\mu(e^\lambda - 1)}$, so $\mathbb{E}[e^{\lambda X_i - (e^\lambda - 1)\mu_i} | \mathcal{F}_{i-1}] \leq 1$, thus $\{Y_n\}_{n \geq 0}$ is a super-martingale. Let $\tau = \min\{n : \sum_{i=1}^n X_i \geq 4c \ln(4/\delta)\}$ be a stopping time, then we have $|Y_{\min\{\tau, n\}}| \leq e^{\lambda(4c \ln(4/\delta) + 1)} < +\infty$ almost surely for every $n \geq 0$. Therefore, by the optional stopping theorem, we have $\mathbb{E}[Y_\tau] \leq 1$. Finally, we have

$$\begin{aligned} \Pr\left[\exists n \geq 1 : \sum_{i=1}^n X_i \geq 4c \ln \frac{4}{\delta}, \sum_{i=1}^n \mu_i \leq c \ln \frac{4}{\delta}\right] &\leq \Pr\left[\sum_{i=1}^{\tau} \mu_i \leq c \ln \frac{4}{\delta}\right] \\ &\leq \Pr\left[Y_\tau \geq \exp\left(\lambda \sum_{i=1}^{\tau} X_i - (e^\lambda - 1) c \ln \frac{4}{\delta}\right)\right] \\ &\leq \Pr\left[Y_\tau \geq \exp\left(\lambda(4c \ln \frac{4}{\delta} - 1) - (e^\lambda - 1) c \ln \frac{4}{\delta}\right)\right] \\ &\leq \exp\left(\lambda(1 - 4c \ln \frac{2}{\delta}) + (e^\lambda - 1) c \ln \frac{4}{\delta}\right) \\ &= e^\lambda e^{(e^\lambda - 1 - 4\lambda) c \ln(4/\delta)}. \end{aligned}$$

Choosing $\lambda = 1$, we have

$$e^\lambda e^{(e^\lambda - 1 - 4\lambda) c \ln(4/\delta)} \leq e \cdot e^{-2c \ln(4/\delta)} = e\left(\frac{\delta}{4}\right)^c \leq \frac{e}{4}\delta \leq \delta,$$

which concludes the proof. $\square$

Lemma 9. *Let $\{\mathcal{F}_i\}_{i \geq 0}$ be a filtration. Let $\{X_i\}_{i=1}^n$ be a sequence of random variables such that $|X_i| \leq 1$ almost surely, that X_i is $\mathcal{F}_i$ -measurable. For every $\delta \in (0, 1)$, we have*

$$\Pr\left[\sum_{i=1}^n \mathbb{E}[X_i^2 | \mathcal{F}_{i-1}] \geq \sum_{i=1}^n 8X_i^2 + 4 \ln \frac{4}{\delta}\right] \leq (\lceil \log_2 n \rceil + 1)\delta.$$

Proof. Let $Y = \sum_{i=1}^n \mathbb{E}[X_i^2 | \mathcal{F}_{i-1}]$, $Z = \sum_{i=1}^n X_i^2$. Applying Lemma 7 with the sequence $\{X_i^2\}_{i=1}^n$, we have for every $c \geq 1$,

$$\Pr\left[Y \geq 4c \ln \frac{4}{\delta}, Z \leq c \ln \frac{4}{\delta}\right] \leq \delta.$$

Therefore, we have

$$\begin{aligned}
& \Pr\left[Y \geq 8Z + 4 \ln \frac{4}{\delta}\right] \\
& \leq \sum_{j=1}^{\lfloor \log_2 n \rfloor} \Pr\left[Y \geq 8Z + 4 \ln \frac{4}{\delta}, 2^{j-1} \ln \frac{4}{\delta} \leq Z \leq 2^j \ln \frac{4}{\delta}\right] + \Pr\left[Y \geq 8Z + 4 \ln \frac{4}{\delta}, Z \leq \ln \frac{4}{\delta}\right] \\
& \leq \sum_{j=1}^{\lfloor \log_2 n \rfloor} \Pr\left[Y \geq 8Z, 2^{j-1} \ln \frac{4}{\delta} \leq Z \leq 2^j \ln \frac{4}{\delta}\right] + \Pr\left[Y \geq 4 \ln \frac{4}{\delta}, Z \leq \ln \frac{4}{\delta}\right] \\
& \leq \sum_{j=1}^{\lfloor \log_2 n \rfloor} \Pr\left[Y \geq 8 \cdot 2^{j-1} \ln \frac{4}{\delta}, 2^{j-1} \ln \frac{4}{\delta} \leq Z \leq 2^j \ln \frac{4}{\delta}\right] + \Pr\left[Y \geq 4 \ln \frac{4}{\delta}, Z \leq \ln \frac{4}{\delta}\right] \\
& \leq \sum_{j=1}^{\lfloor \log_2 n \rfloor} \Pr\left[Y \geq 4 \cdot 2^j \ln \frac{4}{\delta}, Z \leq 2^j \ln \frac{4}{\delta}\right] + \Pr\left[Y \geq 4 \ln \frac{4}{\delta}, Z \leq \ln \frac{4}{\delta}\right] \\
& \leq (\lfloor \log_2 n \rfloor + 1)\delta
\end{aligned}$$

as desired. $\square$

Lemma 10. Let $\{\mathcal{F}_i\}_{i \geq 0}$ be a filtration. Let $\{X_i\}_{i=1}^n$ be a sequence of random variables such that $|X_i| \leq 1$ almost surely, that X_i is $\mathcal{F}_i$ -measurable. For every $\delta \in (0, 1)$, we have

$$\Pr\left[\sum_{i=1}^n X_i^2 \geq \sum_{i=1}^n 8\mathbb{E}[X_i^2 \mid \mathcal{F}_{i-1}] + 4 \ln \frac{4}{\delta}\right] \leq (\lfloor \log_2 n \rfloor + 1)\delta.$$

Proof. Let $Y = \sum_{i=1}^n X_i^2$, $Z = \sum_{i=1}^n \mathbb{E}[X_i^2 \mid \mathcal{F}_{i-1}]$. Applying Lemma 8 with the sequence $\{X_i^2\}_{i=1}^n$, we have for every $c \geq 1$,

$$\Pr\left[Y \geq 4c \ln \frac{4}{\delta}, Z \leq c \ln \frac{4}{\delta}\right] \leq \delta.$$

Therefore, we have

$$\begin{aligned}
& \Pr\left[Y \geq 8Z + 4 \ln \frac{4}{\delta}\right] \\
& \leq \sum_{j=1}^{\lfloor \log_2 n \rfloor} \Pr\left[Y \geq 8Z + 4 \ln \frac{4}{\delta}, 2^{j-1} \ln \frac{4}{\delta} \leq Z \leq 2^j \ln \frac{4}{\delta}\right] + \Pr\left[Y \geq 8Z + 4 \ln \frac{4}{\delta}, Z \leq \ln \frac{4}{\delta}\right] \\
& \leq \sum_{j=1}^{\lfloor \log_2 n \rfloor} \Pr\left[Y \geq 8Z, 2^{j-1} \ln \frac{4}{\delta} \leq Z \leq 2^j \ln \frac{4}{\delta}\right] + \Pr\left[Y \geq 4 \ln \frac{4}{\delta}, Z \leq \ln \frac{4}{\delta}\right] \\
& \leq \sum_{j=1}^{\lfloor \log_2 n \rfloor} \Pr\left[Y \geq 8 \cdot 2^{j-1} \ln \frac{4}{\delta}, 2^{j-1} \ln \frac{4}{\delta} \leq Z \leq 2^j \ln \frac{4}{\delta}\right] + \Pr\left[Y \geq 4 \ln \frac{4}{\delta}, Z \leq \ln \frac{4}{\delta}\right] \\
& \leq \sum_{j=1}^{\lfloor \log_2 n \rfloor} \Pr\left[Y \geq 4 \cdot 2^j \ln \frac{4}{\delta}, Z \leq 2^j \ln \frac{4}{\delta}\right] + \Pr\left[Y \geq 4 \ln \frac{4}{\delta}, Z \leq \ln \frac{4}{\delta}\right] \\
& \leq (\lfloor \log_2 n \rfloor + 1)\delta
\end{aligned}$$

as desired. $\square$

Lemma 11 ([Zhang et al., 2020c], Lemma 11). Let $(M_n)_{n \geq 0}$ be a martingale such that $M_0 = 0$ and $|M_n - M_{n-1}| \leq b$ almost surely for every $n \geq 1$. For each $n \geq 0$, let $\mathcal{F}_n = \sigma(M_0, \dots, M_n)$ and let $\text{Var}_n = \sum_{i=1}^n \mathbb{E}[(M_i - M_{i-1})^2 \mid \mathcal{F}_{i-1}]$. Then for any $n \geq 1$ and $\epsilon, \delta > 0$, we have

$$\Pr\left[|M_n| \geq 2\sqrt{2\text{Var}_n \ln(1/\delta)} + 2\sqrt{\epsilon \ln(1/\delta)} + 2b \ln(1/\delta)\right] \leq 2(\log_2(b^2 n / \epsilon) + 1)\delta.$$

Lemma 12. Let $\lambda_1, \lambda_2, \lambda_4 > 0$, $\lambda_3 \geq 1$ and $\kappa = \max\{\log_2(\lambda_1), 1\}$. Let $a_1, a_2, \dots, a_\kappa$ be non-negative reals such that $a_i \leq \lambda_1$ and $a_i \leq \lambda_2 \sqrt{a_i + a_{i+1} + 2^{i+1} \lambda_3} + \lambda_4$ for any $1 \leq i \leq \kappa$ (with $a_{\kappa+1} = \lambda_1$). Then we have that

$$a_1 \leq 22\lambda_2^2 + 6\lambda_4 + 4\lambda_2 \sqrt{2\lambda_3}.$$

Proof. Note that

$$a_i \leq \lambda_2 \sqrt{a_i} + \lambda_2 \sqrt{a_{i+1} + 2^{i+1} \lambda_3} + \lambda_4,$$

so we have

$$a_i \leq \left(\lambda_2 + \sqrt{\lambda_2 \sqrt{a_{i+1} + 2^{i+1} \lambda_3} + \lambda_4} \right)^2 \leq 2\lambda_2^2 + 2\lambda_2 \sqrt{a_{i+1} + 2^{i+1} \lambda_3} + 2\lambda_4.$$

By Lemma 11 in [Zhang et al., 2020a], we have

$$\begin{aligned} a_1 &\leq \max \left\{ \left(2\lambda_2 + \sqrt{(2\lambda_2)^2 + (2\lambda_2^2 + 2\lambda_4)} \right)^2, 2\lambda_2 \sqrt{8\lambda_3} + 2\lambda_2^2 + 2\lambda_4 \right\} \\ &\leq \max\{20\lambda_2^2 + 4\lambda_4, 2\lambda_2 \sqrt{8\lambda_3} + 2\lambda_2^2 + 2\lambda_4\} \leq 22\lambda_2^2 + 6\lambda_4 + 4\lambda_2 \sqrt{2\lambda_3}, \end{aligned}$$

which concludes the proof. $\square$

B Limitations of Previous Approaches

In the example in Section 1, if we know $x_i \leq \sqrt{\frac{1}{K}}$ for $1 \leq i \leq K$, the best confidence region for θ^* should be $\Theta_t = \{\theta \mid \|\theta - \hat{\theta}_t\|_{\Lambda_{t-1}} \leq C(\sigma\sqrt{d} + \lambda^{1/2})\}$, and we can obtain a variance-aware regret bound by letting $\lambda = \sigma^2$. However, if we let $x_{K+1} = 1$ and use the same concentration inequality as before, the confidence region would be $\Theta_{K+1} = \{\theta \mid \|\theta - \hat{\theta}_t\|_{\Lambda_{t-1}} \leq C(\sigma\sqrt{d} + 1 + \lambda^{1/2})\}$.

We present the detailed computation as below. Choose $\theta^* = \Theta(1)$. $\theta^* - \hat{\theta}_{K+1} = -\frac{\sum_{i=1}^{K+1} x_i \epsilon_i}{\lambda + \sum_{i=1}^{K+1} x_i^2} + \frac{\lambda \theta^*}{\lambda + \sum_{i=1}^{K+1} x_i^2}$. When ϵ_i is bounded in $[-1, 1]$ with variance σ^2 , following Bernstein inequality, we have that $\left| \frac{\sum_{i=1}^{K+1} x_i \epsilon_i}{\lambda + \sum_{i=1}^{K+1} x_i^2} \right| \leq \frac{\sqrt{\sigma^2 \sum_{i=1}^{K+1} x_i^2} + \max_i x_i}{\lambda + \sum_{i=1}^{K+1} x_i^2}$. Therefore, the best confidence interval we have is

$$\begin{aligned} \|\theta^* - \hat{\theta}_{K+1}\|_{\Lambda_K} &\lesssim \sqrt{\frac{\sigma^2 \sum_{i=1}^{K+1} x_i^2}{\lambda + \sum_{i=1}^{K+1} x_i^2}} + \frac{\max_i x_i}{\sqrt{\lambda + \sum_{i=1}^{K+1} x_i^2}} + \frac{\lambda \theta^*}{\sqrt{\lambda + \sum_{i=1}^{K+1} x_i^2}} \\ &= \Theta \left(\sqrt{\frac{\sigma^2}{\lambda + 1}} + \frac{1 + \lambda}{\sqrt{1 + \lambda}} \right), \end{aligned}$$

i.e., $\|\theta^* - \hat{\theta}_{K+1}\|_{\Lambda_K} \lesssim \Theta(\sigma + \lambda^{1/2} + 1)$. Therefore, to maintain a confidence region for the general case following methods in [Zhou et al., 2020a, Fauray et al., 2020], the term $1 + \lambda^{1/2}$ is unavoidable.

This counter example highlights the necessity of our peeling step in the algorithm.

Remark 4. We note that for σ -sub-Gaussian noise (instead of σ^2 variance and 1-sub-Gaussian), one can ensure that $\left| \frac{\sum_{i=1}^{K+1} x_i \epsilon_i}{\lambda + \sum_{i=1}^{K+1} x_i^2} \right| \leq \frac{\sqrt{\sigma^2 \sum_{i=1}^{K+1} x_i^2}}{\lambda + \sum_{i=1}^{K+1} x_i^2}$, which help to reduce the width of confidence interval and obtain $\|\theta^* - \hat{\theta}_{K+1}\|_{\Lambda_K} \leq O(\sigma + \lambda^{1/2})$.

C Proof of Lemma 1

In this section, we present the proof of Lemma 1.

Restatement of Lemma 1 Let $f(x) \geq 0$ be a convex function over $\mathbb{R}$ such that $\frac{f(x)}{x^2} \leq \frac{f(y)}{y^2} \leq 1$ and $f(x) \geq f(y)$ if $x^2 \geq y^2 > 0$. Fix $\ell \in (0, 1]$. For any $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_t \in \mathbb{B}_2^d(1)$ and $\boldsymbol{\mu}_1, \boldsymbol{\mu}_2, \dots, \boldsymbol{\mu}_t \in \mathbb{B}_2^d(1)$, we have that

$$\sum_{i=1}^t \min \left\{ \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^{i-1} f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2}, 1 \right\} \leq O(d^4 \log(Cdt/\ell)). \quad (8)$$

Let $f(x)$ and ℓ be fixed. To prove Lemma 1, we have the lemmas below.

Lemma 13. For any $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_t \in \mathbb{B}_2^d(1)$ and $\boldsymbol{\mu}_1, \boldsymbol{\mu}_2, \dots, \boldsymbol{\mu}_n \in \mathbb{B}_2^d(1)$, we have that

$$\sum_{i=1}^t \min \left\{ \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2}, 1 \right\} \leq O(d \log(Cdt/\ell)). \quad (9)$$

Lemma 14. Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_t \in \mathbb{B}_2^d(1)$ be a sequence of vectors. If there exists a sequence $0 = \tau_0 < \tau_1 < \tau_2 < \dots < \tau_z = t$ such that for each $1 \leq \zeta \leq z$, there exists $\boldsymbol{\mu}_\zeta \in \mathbb{B}_2^d(1)$ such that

$$\sum_{i=1}^{\tau_\zeta} f(\mathbf{x}_i \boldsymbol{\mu}_\zeta) + \ell^2 > 4(d+2)^2 \times \left(\sum_{i=1}^{\tau_{\zeta-1}} f(\mathbf{x}_i \boldsymbol{\mu}_\zeta) + \ell^2 \right), \quad (10)$$

then $z \leq O(d \log^2(dt/\ell))$.

We present the proofs of Lemma 13 and 14 respectively in Section C.1 and C.2. Given these two lemmas, we continue analysis as below.

Let $\tau_0 = 0$ and for $i \geq 1$, we let

$$\tau_i = \min\{t+1\} \cup \left\{ \tau \mid \exists \tau_{i-1} \leq \tau' < \tau, \sum_{j=1}^{\tau} f(\mathbf{x}_j \boldsymbol{\mu}_{\tau'}) + \ell^2 > 4(d+2)^2 \left(\sum_{j=1}^{\tau'} f(\mathbf{x}_j \boldsymbol{\mu}_{\tau'}) + \ell^2 \right) \right\}.$$

Let $k = \min\{i \mid \tau_i = t+1\}$. Then k is well-defined and $k \leq O(d \log^2(dt))$ by Lemma 14. Furthermore, for any $\kappa < k$ and any $\tau_\kappa \leq i_1 < i_2 < \tau_{\kappa+1}$, we have

$$\sum_{j=1}^{i_2} f(\mathbf{x}_j \boldsymbol{\mu}_{i_1}) + \ell^2 \leq 4(d+2)^2 \left(\sum_{j=1}^{i_1} f(\mathbf{x}_j \boldsymbol{\mu}_{i_1}) + \ell^2 \right). \quad (11)$$

Now we are ready to prove Lemma 1. We have

$$\begin{aligned} \sum_{i=1}^t \min \left\{ \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^{i-1} f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2}, 1 \right\} &\leq 2 \sum_{i=1}^t \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^i f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} \\ &\leq 8(d+2)^2 \sum_{\kappa=1}^k \left(\sum_{i=\tau_{\kappa-1}}^{\tau_\kappa-1} \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^{\tau_\kappa-1} f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} \right) \end{aligned} \quad (12)$$

$$\begin{aligned} &\leq 8(d+2)^2 \sum_{\kappa=1}^k \left(\sum_{i=\tau_{\kappa-1}}^{\tau_\kappa-1} \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=\tau_{\kappa-1}}^{\tau_\kappa-1} f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} \right), \\ &\leq k \times O(d^2) \times O(d \log(t/\ell)) \leq O(d^4 \log^3(dt)), \end{aligned} \quad (13)$$

where (12) uses (11) and (13) uses Lemma 13.

C.1 Proof of Lemma 13

Restatement of Lemma 13 For any $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_t \in \mathbb{B}_2^d(1)$ and $\boldsymbol{\mu}_1, \boldsymbol{\mu}_2, \dots, \boldsymbol{\mu}_n \in \mathbb{B}_2^d(1)$, we have that

$$\sum_{i=1}^t \min \left\{ \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2}, 1 \right\} \leq O(d \log(Cdt/\ell)). \quad (14)$$

Proof. Let S_t be the permutation group over $[t]$. We claim that if

$$\sum_{i=1}^t \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} = \max_{\xi \in S_t} \sum_{i=1}^t \frac{f(\mathbf{x}_{\xi(i)} \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_{\xi(j)} \boldsymbol{\mu}_i) + \ell^2}, \quad (15)$$

then there exists some i such that $(\mathbf{x}_i \boldsymbol{\mu}_i)^2 \geq (\mathbf{x}_j \boldsymbol{\mu}_i)^2$ for any $j \in [t]$. Otherwise, we construct a directed graph $G = (V, E)$ where $V = [t]$ and edge (i, j) with $i \neq j$ is in E if and only if $(\mathbf{x}_j \boldsymbol{\mu}_i)^2 \geq (\mathbf{x}_{j'} \boldsymbol{\mu}_i)^2$ for any $j' \in [t]$. Let $d(i)$ be the out degree of i . By assuming $\{(\mathbf{x}_i \boldsymbol{\mu}_i)^2 \geq (\mathbf{x}_j \boldsymbol{\mu}_i)^2, \forall j \in [t]\}$ fails to hold, we learn that $d(i) \geq 1$ for every i , so there exists a circle $(i_1, i_2, \dots, i_k)$ in G . Consider the permutation ξ such that $\xi(i_j) = i_{j+1}$ for $j \in [k]$ (with $i_{k+1} := i_1$) and $\xi(i) = i$ for $i \notin \{i_1, \dots, i_k\}$. By definition, we have $(\boldsymbol{\mu}_{i_j} \mathbf{x}_{\xi(i_j)})^2 > (\boldsymbol{\mu}_{i_j} \mathbf{x}_{i_j})^2$ for $j \in [k]$, which implies that $f(\boldsymbol{\mu}_{i_j} \mathbf{x}_{\xi(i_j)}) > f(\boldsymbol{\mu}_{i_j} \mathbf{x}_{i_j})$ for $j \in [k]$. Therefore

$$\sum_{i=1}^t \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} < \sum_{i=1}^t \frac{f(\mathbf{x}_{\xi(i)} \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} = \sum_{i=1}^t \frac{f(\mathbf{x}_{\xi(i)} \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_{\xi(j)} \boldsymbol{\mu}_i) + \ell^2},$$

which leads to contradiction.

We assume that (15) holds, otherwise we can bound an upper bound of the original quantity. Therefore, we can find an index i such that $(\mathbf{x}_i \boldsymbol{\mu}_i)^2 \geq (\mathbf{x}_j \boldsymbol{\mu}_i)^2$ for any $j \in [t]$. Without loss of generality, we assume $i = 1$. Because $\frac{f(x)}{x^2}$ is decreasing in x , so we have

$$\frac{f(\mathbf{x}_1 \boldsymbol{\mu}_1)}{(\mathbf{x}_1 \boldsymbol{\mu}_1)^2} \leq \frac{f(\mathbf{x}_j \boldsymbol{\mu}_1)}{(\mathbf{x}_j \boldsymbol{\mu}_1)^2}$$

for any $j \in [t]$, which implies

$$\frac{f(\mathbf{x}_1 \boldsymbol{\mu}_1)}{\sum_{j=1}^t f(\mathbf{x}_j \boldsymbol{\mu}_1) + \ell^2} = \frac{(\mathbf{x}_1 \boldsymbol{\mu}_1)^2}{\left(\sum_{j=1}^t f(\mathbf{x}_j \boldsymbol{\mu}_1) + \ell^2\right) \cdot \frac{(\mathbf{x}_1 \boldsymbol{\mu}_1)^2}{f(\mathbf{x}_1 \boldsymbol{\mu}_1)}} \leq \frac{(\mathbf{x}_1 \boldsymbol{\mu}_1)^2}{\sum_{j=1}^t (\mathbf{x}_j \boldsymbol{\mu}_1)^2 + \ell^2}. \quad (16)$$

Therefore, we have

$$\begin{aligned} \sum_{i=1}^t \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_i \boldsymbol{\mu}_i) + \ell^2} &\leq \frac{(\mathbf{x}_1 \boldsymbol{\mu}_1)^2}{\sum_{j=1}^t (\mathbf{x}_j \boldsymbol{\mu}_1)^2 + \ell^2} + \sum_{i=2}^t \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_i \boldsymbol{\mu}_i) + \ell^2} \\ &\leq \frac{(\mathbf{x}_1 \boldsymbol{\mu}_1)^2}{\sum_{j=1}^t (\mathbf{x}_j \boldsymbol{\mu}_1)^2 + \ell^2} + \sum_{i=2}^t \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=2}^t f(\mathbf{x}_i \boldsymbol{\mu}_i) + \ell^2}. \end{aligned} \quad (17)$$

Similarly, we can show that there exists a permutation $\xi^* \in S_t$ such that

$$\sum_{i=1}^t \frac{f(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^t f(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} \leq \sum_{i=1}^t \frac{(\mathbf{x}_{\xi^*(i)} \boldsymbol{\mu}_i)^2}{\sum_{j=i}^t (\mathbf{x}_{\xi^*(j)} \boldsymbol{\mu}_i)^2 + \ell^2}. \quad (18)$$

Finally, by Lemma 15, we have that

$$\sum_{i=1}^t \frac{(\mathbf{x}_{\xi^*(i)} \boldsymbol{\mu}_i)^2}{\sum_{j=i}^t (\mathbf{x}_{\xi^*(j)} \boldsymbol{\mu}_i)^2 + \ell^2} = \sum_{i=1}^t \min \left\{ \frac{(\mathbf{x}_{\xi^*(i)} \boldsymbol{\mu}_i)^2}{\sum_{j=i}^t (\mathbf{x}_{\xi^*(j)} \boldsymbol{\mu}_i)^2 + \ell^2}, 1 \right\} \leq O(d \log(t/\ell)).$$

□

C.2 Proof of Lemma 14

Restatement of Lemma 14 Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_t \in \mathbb{B}_2^d(1)$ be a sequence of vectors. If there exists a sequence $0 = \tau_0 < \tau_1 < \tau_2 < \dots < \tau_z = t$ such that for each $1 \leq \zeta \leq z$, there exists $\boldsymbol{\mu}_\zeta \in \mathbb{B}_2^d(1)$ such that

$$\sum_{i=1}^{\tau_\zeta} f(\mathbf{x}_i \boldsymbol{\mu}_\zeta) + \ell^2 > 4(d+2)^2 \times \left(\sum_{i=1}^{\tau_\zeta-1} f(\mathbf{x}_i \boldsymbol{\mu}_\zeta) + \ell^2 \right), \quad (19)$$

then $z \leq O(d \log^2(dt/\ell))$.

Proof. If $f(1) \leq \ell^2/t$, then the conclusion holds trivially because $0 \leq f(x) \leq f(1) \leq \ell^2/t$ for all $x \in [-1, 1]$. Suppose $f(1) > \ell^2/t$. Since $\frac{f(x)}{x^2} \leq \frac{f(y)}{y^2} \leq 1$ for all $x^2 \geq y^2$, we have that for $0 < \lambda \leq 1$ and any $x \in \mathbb{R}$, $f(\lambda x) \geq \lambda^2 f(x)$.

Let $\mathbf{e}_i = [0, \dots, 1, \dots, 0]$ be the one-hot vector whose only 1 entry is at its i -th coordinate. Noting that $f(x) \leq x^2$, $|\mathbf{x}_i \boldsymbol{\mu}_\zeta| \leq \|\boldsymbol{\mu}_\zeta\|_2$ and

$$\sum_{i=1}^{\tau_\zeta} f(\mathbf{x}_i \boldsymbol{\mu}_\zeta) > 4(d+2)^2 \times \left(\sum_{i=1}^{\tau_\zeta-1} f(\mathbf{x}_i \boldsymbol{\mu}_\zeta) + \ell^2 \right) - \ell^2 \geq 4d^2 \ell^2$$

we have that $\|\boldsymbol{\mu}_\zeta\|_2 \geq \sqrt{\frac{4d^2 \ell^2}{t}}$. Define $E_\tau(\boldsymbol{\mu}) = \sum_{i=1}^t f(\mathbf{x}_i \boldsymbol{\mu}) + \frac{\ell^2}{d} \sum_{i=1}^d f(\mathbf{e}_i \boldsymbol{\mu})$. Then $E_\tau(\boldsymbol{\mu})$ is convex in $\boldsymbol{\mu}$ because $f(x)$ is convex in x . By definition, we have that

$$E_\tau(\boldsymbol{\mu}) \leq \sum_{i=1}^{\tau} f(\mathbf{x}_i \boldsymbol{\mu}) + \ell^2.$$

By (19), we have that

$$E_{\tau_\zeta}(\boldsymbol{\mu}_\zeta) \geq \sum_{i=1}^{\tau_\zeta} f(\mathbf{x}_i \boldsymbol{\mu}_\zeta) \geq 4d^2 \left(\sum_{i=1}^{\tau_\zeta-1} f(\mathbf{x}_i \boldsymbol{\mu}_\zeta) + \ell^2 \right) \geq 4d^2 E_{\tau_\zeta-1}(\boldsymbol{\mu}_\zeta). \quad (20)$$

Define

$$\Lambda = \{i \in \mathbb{Z} : \lfloor \log_2(d\ell^4/t^2) + 2 \rfloor \leq i \leq 2 \lfloor \log_2 t + 2 \rfloor\}.$$

We consider the convex set $D_{\tau,i} = \{\boldsymbol{\mu} : E_\tau(\boldsymbol{\mu}) \leq 2^i\}$ for $i \in \Lambda$. Let ζ be fixed. Because $\|\boldsymbol{\mu}_\zeta\| \geq \sqrt{\frac{4d^2 \ell^2}{t}}$ and $\sup_i f(\mathbf{e}_i \boldsymbol{\mu}) \geq \frac{4d\ell^2}{t} \cdot f(1) \geq \frac{4d\ell^4}{t^2}$, we have that $\frac{4d\ell^4}{t^2} \leq E_\tau(\boldsymbol{\mu}_\zeta) \leq t + \ell^2 \leq t+1$ for any $1 \leq \tau \leq t$. Then we can find $i_\zeta \in \Lambda$ such that $E_{\tau_{\zeta-1}}(\boldsymbol{\mu}_\zeta) \in (2^{i_\zeta-1}, 2^{i_\zeta}]$, which means that $\boldsymbol{\mu}_\zeta \in D_{\tau_{\zeta-1}, i_\zeta}$. Note that for $0 \leq \lambda \leq 1$, $f(\lambda x) \geq \lambda^2 f(x)$ for any x , it then follows that $E_t(\lambda \boldsymbol{\mu}) \geq \lambda^2 E_t(\boldsymbol{\mu})$ for any $t, \boldsymbol{\mu}$. Choosing $\lambda = \frac{1}{d}$, we have that $E_{\tau_\zeta}(\frac{\boldsymbol{\mu}_\zeta}{d}) \geq \frac{1}{d^2} E_{\tau_\zeta}(\boldsymbol{\mu}_\zeta) \geq 4E_{\tau_{\zeta-1}}(\boldsymbol{\mu}_\zeta) \geq 2^{i_\zeta}$. Therefore, $\frac{\boldsymbol{\mu}_\zeta}{d} \notin D_{\tau_\zeta, i_\zeta}$. In words, the intercept of D_{τ_ζ, i_ζ} in the direction $\boldsymbol{\mu}_\zeta$ is at most $1/d$ times of that of $D_{\tau_{\zeta-1}, i_\zeta}$.

Note that $D_{t,i}$ is decreasing in t for any i , so by Lemma 16, we have

$$\text{Volume}(D_{\tau_\zeta, i_\zeta}) \leq \frac{6}{7} \text{Volume}(D_{\tau_{\zeta-1}, i_\zeta}).$$

Also note that $\text{Volume}(D_{0,i}) \leq (\frac{2t}{\ell})^d$ and $\text{Volume}(D_{t,i}) \geq (\frac{1}{dt^3})^d$, so we conclude that $z \leq d|\Lambda| \log_{7/6}(2dt^4/\ell) \leq O(d \log^2(td/\ell))$.

□

C.3 Other Lemmas and Proofs

Lemma 15. Fix $\ell \in (0, 1]$. Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_t \in \mathbb{B}_2^d(1)$ and $\boldsymbol{\mu}_1, \boldsymbol{\mu}_2, \dots, \boldsymbol{\mu}_t \in \mathbb{B}_2^d(1)$ be two sequences of vectors. Then we have

$$\sum_{i=1}^t \mathbb{I} \left\{ (\mathbf{x}_i \boldsymbol{\mu}_i)^2 > \sum_{j=1}^{i-1} (\mathbf{x}_j \boldsymbol{\mu}_i)^2 + \ell^2 \right\} \leq \sum_{i=1}^t \min \left\{ \frac{(\mathbf{x}_i \boldsymbol{\mu}_i)^2}{\sum_{j=1}^{i-1} (\mathbf{x}_j \boldsymbol{\mu}_i)^2 + \ell^2}, 1 \right\} \leq O(d \log \frac{t}{\ell}). \quad (21)$$

Proof. The first inequality in (21) holds clearly. To prove the second inequality, we define $\mathbf{U}_0 = \ell^2 \mathbf{I}$ and $\mathbf{U}_i = \ell^2 \mathbf{I} + \sum_{j=1}^i \mathbf{x}_j \mathbf{x}_j^\top$ for $i \geq 1$. Note that

$$\frac{(\mathbf{x}_i \boldsymbol{\mu}_i)^2}{\sum_{j=1}^{i-1} (\mathbf{x}_j \boldsymbol{\mu}_i)^2 + \ell^2} \leq \frac{(\mathbf{x}_i \boldsymbol{\mu}_i)^2}{\boldsymbol{\mu}_i^\top \mathbf{U}_{i-1} \boldsymbol{\mu}_i} \leq \mathbf{x}_i^\top \mathbf{U}_{i-1}^{-1} \mathbf{x}_i,$$

where the first inequality is because $\|\mu_i\|_2 \leq 1$ and the second inequality uses the Cauchy's inequality, so we have

$$\sum_{i=1}^t \min \left\{ \frac{(\mathbf{x}_i \mu_i)^2}{\sum_{j=1}^{i-1} (\mathbf{x}_j \mu_i)^2 + \ell^2}, 1 \right\} \leq \sum_{i=1}^t \min \{ \mathbf{x}_i^\top \mathbf{U}_{i-1}^{-1} \mathbf{x}_i, 1 \} \leq 2d \ln(t/\ell^2) \leq 4d \ln(t/\ell),$$

where the second-to-third inequality uses the elliptical potential lemma. $\square$

Lemma 16. *Given $x \in \mathbb{R}^d$, we use $(u(x), l(x))$ to denote the polar coordinate of x where $\|x\|_2 = \frac{x}{\|x\|_2}$ is the direction and $l(x) = \|x\|_2$. We also use (u, ℓ) to denote the unique element x in $\mathbb{R}^d$ such that $(u(x), l(x)) = (u, \ell)$. Let D be a bounded symmetric convex subset of $\mathbb{R}^d$ with $d \geq 2$. Given any direction $\mu \in \partial \mathbb{B}_d$, there exists a unique $l(u) \in \mathbb{R}$ such that $(u, l(u)), (-u, l(u)) \in \partial D$ are on its boundary. Let D' be a bounded symmetric convex subset of $\mathbb{R}^d$ containing $D \subseteq D'$ such that $(u, d \cdot l(u)) \in D'$ for some direction $u \in \partial \mathbb{B}_d$. Then we have that*

$$\text{Volume}(D') \geq \frac{7}{6} \text{Volume}(D).$$

Proof. Let $A = (u, l(u))$ and $B = (u, d \cdot l(u))$. Since A is on the boundary of D , we can find a hyperplane h_1 such that $A \in h_1$ and h_1 is tangent to D . Let h_2 be the parallel hyperplane of h_1 containing the origin $O \in h_2$. Define

$$H = \left\{ x \in \mathbb{R}^d \mid d(x, h_1) + d(x, h_2) = d(h_1, h_2), \exists y \in D, \lambda \in \mathbb{R}, (B - y) = \lambda(B - x) \right\}$$

It is obvious that $\text{Volume}(H) \geq \frac{1}{2} \text{Volume}(D)$ since for each $x \in D$ lying between h_1 and h_2 , $x \in H$. Define

$$U = \left\{ x \in \mathbb{R}^d \mid d(x, h_2) = d(x, h_1) + d(h_1, h_2), \exists y \in H, \lambda \in [0, 1], x = \lambda y + (1 - \lambda)B \right\}.$$

We claim that

$$\text{Volume}(U) = \left(1 - \frac{1}{d}\right)^d \text{Volume}(U \cup H) = \left(1 - \frac{1}{d}\right)^d (\text{Volume}(U) + \text{Volume}(H)). \quad (22)$$

To see the first equality, we note that U and $U \cup H$ are both d -dimensional pyramids. It then follows from the volume formula and the relation $d(B, O) = d \times d(A, O)$. The second equality is because by their definitions, U, H are separated by the hyperplane h_1 , and thus they are disjoint. Finally, by (22), we have

$$\begin{aligned} \text{Volume}(D') &\geq \text{Volume}(U) + \text{Volume}(H) = \left(1 + \frac{1}{1 - (1 - 1/d)^d}\right) \text{Volume}(H) \\ &\geq \frac{1}{2} \left(1 + \frac{1}{(1 - (1 - 1/d)^d)}\right) \text{Volume}(D) \geq \frac{7}{6} \text{Volume}(D). \end{aligned}$$

$\square$

D Missing Proofs in Section 4

D.1 Application of the General Potential Lemma

As an application of Lemma 1 on linear bandit and linear RL, we have the lemma as below

Lemma 17. *Fix $\ell \in (0, 1]$. Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_t \in \mathbb{B}_2^d(1)$ be a sequence of vectors, and $\mu_1, \mu_2, \dots, \mu_t \in \mathbb{B}_2^d(1)$ be another sequence of vectors. Then we have*

$$\sum_{i=1}^t \frac{\text{clip}^2(\mathbf{x}_i \mu_i, \ell)}{\sum_{j=1}^{i-1} \text{clip}(\mathbf{x}_j \mu_i, \ell) \mathbf{x}_j^\top \mu_i + \ell^2} \leq O(d^4 \log^3(dt)). \quad (23)$$

Proof. Let

$$f_\ell(x) = \begin{cases} x^2, & |x| \leq \ell, \\ 2\ell x - \ell^2, & x > \ell, \\ -2\ell x - \ell^2, & x < -\ell \end{cases}$$

be a convex relaxation of the function $x \mapsto \text{clip}(x, \ell)x$. It is easy to see that $f_\ell(x)$ is convex in x and for any $x \in \mathbb{R}, \ell > 0$,

$$\text{clip}(x, \ell)x \leq f_\ell(x) \leq 2\text{clip}(x, \ell)x \leq 2x^2. \quad (24)$$

Let $h(x) = \frac{f_\ell(x)}{2}$. It is easy to see that if $x^2 \geq y^2$, $\frac{h(x)}{x^2} = \frac{\text{clip}(x, \ell)}{2x} \leq \frac{\text{clip}(y, \ell)}{2y} = \frac{h(y)}{y^2} \leq 1$. By Lemma 1 with $f(x) = h(x)$, we have that

$$\sum_{i=1}^t \frac{h(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^{i-1} h(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} \leq O(d^4 \log^3(dt)).$$

By (24), we obtain that

$$\begin{aligned} \sum_{i=1}^t \frac{\text{clip}^2(\mathbf{x}_i \boldsymbol{\mu}_i, \ell)}{\sum_{j=1}^{i-1} \text{clip}(\mathbf{x}_j \boldsymbol{\mu}_i, \ell) \mathbf{x}_j^\top \boldsymbol{\mu}_i + \ell^2} &\leq \sum_{i=1}^t \frac{\text{clip}(\mathbf{x}_i \boldsymbol{\mu}_i, \ell) \mathbf{x}_i \boldsymbol{\mu}_i}{\sum_{j=1}^{i-1} \text{clip}(\mathbf{x}_j \boldsymbol{\mu}_i, \ell) \mathbf{x}_j^\top \boldsymbol{\mu}_i + \ell^2} \\ &\leq 4 \sum_{i=1}^t \frac{h(\mathbf{x}_i \boldsymbol{\mu}_i)}{\sum_{j=1}^{i-1} h(\mathbf{x}_j \boldsymbol{\mu}_i) + \ell^2} \\ &\leq O(d^4 \log^3(dt)). \end{aligned}$$

The proof is completed. $\square$

D.2 Proof of Theorem 3

D.2.1 Optimism

The equation (1) accounts for the main novelty of our algorithm. We note that our confidence set is different from all previous ones [Dani et al., 2008, Abbasi-Yadkori et al., 2011]. Our confidence set is built based on the following new inequality, which may be of independent interest.

With Lemma 4 in hand, we can easily prove that the optimal $\boldsymbol{\theta}^*$ is always in our confidence set with high probability. The proof details can be found in Appendix D.3.

Lemma 18. *With probability at least $1 - O(\delta \log K)$, we have $\boldsymbol{\theta}^* \in \Theta_k$ for all $k \in [K]$.*

D.2.2 Bounding the Regret

We bound the regret under the event specified in Lemma 18. We have

$$\begin{aligned} \mathfrak{R}^K &= \sum_{k=1}^K (\max_{\mathbf{x} \in \mathcal{A}_k} \mathbf{x} \boldsymbol{\theta}^* - \mathbf{x}_k \boldsymbol{\theta}^*) \\ &\leq \sum_{k=1}^K \left(\max_{\mathbf{x} \in \mathcal{A}_k, \boldsymbol{\theta} \in \Theta_k} \mathbf{x} \boldsymbol{\theta} - \mathbf{x}_k \boldsymbol{\theta}^* \right) \leq \sum_{k=1}^K \mathbf{x}_k (\boldsymbol{\theta}_k - \boldsymbol{\theta}^*) = \sum_k \mathbf{x}_k \boldsymbol{\mu}_k, \end{aligned}$$

where second inequality follows from Lemma 18. Therefore, it suffices to bound $\sum_k \mathbf{x}_k \boldsymbol{\mu}_k$, for which we have the following lemma.

Lemma 19. *With probability $1 - O(\delta \log K)$, we have*

$$\sum_k \mathbf{x}_k \boldsymbol{\mu}_k \leq O \left(d^{4.5} (\log^4 dK) (\log \frac{dK}{\delta}) \left(\sqrt{d} + \sqrt{\sum_{k=1}^K \sigma_k^2} \right) \right).$$

Since this lemma is one of our main technical contribution, we provide more proof details.

Proof. First, we define the desired event $\mathcal{E} = \mathcal{E}_1 \cap \mathcal{E}_2$, where

$$\mathcal{E}_1 = \{\forall k \in [K] : \boldsymbol{\theta}^* \in \Theta_k\}, \quad \mathcal{E}_2 = \left\{ \sum_{k=1}^K \eta_k(\boldsymbol{\theta}^*) \leq \sum_{k=1}^K 8\sigma_k^2 + 4 \ln \frac{4}{\delta} \right\}.$$

By Lemma 18, we have $\Pr[\mathcal{E}_1] \geq 1 - O(\delta)$. By Lemma 10, we have $\Pr[\mathcal{E}_2] \geq 1 - O(\delta \log K)$. Therefore, by union bound, we have $\Pr[\mathcal{E}] \geq 1 - O(\delta \log K)$.

Now we bound $\sum_k \mathbf{x}_k \boldsymbol{\mu}_k$ under the event $\mathcal{E}$ to prove the lemma. Recall that

$$\Phi_k^j(\boldsymbol{\mu}) = \sum_{v=1}^{k-1} \text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}) \mathbf{x}_v \boldsymbol{\mu} + \ell_j^2, \quad \Psi_k^j(\boldsymbol{\mu}) = \sum_{v=1}^{k-1} \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}) \eta_v(\boldsymbol{\theta}^*). \quad (25)$$

Recall the definition of $\{\mathcal{K}_j\}_{j=1}^{L_2+1}$ in Section 4.2

To proceed, we need the following claim.

Claim 20. *We have*

$$\begin{aligned} \sum_k \mathbf{x}_k \boldsymbol{\mu}_k &= \sum_{k \in \mathcal{K}_{L_2+1}} \mathbf{x}_k \boldsymbol{\mu}_k + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k \\ &\leq 1 + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k \times \frac{3\sqrt{\Psi_k^j(\boldsymbol{\mu}_k)\ell} + \sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k)(\mathbf{x}_v \boldsymbol{\mu}_k)^2\ell} + 3\ell_j\ell}{\Phi_k^j(\boldsymbol{\mu}_k)}. \end{aligned} \quad (26)$$

We defer the proof of the claim to Appendix D.4 and continue to bound the three terms in (26). For the second term, we have

$$\begin{aligned} &\sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k \frac{\sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k)(\mathbf{x}_v \boldsymbol{\mu}_k)^2\ell}}{\Phi_k^j(\boldsymbol{\mu}_k)} \\ &\leq \frac{1}{2} \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k \mathbb{I} \left\{ \frac{\sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k)(\mathbf{x}_v \boldsymbol{\mu}_k)^2\ell}}{\Phi_k^j(\boldsymbol{\mu}_k)} > \frac{1}{2} \right\}. \end{aligned} \quad (27)$$

We note that

$$\begin{aligned} \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k \mathbb{I} \left\{ \frac{\sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k)(\mathbf{x}_v \boldsymbol{\mu}_k)^2\ell}}{\Phi_k^j(\boldsymbol{\mu}_k)} > \frac{1}{2} \right\} &\leq \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k \mathbb{I} \left\{ \Phi_k^j(\boldsymbol{\mu}_k) \leq 4\ell_j\ell \right\} \\ &\leq \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k \frac{4\ell_j\ell}{\Phi_k^j(\boldsymbol{\mu}_k)} \\ &\leq \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \frac{4\text{clip}_j^2(\mathbf{x}_k \boldsymbol{\mu}_k)\ell}{\Phi_k^j(\boldsymbol{\mu}_k)} \\ &\leq O(d^4 |\Lambda_2| \ell \log^3(dK)), \end{aligned} \quad (28)$$

where the last inequality uses Lemma 17. Collecting (26), (27) and (28), we have

$$\sum_k \mathbf{x}_k \boldsymbol{\mu}_k \leq 1 + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} 3\mathbf{x}_k \boldsymbol{\mu}_k \times \frac{\sqrt{\Psi_k^j(\boldsymbol{\mu}_k)\ell} + \ell_j\ell}{\Phi_k^j(\boldsymbol{\mu}_k)} + \frac{1}{2} \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \mathbf{x}_k \boldsymbol{\mu}_k + O(d^4 |\Lambda_2| \ell \log^3(dK)).$$

Solving $\sum_k \mathbf{x}_k \boldsymbol{\mu}_k$, we obtain

$$\begin{aligned} \sum_k \mathbf{x}_k \boldsymbol{\mu}_k &\leq O(d^4 |\Lambda_2| \ell \log^3(dK)) + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} 6\mathbf{x}_k \boldsymbol{\mu}_k \times \frac{\sqrt{\Psi_k^j(\boldsymbol{\mu}_k)\ell} + \ell_j\ell}{\Phi_k^j(\boldsymbol{\mu}_k)} \\ &\leq O(d^4 |\Lambda_2| \ell \log^3(dK)) + \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} 12\mathbf{x}_k \boldsymbol{\mu}_k \times \frac{\sqrt{\Psi_k^j(\boldsymbol{\mu}_k)\ell}}{\Phi_k^j(\boldsymbol{\mu}_k)}, \end{aligned} \quad (29)$$

where (29) uses the last two steps in (28). The remaining term in (29) can be bounded as

$$\sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} 12 \mathbf{x}_k \boldsymbol{\mu}_k \times \frac{\sqrt{\Psi_k^j(\boldsymbol{\mu}_k)_\ell}}{\Phi_k^j(\boldsymbol{\mu}_k)} \leq \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} 12 \mathbf{x}_k \boldsymbol{\mu}_k \ell_j \frac{\sqrt{\sum_{v=1}^{k-1} \eta_v(\boldsymbol{\theta}^*)_\ell}}{\Phi_k^j(\boldsymbol{\mu}_k)} \quad (30)$$

$$\leq \sum_{j=1}^{L_2} \sum_{k \in \mathcal{K}_j} \frac{12 \mathbf{x}_k \boldsymbol{\mu}_k \ell_j}{\Phi_k^j(\boldsymbol{\mu}_k)} \sqrt{\sum_{k=1}^K \eta_k(\boldsymbol{\theta}^*)_\ell} \leq O(d^4 |\Lambda_2| \log^3(dK)) \times \sqrt{\sum_{k=1}^K \eta_k(\boldsymbol{\theta}^*)_\ell} \quad (31)$$

$$\leq O(d^4 |\Lambda_2| \log^3(dK)) \times \sqrt{\left(\ln \frac{1}{\delta} + \sum_{k=1}^K \sigma_k^2 \right) \ell}, \quad (32)$$

where (30) uses the definition of $\Psi_k^j(\cdot)$, (31) again uses the last two steps in (28), and (32) uses the event $\mathcal{E}_2$. $\square$

Now we can finish the proof of Theorem 3. We choose $\delta = O((K \log K)^{-1})$. Since on the event $\mathcal{E}^C$, we have $\mathfrak{R}^K \leq K$. Therefore, together with the bound on $\mathcal{E}$ from Lemma 19, we conclude that the expected regret is bounded by $\mathbb{E}[\mathfrak{R}^K] \leq \tilde{O}(d^{4.5} \sqrt{\sum_{k=1}^K \sigma_k^2} + d^5)$.

Proof. It suffices to prove the theorem for $b = 1$, because otherwise we can apply $\{X_i/b\}_{i=1}^n$ to the $b = 1$ case. By Lemma 11 with $\epsilon = 1$ and $\delta < 1/e$, we have

$$\Pr \left[\left| \sum_{i=1}^n X_i \right| \geq 2 \sqrt{\sum_{i=1}^n 2 \mathbb{E}[X_i^2 \mid \mathcal{F}_{i-1}] \ln \frac{1}{\delta} + 4 \ln \frac{1}{\delta}} \right] \leq 4\delta \log_2 n. \quad (33)$$

By Lemma 9, we have

$$\Pr \left[\sum_{i=1}^n \mathbb{E}[X_i^2 \mid \mathcal{F}_{i-1}] \geq \sum_{i=1}^n 8X_i^2 + 4 \ln \frac{4}{\delta} \right] \leq ([\log_2 n] + 1)\delta. \quad (34)$$

Therefore, by a union bound over (33) and (34), we have with probability at least $1 - 6\delta \log_2 n$,

$$\begin{aligned} \left| \sum_{i=1}^n X_i \right| &\leq \sqrt{\sum_{i=1}^n 8 \mathbb{E}[X_i^2 \mid \mathcal{F}_{i-1}] \ln \frac{1}{\delta} + 4 \ln \frac{1}{\delta}} \\ &\leq \sqrt{8 \left(\sum_{i=1}^n 8X_i^2 + 4 \ln \frac{4}{\delta} \right) \ln \frac{1}{\delta} + 4 \ln \frac{1}{\delta}} \leq 8 \sqrt{\sum_{i=1}^n X_i^2 \ln \frac{1}{\delta} + 16 \ln \frac{1}{\delta}}, \end{aligned}$$

which concludes the proof. $\square$

D.3 Proof of Lemma 18

Proof. Let $\delta' = e^{-\ell}$. We define the desired event $\mathcal{E} = \bigcap_{k \in [K], j \in \Lambda_2} \mathcal{E}_k^j$, where

$$\mathcal{E}_k^j = \left\{ \left| \sum_{v=1}^k \text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}) \epsilon_v(\boldsymbol{\theta}^*) \right| \leq \sqrt{\sum_{v=1}^k \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}) \eta_v(\boldsymbol{\theta}^*)_\ell} + \ell_j \ell, \forall \boldsymbol{\mu} \in \mathcal{B} \right\}.$$

Note that for each v , we have that $|\text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}) \epsilon_v(\boldsymbol{\theta}^*)| \leq \ell_j$ and that $(\text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}) \epsilon_v(\boldsymbol{\theta}^*))^2 = \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}) \eta_v(\boldsymbol{\theta}^*)$, so by Theorem 4, we have

$$\begin{aligned} \Pr \left[\left| \sum_{v=1}^k \text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}) \epsilon_v \right| \leq \sqrt{\sum_{v=1}^k \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}) \text{Var}(\epsilon_v \mid \mathcal{F}_v) \ell + \ell_j \ell} \right] \\ \geq 1 - O \left(e^{-\frac{\ell}{\log_2 \log_2 K}} \right) \\ \geq 1 - O \left(\frac{\delta}{K |\mathcal{B}| |\Lambda_2|} \log K \right), \end{aligned}$$

where $\mathcal{F}_v$ is as defined in Section 3. Finally, using a union bound over $(\boldsymbol{\mu}, j, k) \in \mathcal{B} \times \Lambda_2 \times [K]$, we have $\Pr[\mathcal{E}] \geq 1 - O(\delta \log K)$. $\square$

D.4 Proof of Claim 20

Proof. We elaborate on (26). We will prove it by showing that the numerator is always greater than the denominator in the fraction in (26), so each term $\mathbf{x}_k \boldsymbol{\mu}_k$ is multiplied by a number greater than 1. We have for every $j \in \Lambda_2$,

$$\begin{aligned} \Phi_k^j(\boldsymbol{\mu}_k) &= \sum_{v=1}^{k-1} \text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}_k) \mathbf{x}_v \boldsymbol{\mu}_k + \ell_j^2 \\ &\leq \left| \sum_{v=1}^{k-1} \text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}_k) \epsilon_v(\boldsymbol{\theta}^*) \right| + \left| \sum_{v=1}^{k-1} \text{clip}_j(\mathbf{x}_v \boldsymbol{\mu}_k) \epsilon_v(\boldsymbol{\theta}_k) \right| + \ell_j^2 \end{aligned} \quad (35)$$

$$\leq \sqrt{\Psi_k^j(\boldsymbol{\mu}_k) \ell} + \sqrt{\sum_{v=1}^{k-1} \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k) \eta_v(\boldsymbol{\theta}_k) \ell} + \ell_j^2 + \frac{1}{HK} + 2\ell_j \ell \quad (36)$$

$$\begin{aligned} &\leq \sqrt{\Psi_k^j(\boldsymbol{\mu}_k) \ell} + \sqrt{\sum_{v=1}^{k-1} \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k) \eta_v(\boldsymbol{\theta}^*) \ell} + \sqrt{\sum_{v=1}^{k-1} \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k) |\eta_v(\boldsymbol{\theta}_k) - \eta_v(\boldsymbol{\theta}^*)| \ell} + 3\ell_j \ell \\ &= 2\sqrt{\Psi_k^j(\boldsymbol{\mu}_k) \ell} + \sqrt{\sum_{v=1}^{k-1} \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k) |\eta_v(\boldsymbol{\theta}_k) - \eta_v(\boldsymbol{\theta}^*)| \ell} + 3\ell_j \ell \end{aligned}$$

$$\leq 2\sqrt{\Psi_k^j(\boldsymbol{\mu}_k) \ell} + \sqrt{\sum_{v=1}^{k-1} \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k) (\eta_v(\boldsymbol{\theta}^*) + 2(\mathbf{x}_v \boldsymbol{\mu}_k)^2) \ell} + 3\ell_j \ell \quad (37)$$

$$\begin{aligned} &\leq 2\sqrt{\Psi_k^j(\boldsymbol{\mu}_k) \ell} + \sqrt{\sum_{v=1}^{k-1} \text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k) \eta_v(\boldsymbol{\theta}^*) \ell} + \sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k) (\mathbf{x}_v \boldsymbol{\mu}_k)^2 \ell} + 3\ell_j \ell, \\ &= 3\sqrt{\Psi_k^j(\boldsymbol{\mu}_k) \ell} + \sqrt{\sum_{v=1}^{k-1} 2\text{clip}_j^2(\mathbf{x}_v \boldsymbol{\mu}_k) (\mathbf{x}_v \boldsymbol{\mu}_k)^2 \ell} + 3\ell_j \ell, \end{aligned} \quad (38)$$

where (35) uses $\epsilon_v(\boldsymbol{\theta}_k) - \epsilon_v(\boldsymbol{\theta}^*) = \mathbf{x}_v(\boldsymbol{\theta}_k - \boldsymbol{\theta}^*) = \mathbf{x}_v \boldsymbol{\mu}_k$, (36) uses that $\boldsymbol{\theta}^*, \boldsymbol{\theta}_k \in \Theta_k$, the definition of Θ_k in (1) and the fact that $\mathcal{B}$ is an K^{-3} -net of $\mathbb{B}_2^d(2)$, and (37) uses

$$\begin{aligned} |\eta_v(\boldsymbol{\theta}_k) - \eta_v(\boldsymbol{\theta}^*)| &= |(\epsilon_v(\boldsymbol{\theta}^*) - \mathbf{x}_v \boldsymbol{\mu}_k)^2 - (\epsilon_v(\boldsymbol{\theta}^*))^2| \\ &\leq 2|\epsilon_v(\boldsymbol{\theta}^*)| \mathbf{x}_v \boldsymbol{\mu}_k + (\epsilon_v(\boldsymbol{\theta}^*))^2 \leq (\mathbf{x}_v \boldsymbol{\mu}_k)^2 + 2(\epsilon_v(\boldsymbol{\theta}^*))^2. \end{aligned}$$

Since (38) holds for every $j \in \Lambda_2$, it holds for $j = j_k$, and thus (26) follows. $\square$

E Missing Proofs in Section 5

E.1 Proof of Theorem 5

Before introducing our proof, we make some definitions. We let $\theta_{k,h}^m = \arg \max_{\theta \in \Theta_k} |\mathbf{x}_{k,h}^m(\theta - \theta^*)|$ and $\mu_{k,h}^m = \theta_{k,h}^m - \theta^*$. Recall that $\mathcal{T}_k^{m,i}$ is defined in Algorithm 2. We define

$$\Phi_k^{m,i,j}(\mu) = \sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j(\mathbf{x}_{v,u}^m \mu) \mathbf{x}_{v,u}^m \mu + \ell_j^2, \quad \Psi_k^{m,i,j}(\mu) = \sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j^2(\mathbf{x}_{v,u}^m \mu) \eta_{v,u}^m. \quad (39)$$

Note that our definitions in (39) are similar to those for linear bandits in (25). The main differences are: 1) we define $\Phi(\cdot), \Psi(\cdot)$ also for higher moments, as indicated by the index m in their superscripts; 2) we add the variance layer, so that we only use samples from $\mathcal{T}_k^{m,i}$; 3) since we can now estimate variance, we use the upper bound of estimated variance in lieu of the empirical variance. For $h \in [H + 1]$, we further define

$$I_h^k = \mathbb{I}\{\forall u \leq h, m, i, j : \Phi_{k,u}^{m,i,j}(\mu_{k,u}^m) \leq 4(d+2)^2 \Phi_k^{m,i,j}(\mu_{k,u}^m)\}, \quad (40)$$

where $I_h^k = 1$ indicates that for every $u \leq h$, the confidence set using data prior to the time step (k, u) can be properly approximated by the confidence set with data prior to the episode k . We define I_h^k in this way to ensure that it is $\mathcal{F}_h^k$ -measurable. The following lemma ensures that Q_h^k is optimistic with high probability. Its proof is deferred to Appendix E.3.

Lemma 21. $\Pr[\forall k, h, s, a : Q_h^k(s, a) \geq Q_h^*(s, a)] \geq \Pr[\forall k \in [K] : \theta^* \in \Theta_k] \geq 1 - O(\delta)$.

When the event specified in Lemma 21 holds, the regret can be decomposed as

$$\mathfrak{R}^K = \sum_{k=1}^K (V_1^*(s_1^k) - V_1^{\pi_k}(s_1^k)) \leq \sum_{k=1}^K (V_1^k(s_1^k) - V_1^{\pi_k}(s_1^k)) \leq \check{\mathfrak{R}}_1 + \check{\mathfrak{R}}_2 + \mathfrak{R}_3 + \sum_{k,h} (I_h^k - I_{h+1}^k),$$

where

$$\begin{aligned} \check{\mathfrak{R}}_1 &= \sum_{k,h} (P_{s_h^k, a_h^k} V_{h+1}^k - V_{h+1}^k(s_{h+1}^k)) I_h^k, & \check{\mathfrak{R}}_2 &= \sum_{k,h} (V_h^k(s_h^k) - r_h^k - P_{s_h^k, a_h^k} V_{h+1}^k) I_h^k, \\ \mathfrak{R}_3 &= \sum_{k=1}^K \left(\sum_{h=1}^H r_h^k - V_1^{\pi_k}(s_1^k) \right). \end{aligned}$$

Next we analyze these terms. First, we observe that $\mathfrak{R}_3$ is a sum of a martingale difference sequence, so by Lemma 6, we have $\mathfrak{R}_3 \leq O(\sqrt{K \log(1/\delta)})$ with probability at least $1 - \delta$. Next, we use the following lemma to bound $\sum_{k,h} (I_h^k - I_{h+1}^k)$. We defer its proof to Appendix E.4.

Lemma 22. $\sum_{k,h} (I_h^k - I_{h+1}^k) \leq O(d \log^5(dHK))$.

To bound $\check{\mathfrak{R}}_1$ and $\check{\mathfrak{R}}_2$, we need to define the following quantities. First, we denote $\check{\mathbf{x}}_{k,h} = \mathbf{x}_{k,h} I_h^k$ and $\check{\eta}_{k,h}^m = \eta_{k,h}^m I_h^k$. Next, for $m \in \Lambda_0$, we define

$$\check{R}_m = \sum_{k,h} |\check{\mathbf{x}}_{k,h}^m \mu_{k,h}^m|, \quad \check{M}_m = \sum_{k,h} \left(P_{s_h^k, a_h^k} (V_{h+1}^k)^{2^m} - (V_{h+1}^k(s_{h+1}^k))^{2^m} \right) I_h^k.$$

Intuitively, $\check{R}_m$ represents the “regret” of 2^m -th moment prediction and $\check{M}_m$ represents the total variance of 2^m -th order value function. We have $\check{\mathfrak{R}}_1 = \check{M}_0$ by definition and using that

$$Q_h^k(s, a) - r(s, a) - P_{s,a} V_{h+1}^k \leq \max_{\theta \in \Theta_k} \mathbf{x}_{k,h}^0(\theta - \theta^*),$$

we have $\check{\mathfrak{R}}_2 \leq \check{R}_0$. So it suffices to bound $\check{R}_0 + \check{M}_0$, which is done by the following lemma.

Lemma 23. *With probability at least $1 - \delta$, we have*

$$\check{R}_0 + |\check{M}_0| \leq O\left(d^{4.5} \sqrt{K \log^5(dHK) \log(1/\delta)} + d^9 \log^6(dHK) \log(1/\delta)\right).$$

Lemma 23 is the main technical part of our result in Section 5, so we sketch its proof in the next subsection. With the lemma in hand, we have with probability $1 - \delta$ that $\mathfrak{R}^K \leq \tilde{O}(d^{4.5} \sqrt{K} + d^9)$. Finally, We conclude the proof to Theorem 5 by choosing $\delta = 1/K$ and noting that $\mathfrak{R}^K \leq K$.

E.2 Bounding $\check{R}$ and $\check{M}$

We sketch the proof for Lemma 23. The first step to bound $\check{R}_m$ is to relate it to the variance $\check{\eta}^m$.

Lemma 24. *With probability at least $1 - \delta$, we have $\check{R}_m \leq O(d^4 \sqrt{\sum_{k,h} \check{\eta}_{k,h}^m} \iota \log^7(dHK)) + d^6 \iota \log^5(dHK)$.*

We defer the proof to Appendix E.5. The proof is spiritually similar to proof of Lemma 19. The main difference is that we use the peeling technique to the magnitude of the variance.

Based on Lemma 24, we use the following recursion lemma to relate $\check{R}_m, \check{M}_m$ to $\check{R}_{m+1}, \check{M}_{m+1}$. We defer the proof to Appendix E.6. It mainly uses similar ideas in Zhang et al. [2020a].

Lemma 25 (Recursions). *With probability at least $1 - \delta$, we have*

$$\begin{aligned} \check{R}_m &\leq O\left(d^4 \sqrt{(\check{M}_{m+1} + 2^{m+1}(K + \check{R}_0) + \check{R}_{m+1} + \check{R}_m) \iota \log^7(dHK)} + d^6 \iota \log^5(dHK)\right), \\ |\check{M}_m| &\leq O\left(\sqrt{(\check{M}_{m+1} + O(d \log^5(dHK)) + 2^{m+1}(K + \check{R}_0)) \log(1/\delta)} + \log(1/\delta)\right). \end{aligned}$$

Finally, we can prove Lemma 23 by collecting Lemma 24, 25 and using a technical lemma about recursion (Lemma 12). The details are in Appendix E.7.

E.3 Proof of Lemma 21

Proof. The lemma consists of two inequalities. The first inequality is proved using backward induction, where the induction step is given as

$$\begin{aligned} Q_h^k(s, a) &= \min\{1, r(s, a) + \max_{\theta \in \Theta_k} \sum_{i=1}^d \theta_i P_{s,a}^i V_{h+1}^k\} \\ &\geq \min\{1, r(s, a) + \sum_{i=1}^d \theta_i^* P_{s,a}^i V_{h+1}^k\} \geq \min\{1, r(s, a) + \sum_{i=1}^d \theta_i^* P_{s,a}^i V_{h+1}^*\} = Q_h^*(s, a), \\ V_h^k(s) &= \max_a Q_h^k(s, a) \geq \max_a Q_h^*(s, a) = V_h^*(s). \end{aligned}$$

We now prove the second inequality. Let $\delta' = e^{-\iota}$. We define the desired event $\mathcal{E} = \bigcap_{k,m,i,j} \mathcal{E}_k^{m,i,j}$, where

$$\mathcal{E}_k^{m,i,j} = \left\{ \left| \sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j(\mathbf{x}_{v,u}^m \boldsymbol{\mu}) \varepsilon_{\kappa,h}^m \right| \leq 4 \sqrt{\sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j^2(\mathbf{x}_{v,u}^m \boldsymbol{\mu}) \text{Var}(\varepsilon_{v,u}^m \mid \mathcal{F}_u^v) \ln \frac{1}{\delta'}} + 4\ell_j \ln \frac{1}{\delta'}, \forall \boldsymbol{\mu} \in \mathcal{B} \right\}.$$

Note that for a fixed k , we have that $|\text{clip}_j(\mathbf{x}_{v,u}^m \boldsymbol{\mu}) \varepsilon_{v,u}^m| \leq \ell_j \leq 1$ and that

$$\text{Var}\left(\text{clip}_j(\mathbf{x}_{v,u}^m \boldsymbol{\mu}) \varepsilon_{v,u}^m \mathbb{I}\{(v, u) \in \mathcal{T}_k^{m,i}\} \mid \mathcal{F}_u^v\right) = \text{clip}_j^2(\mathbf{x}_{k,h}^m \boldsymbol{\mu})^2 \mathbb{I}\{(v, u) \in \mathcal{T}_k^{m,i}\} \text{Var}(\varepsilon_{v,u}^m \mid \mathcal{F}_u^v),$$

so by Lemma 11 with $b = \ell_j, \epsilon = 1$, we have

$$\begin{aligned} \Pr\left[\left| \sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j(\mathbf{x}_{v,u}^m \boldsymbol{\mu}) \varepsilon_{v,u}^m \right| \geq 4 \sqrt{\sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j^2(\mathbf{x}_{v,u}^m \boldsymbol{\mu}) \text{Var}(\varepsilon_{v,u}^m \mid \mathcal{F}_u^v) \ln \frac{1}{\delta'}} + 4\ell_j \ln \frac{1}{\delta'}\right] \\ \leq 4\delta' \log_2(HK). \end{aligned}$$

Using a union bound over $(\boldsymbol{\mu}, m, i, j, k) \in \mathcal{B} \times \Lambda_0 \times \Lambda_1 \times \Lambda_2 \times [K]$, we have $\Pr[\mathcal{E}] \geq 1 - O(\delta' K |\mathcal{B}| \log^4(HK)) \geq 1 - O(\delta)$.

Next we show that the event $\mathcal{E}$ implies that $\boldsymbol{\theta}^* \in \Theta_k$ for every $k \in [K]$. We show by induction over k . For $k = 1$ it is clear. For $k \geq 1$, since $\boldsymbol{\theta}^* \in \Theta_k$, for every $h \in [H]$, we have $\eta_{k,h}^m = \max_{\boldsymbol{\theta} \in \Theta_k} \{\boldsymbol{\theta} \mathbf{x}_{k,h}^{m+1} - (\boldsymbol{\theta} \mathbf{x}_{k,h}^m)^2\} \geq \boldsymbol{\theta}^* \mathbf{x}_{k,h}^{m+1} - (\boldsymbol{\theta}^* \mathbf{x}_{k,h}^m)^2 \geq \text{Var}(\varepsilon_{k,h}^m \mid \mathcal{F}_h^k)$, which, together with the event $\bigcap_{m,i,j} \mathcal{E}_{k+1}^{m,i,j}$, implies that $\boldsymbol{\theta}^* \in \Theta_{k+1}$. $\square$

E.4 Proof of Lemma 22

Proof. We define

$$I_{k,h}^{m,i,j} = \mathbb{I}\{\forall u \leq h : \Phi_{k,u}^{m,i,j}(\boldsymbol{\mu}_{k,u}^m) \leq 4(d+2)^2 \Phi_k^{m,i,j}(\boldsymbol{\mu}_{k,u}^m)\}.$$

Then we have $I_h^k = \prod_{m,i,j} I_{k,h}^{m,i,j}$. Also we have

$$\sum_h (I_h^k - I_{h+1}^k) \leq \sum_{m,i,j} \sum_h (I_{k,h}^{m,i,j} - I_{k,h+1}^{m,i,j}).$$

Note that $I_h^k \geq I_{h+1}^k$ and $I_{k,h}^{m,i,j} \geq I_{k,h+1}^{m,i,j}$. For each fixed m, i, j , if $\sum_h (I_{k,h}^{m,i,j} - I_{k,h+1}^{m,i,j}) = 1$, then there exists $h \in [H]$, such that for the time step (k, h) , we have $\Phi_{k,h}^{m,i,j}(\boldsymbol{\mu}) > 4(d+2)^2 \Phi_k^{m,i,j}(\boldsymbol{\mu})$ for some $\boldsymbol{\mu}$. By Lemma 14 with $f(x) = \text{clip}(x, \ell_j)x$ and $\ell = \ell_j$, there are at most $O(d \log^2(dHK))$ such time steps. We conclude by noting that we have $|\Lambda_0 \times \Lambda_1 \times \Lambda_2| \leq O(\log^3(dHK))$ possible m, i, j pairs. $\square$

E.5 Proof of Lemma 24

To prove this lemma, we define the index sets to help us apply the peeling technique. We denote

$$\begin{aligned} \mathcal{T}_k^{m,i,j} &= \{(v, u) \in \mathcal{T}_k^{m,i} : |\mathbf{x}_{v,u}^m \boldsymbol{\mu}_{v,u}^m| \in (\ell_{j+1}, \ell_j]\}, \\ \mathcal{T}_k^{m,i,L_2+1} &= \{(v, u) \in \mathcal{T}_k^{m,i} : |\mathbf{x}_{v,u}^m \boldsymbol{\mu}_{v,u}^m| \in [0, \ell_{L_2+1}]\}, \end{aligned}$$

and $\tilde{\mathcal{T}}_k^{m,i,j} = \{(v, u) \in \mathcal{T}_k^{m,i,j} : I_u^v = 1\}$. We also denote $\mathcal{T}^{m,i,j} = \mathcal{T}_{K+1}^{m,i,j}$, $\tilde{\mathcal{T}}^{m,i,j} = \tilde{\mathcal{T}}_{K+1}^{m,i,j}$.

Proof. Since $\boldsymbol{\theta}_{k,h}^m \in \Theta_k \subseteq \Theta_k^{m,i,j}$, choosing $\boldsymbol{\mu} = \boldsymbol{\mu}_{k,h}^m$ in the confidence set definition and using that $\mathbf{x}_{v,u}^m \boldsymbol{\mu}_{k,h}^m = \epsilon_{v,u}^m(\boldsymbol{\theta}^*) - \epsilon_{v,u}^m(\boldsymbol{\theta}_{k,h}^m)$, we have

$$\begin{aligned} \Phi_k^{m,i,j}(\boldsymbol{\mu}_{k,h}^m) &= \sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j(\mathbf{x}_{v,u}^m \boldsymbol{\mu}_{k,h}^m) \mathbf{x}_{v,u}^m \boldsymbol{\mu}_{k,h}^m + \ell_j^2 \\ &\leq \left| \sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j(\mathbf{x}_{v,u}^m \boldsymbol{\mu}_{k,h}^m) \epsilon_{v,u}^m(\boldsymbol{\theta}^*) \right| + \left| \sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j(\mathbf{x}_{v,u}^m \boldsymbol{\mu}_{k,h}^m) \epsilon_{v,u}^m(\boldsymbol{\theta}_{k,h}^m) \right| + \ell_j^2 \\ &\leq 8 \sqrt{\sum_{(v,u) \in \mathcal{T}_k^{m,i}} \text{clip}_j(\mathbf{x}_{v,u}^m \boldsymbol{\mu}_{k,h}^m) \eta_{v,u}^m \ell} + 8 \ell_j \ell + \ell_j^2 \\ &\leq 8 \sqrt{\Psi_k^{m,i,j}(\boldsymbol{\mu}_{k,h}^m) \ell} + 16 \ell_j \ell. \end{aligned} \tag{41}$$

Therefore, when $I_k^h = 0$, we have

$$\frac{\Phi_{k,h}^{m,i,j}(\boldsymbol{\mu}_{k,h}^m)}{4(d+2)^2} \leq \Phi_k^{m,i,j}(\boldsymbol{\mu}_{k,h}^m) \leq 16(\sqrt{\Psi_k^{m,i,j}(\boldsymbol{\mu}_{k,h}^m) \ell} + \ell_j \ell).$$

Next we analyze the sum. Using the fact that

$$\frac{64(d+2)^2 \left(\sqrt{\Psi_{k,h}^{m,i,j}(\boldsymbol{\mu}_{k,h}^m) \ell} + \ell_j \ell \right)}{\Phi_{k,h}^{m,i,j}(\boldsymbol{\mu}_{k,h}^m)} \geq 1,$$

we obtain

$$\sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} |\mathbf{x}_{k,h}^m \boldsymbol{\mu}_{k,h}^m| \leq \sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} |\mathbf{x}_{k,h}^m \boldsymbol{\mu}_{k,h}^m| \frac{64(d+2)^2 \left(\sqrt{\Psi_{k,h}^{m,i,j}(\boldsymbol{\mu}_{k,h}^m) \ell} + \ell_j \ell \right)}{\Phi_{k,h}^{m,i,j}(\boldsymbol{\mu}_{k,h}^m)} \tag{42}$$

$$\leq 64(d+2)^2 \sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} \left(\frac{|\mathbf{x}_{k,h}^m \boldsymbol{\mu}_{k,h}^m| \sqrt{\ell_i \ell}}{\sqrt{\Phi_{k,h}^{m,i,j}(\boldsymbol{\mu}_{k,h}^m)}} + \frac{|\mathbf{x}_{k,h}^m \boldsymbol{\mu}_{k,h}^m| \ell_j \ell}{\Phi_{k,h}^{m,i,j}(\boldsymbol{\mu}_{k,h}^m)} \right), \tag{43}$$

where the last inequality uses that for every μ , we have

$$\Psi_{k,h}^{m,i,j}(\mu) = \sum_{(v,u) \in \mathcal{T}_{k,h}^{m,i}} \text{clip}_j^2(x_{v,u}^m \mu) \eta_{v,u}^m \leq \ell_i \sum_{(v,u) \in \mathcal{T}_{k,h}^{m,i}} \text{clip}_j(x_{k,h}^m \mu) x_{k,h}^m \mu \leq \ell_i \Phi_{k,h}^{m,i,j}(\mu). \quad (44)$$

In (44), the first inequality uses that $\eta_{v,u}^m \leq \ell_i$ for $(v,u) \in \mathcal{T}_{k,h}^{m,i}$ and that $\text{clip}_j^2(\alpha) \leq \text{clip}_j(\alpha)\alpha$ for $\alpha \in \mathbb{R}$, and the second inequality uses the definition of $\Phi_{k,h}^{m,i,j}(\mu)$. Next we bound the two terms in (43). To bound the first term, we note that

$$\sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} \frac{|x_{k,h}^m \mu_{k,h}^m|}{\sqrt{\Phi_{k,h}^{m,i,j}(\mu_{k,h}^m)}} \leq \sqrt{|\tilde{\mathcal{T}}^{m,i,j}|} \sqrt{\sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} \frac{(x_{k,h}^m \mu_{k,h}^m)^2}{\Phi_{k,h}^{m,i,j}(\mu_{k,h}^m)}} \quad (45)$$

$$\leq \sqrt{|\tilde{\mathcal{T}}^{m,i,j}|} \sqrt{\sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} \frac{\text{clip}_j^2(x_{k,h}^m \mu_{k,h}^m)}{\Phi_{k,h}^{m,i,j}(\mu_{k,h}^m)}} \quad (46)$$

$$\leq \sqrt{|\tilde{\mathcal{T}}^{m,i,j}|} \sqrt{\sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} \frac{\text{clip}_j^2(x_{k,h}^m \mu_{k,h}^m)}{\sum_{(v,u) \in \tilde{\mathcal{T}}_{k,h}^{m,i,j}} \text{clip}_j(x_{v,u}^m \mu_{k,h}^m) x_{v,u}^m \mu_{k,h}^m + \ell_j^2}} \quad (47)$$

$$\leq \sqrt{|\tilde{\mathcal{T}}^{m,i,j}|} \times O(\sqrt{d^4 \log^3(dHK)}), \quad (48)$$

where (45) uses Cauchy's inequality, (46) uses that $|x_{k,h}^m \mu_{k,h}^m| \leq \ell_j$ for $(k,h) \in \mathcal{T}^{m,i,j}$, (47) uses the definition of $\Phi_{k,h}^{m,i,j}(\mu)$, and (48) uses Lemma 17. To bound the second term in (43), we have

$$\sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} \frac{|x_{k,h}^m \mu_{k,h}^m| \ell_j}{\Phi_{k,h}^{m,i,j}(\mu_{k,h}^m)} \leq \sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} \frac{2 \text{clip}_j^2(x_{k,h}^m \mu_{k,h}^m)}{\Phi_{k,h}^{m,i,j}(\mu_{k,h}^m)} \leq O(d^4 \log^3(dHK)), \quad (49)$$

where the first inequality uses that $|x_{k,h}^m \mu_{k,h}^m| \geq \ell_j/2$ for $(k,h) \in \mathcal{T}^{m,i,j}$ and the second inequality is the same as what we have shown from (46) to (48). As a result, combining (43), (48) and (49), we have

$$\sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} |x_{k,h}^m \mu_{k,h}^m| \leq 64(d+2)^2 \times O\left(\sqrt{d^4 \ell_i |\tilde{\mathcal{T}}^{m,i,j}| \log^3(dHK)} + d^4 \ell \log^3(dHK)\right) \quad (50)$$

$$\leq O\left(d^4 \sqrt{\ell_i |\tilde{\mathcal{T}}^{m,i,j}| \log^3(dHK)} + d^6 \ell \log^3(dHK)\right). \quad (51)$$

Recall that (51) requires $x_{k,h}^m \mu_{k,h}^m \in [\ell_j/2, \ell_j]$, which would be false for $j = L_2 + 1$. In this corner case, $j = L_2 + 1$, we have

$$\sum_i \sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} |x_{k,h}^m \mu_{k,h}^m| \leq KH \ell_j \leq O(1). \quad (52)$$

Finally, combining (51) and (52), we have

$$\begin{aligned} \sum_{k,h} |x_{k,h}^m \mu_{k,h}^m| &= \sum_{i,j} \sum_{(k,h) \in \tilde{\mathcal{T}}^{m,i,j}} |x_{k,h}^m \mu_{k,h}^m| \\ &\leq O(1) + \sum_{i,j} O\left(d^4 \sqrt{\ell_i |\tilde{\mathcal{T}}^{m,i,j}| \log^3(dHK)} + L_2 d^6 \ell \log^3(dHK)\right) \\ &\leq O\left(d^4 \sqrt{\sum_{k,h} \tilde{\eta}_{k,h}^m \ell \log^7(dHK)} + d^6 \ell \log^5(dHK)\right), \end{aligned} \quad (53)$$

where (53) uses that $\ell_i |\tilde{\mathcal{T}}^{m,i,j}| \leq O(1 + \sum_{k,h} \tilde{\eta}_{k,h}^m)$, which can be proved as follows: for $i \leq L_1$, it is due to $\eta_{k,h}^m \geq \ell_i/2$; for $i = L_1 + 1$, it is due to $1/\ell_i \geq KH \geq |\tilde{\mathcal{T}}^{m,i,j}|$. $\square$

E.6 Proof of Lemma 25

Proof. Define

$$\check{\zeta}_{k,h}^m = (P_{s_h^k, a_h^k}(V_{h+1}^k)^{2^{m+1}} - (P_{s_h^k, a_h^k}(V_{h+1}^k)^{2^m})^2) I_h^k.$$

We note that $\check{M}_m$ is a martingale, so by Lemma 11 with a union bound over m , we have

$$\Pr \left[\forall m \in \Lambda_0 : |\check{M}_m| \leq 2 \sqrt{2 \sum_{k,h} \check{\zeta}_{k,h}^m \ln \frac{1}{\delta}} + 4 \ln \frac{1}{\delta} \right] \geq 1 - O(\delta \log^2(dKH)). \quad (54)$$

By the definition of $\check{\eta}_{k,h}^m$, we have

$$\sum_{k,h} \check{\eta}_{k,h}^m \leq \sum_{k,h} \left(\check{\zeta}_{k,h}^m + \max_{\theta \in \Theta_k} \check{x}_{k,h}^{m+1}(\theta - \theta^*) + 2 \max_{\theta \in \Theta_k} \check{x}_{k,h}^m(\theta^* - \theta) \right) \quad (55)$$

$$\leq \sum_{k,h} \check{\zeta}_{k,h}^m + \check{R}_{m+1} + 2\check{R}_m, \quad (56)$$

We have that

$$\begin{aligned} \sum_{k,h} \check{\zeta}_{k,h}^m &= \sum_{k,h} \left(P_{s_h^k, a_h^k}(V_{h+1}^k)^{2^{m+1}} - (P_{s_h^k, a_h^k}(V_{h+1}^k)^{2^m})^2 \right) I_h^k \\ &\leq \sum_{k,h} \left(P_{s_h^k, a_h^k}(V_{h+1}^k)^{2^{m+1}} - (V_{h+1}^k(s_{h+1}^k))^{2^{m+1}} \right) I_h^k + \sum_{k,h} (V_h^k(s_h^k))^{2^{m+1}} (I_h^k - I_{h+1}^k) \\ &\quad + \sum_{k,h} \left((V_h^k(s_h^k))^{2^{m+1}} - (P_{s_h^k, a_h^k}(V_{h+1}^k)^{2^m})^2 \right) I_h^k \\ &\leq \check{M}_{m+1} + O(d \log^5(dHK)) + \sum_{k,h} \left((V_h^k(s_h^k))^{2^{m+1}} - (P_{s_h^k, a_h^k}(V_{h+1}^k)^{2^m})^2 \right) I_h^k \\ &\leq \check{M}_{m+1} + O(d \log^5(dHK)) + \sum_{k,h} \left((V_h^k(s_h^k))^{2^{m+1}} - (P_{s_h^k, a_h^k} V_{h+1}^k)^{2^{m+1}} \right) \\ &\leq \check{M}_{m+1} + O(d \log^5(dHK)) + 2^{m+1} \sum_{k,h} I_h^k \cdot \max\{V_h^k(s_h^k) - P_{s_h^k, a_h^k} V_{h+1}^k, 0\} \\ &\leq \check{M}_{m+1} + O(d \log^5(dHK)) + 2^{m+1} \sum_{k,h} I_h^k \left(r(s_h^k, a_h^k) + \max_{\theta \in \Theta_k} \mathbf{x}_{k,h}^0(\theta - \theta^*) \right) \\ &\leq \check{M}_{m+1} + O(d \log^5(dHK)) + 2^{m+1}(K + \check{R}_0). \end{aligned} \quad (57)$$

Finally, by (56), (57) and Lemma 24, we have

$$\begin{aligned} \check{R}_m &\leq O \left(d^4 \sqrt{(\check{M}_{m+1} + O(d \log^5(dHK)) + 2^{m+1}(K + \check{R}_0) + \check{R}_{m+1} + 2\check{R}_m) \iota \log^7(dHK)} + d^6 \iota \log^5(dHK) \right) \\ &\leq O \left(d^4 \sqrt{(\check{M}_{m+1} + 2^{m+1}(K + \check{R}_0) + \check{R}_{m+1} + \check{R}_m) \iota \log^7(dHK)} + d^6 \iota \log^5(dHK) \right), \end{aligned} \quad (58)$$

which proves the first part of the lemma. By (54) and (57), we have

$$|\check{M}_m| \leq O \left(\sqrt{(\check{M}_{m+1} + O(d \log^5(dHK)) + 2^{m+1}(K + \check{R}_0)) \log(1/\delta)} + \log(1/\delta) \right), \quad (59)$$

which proves the second part of the lemma. $\square$

E.7 Proof of Lemma 23

Proof. Let $b_m = \check{R}_m + |\check{M}_m|$. By (58) and (59), we can bound b_m recursively as

$$b_m \leq O \left(\sqrt{d^9 \log^5(Td) \log \frac{1}{\delta}} \sqrt{b_m + b_{m+1} + 2^{m+1}(K + \check{R}_0) + d^7 \log^6(Td) \log \frac{1}{\delta}} \right). \quad (60)$$

Note that $b_m \leq 2KH$ for $m \in \Lambda_1$. By Lemma 12 with parameters

$$\lambda_1 = 2KH, \quad \lambda_2 = \sqrt{d^9 \log^5(Td) \log(1/\delta)}, \quad \lambda_3 = K + \check{R}_0, \quad \lambda_4 = d^7 \log^6(Td) \log(1/\delta),$$

we obtain that

$$\check{R}_0 \leq b_0 \leq O\left(\sqrt{d^9(K + \check{R}_0) \log^5(Td) \log(1/\delta)} + d^9 \log^6(Td) \log(1/\delta)\right),$$

which implies

$$b_0 \leq O\left(d^{4.5} \sqrt{K \log^5(Td) \log(1/\delta)} + d^9 \log^6(Td) \log(1/\delta)\right)$$

and completes the proof. □