

Optimal Exploration for Model-Based RL in Nonlinear Systems

Andrew Wagenmaker*

Guanya Shi[†]

Kevin Jamieson[‡]

Abstract

Learning to control unknown nonlinear dynamical systems is a fundamental problem in reinforcement learning and control theory. A commonly applied approach is to first explore the environment (exploration), learn an accurate model of it (system identification), and then compute an optimal controller with the minimum cost on this estimated system (policy optimization). While existing work has shown that it is possible to learn a uniformly good model of the system (Mania et al., 2022), in practice, if we aim to learn a good controller with a low cost on the actual system, certain system parameters may be significantly more critical than others, and we therefore ought to focus our exploration on learning such parameters.

In this work, we consider the setting of nonlinear dynamical systems and seek to formally quantify, in such settings, (a) which parameters are most relevant to learning a good controller, and (b) how we can best explore so as to minimize uncertainty in such parameters. Inspired by recent work in linear systems (Wagenmaker et al., 2021), we show that minimizing the controller loss in nonlinear systems translates to estimating the system parameters in a particular, task-dependent metric. Motivated by this, we develop an algorithm able to efficiently explore the system to reduce uncertainty in this metric, and prove a lower bound showing that our approach learns a controller at a near-instance-optimal rate. Our algorithm relies on a general reduction from policy optimization to optimal experiment design in arbitrary systems, and may be of independent interest. We conclude with experiments demonstrating the effectiveness of our method in realistic nonlinear robotic systems¹.

1 Introduction

Controlling nonlinear dynamical systems is a core problem in robotics, cyber-physical systems, and beyond, and a significant body of work in both the control theory and reinforcement learning communities has sought to address this challenge (Slotine et al., 1991; Åström & Wittenmark, 2013; Sutton & Barto, 2018). In many real-world scenarios (Shi et al., 2019; Ljung, 1998; Nguyen-Tuong & Peters, 2011; Brunke et al., 2022), the dynamics of the system of interest is unknown, or only a coarse model of them is available, which significantly increases the challenge of control—not only must we control such systems, we must *learn* to control them. While a variety of methods exist to address this challenge, a commonly applied approach is to first perform *system identification*, learning an accurate model of the system’s dynamics, and then use this model to obtain a controller. Despite its promising potential, there are still several fundamental questions that must be answered to make this approach practically effective.

*University of Washington, Seattle. Email: ajwagen@cs.washington.edu

[†]University of Washington, Seattle. Email: guanyas@cs.washington.edu

[‡]University of Washington, Seattle. Email: jamieson@cs.washington.edu

¹Code: https://github.com/ajwagen/nonlinear_sysid_for_control

Which parameters are most relevant to learning a good controller? Beyond some special cases, little work has been done characterizing how the estimation error from system identification translates to end-to-end suboptimality in the resulting controller of our nonlinear systems. In particular, certain parameters of the system or regions of the state space may be irrelevant to learning a good controller, and coarse estimates of these parameters would suffice, while other parameters may be critical to learning a good controller, and we must therefore estimate these parameters very accurately in order to effectively control the system. In the context of this work, where nonlinearities are considered, the heterogeneity of the parameters is further accentuated. For instance, around a point of equilibrium, some system parameters might be completely inactive, having no impact on the dynamics (see the example in Section 1.1 for an illustration of this).

How can we best explore so as to minimize uncertainty in relevant parameters? Even if we are able to determine which parameters are most important for obtaining a good controller on the true system, it is not obvious how to use this information. How can we direct our system identification phase in order to focus on learning these parameters as quickly as possible, without spending time estimating the parameters of the system less critical for control? This is fundamentally a question of *exploration*. While it is known in linear systems that random excitation will efficiently explore (Simchowitz et al., 2018), exploration in nonlinear systems is significantly more challenging since, in order to excite all parameters of interest, non-trivial planning may be required to ensure all relevant states are reached (as is the case in the example considered in Section 1.1).

We address both these questions in a particular class of nonlinear systems parameterized as:

$$\mathbf{x}_{h+1} = A_\star \phi(\mathbf{x}_h, \mathbf{u}_h) + \mathbf{w}_h. \quad (1.1)$$

Here $\mathbf{x}_h \in \mathbb{R}^{d_x}$ denotes the state of the system, $\mathbf{u}_h \in \mathcal{U} \subseteq \mathbb{R}^{d_u}$ the input, $\mathbf{w}_h \sim \mathcal{N}(0, \sigma_w^2 \cdot I)$ random noise, $\phi(\cdot, \cdot) \in \mathbb{R}^{d_\phi}$ a (possibly nonlinear, known) feature map, and $A_\star \in \mathbb{R}^{d_x \times d_\phi}$ the (unknown) system parameter. Systems of this form are able to model a variety of real-world settings (Shi et al., 2021a; O’Connell et al., 2022; Boffi et al., 2021; Song & Sun, 2021; Richards et al., 2021)², and have been the subject of recent attention in the reinforcement learning community (Mania et al., 2022; Kakade et al., 2020; Song & Sun, 2021), yet the aforementioned questions have remained unanswered. Towards addressing this, in this work we make the following contributions:

1. For systems of the form (1.1), given some cost of interest which we wish to find a controller to minimize, we (a) formally characterize how estimation error translates into suboptimality in the learned controller, under the *certainty equivalent* control rule and (b) provide a lower bound on the loss of *any* (sufficiently regular) control rule learned from T rounds of interaction with (1.1).
2. Motivated by this characterization, we present an algorithm which achieves the *instance-optimal* rate, with controller loss matching our lower bound. To the best of our knowledge, this is the first statistically optimal algorithm in the setting of nonlinear dynamical systems. Our algorithm relies on a generic reduction from policy optimization to optimal exploration in *arbitrary* dynamical systems (not necessarily of the form (1.1)), which may be of independent interest.

²In real-world settings, ϕ is typically (1) from physics (i.e., the system structure is known but some parameters such as drag coefficient are unknown (Slotine et al., 1991)), (2) learned using representation learning or meta-learning (O’Connell et al., 2022; Richards et al., 2021), and/or (3) from random features (e.g., any sufficiently regular, smooth nonlinear system $\mathbf{f}(\mathbf{x}, \mathbf{u})$ can be modeled by (1.1) using N random features up to a $1/\sqrt{N}$ error (Rahimi & Recht, 2008)).

3. We present numerical experiments on several realistic nonlinear systems which illustrate that our approach—efficiently exploring to reduce uncertainty in parameters most relevant to learning a controller—yields significant gains in practice.

Our work builds on the recent work of [Wagenmaker et al. \(2021\)](#), which addresses a similar set of challenges in the linear dynamical systems setting—we extend this work to the nonlinear setting. To further motivate our approach, we consider the following example.

1.1 Motivating Example

To motivate the need for effective exploration, we consider a simple 1-D system with nonlinear dynamics given by:

$$\mathbf{x}_{h+1} = a_1 \mathbf{x}_h + a_2 \mathbf{u}_h + \sum_{i=1}^{10} a_{i+2} \phi_i(\mathbf{x}_h) + \mathbf{w}_h$$

where $\phi_i(\mathbf{x}) = \max\{1 - 100(\mathbf{x} - c_i)^2, 0\}$ for some c_i . We choose $a_1 = 0.8, a_2 = 1$, and $a_3 = \dots = a_{12} = -3$. We assume $a_{1:12}$ are unknown, $(\phi_i)_{i=1}^{10}$ is known, and set

$$\text{cost}(\mathbf{x}, \mathbf{u}) = (\mathbf{x} - c_1)^2 + 100^{-1} \cdot \mathbf{u}^2.$$

With this choice of cost, the optimal controller

will attempt to direct the state $\mathbf{x}$ to the equilibrium point c_1 and maintain this position. Note that, with our choice of ϕ_i , $\phi_i(\mathbf{x}) \neq 0$ only when $\mathbf{x}$ is very close to c_i . This renders the parameters $a_{4:12}$ irrelevant to learning the optimal controller, since $\phi_2, \dots, \phi_{10}$ will be inactive if we are playing optimally, but learning a_3 is critical to performing optimally, as its value significantly changes the dynamics at the goal state.

We illustrate the result of running on this system in Figure 1, comparing our proposed approach (**Task-Driven Exploration**, Algorithm 1) to the approach which chooses $\mathbf{u}_h \sim \mathcal{N}(0, \sigma_u^2)$ (**Random Exploration**), and the approach proposed in [Mania et al. \(2022\)](#) (**Uniform Exploration**) which seeks to explore so as to estimate $a_{1:12}$ uniformly well. As can be seen, neither of these latter two approaches are able to learn a good controller, while our approach easily finds a near-optimal controller. The failure modes of each of these approaches is somewhat different. Here **Random Exploration** fails since the chance of reaching the point $\mathbf{x}_h \sim c_1$ is extremely small if the input is random noise—reaching c_1 requires playing a particular sequence of actions which are very unlikely to be played if $\mathbf{u}_h$ is chosen randomly. The **Uniform Exploration** approach does, in contrast, plan and, given enough time, is guaranteed to estimate all parameters accurately. However, as it aims to estimate all parameters uniformly well, it will attempt to estimate $a_{4:12}$ accurately despite their irrelevance to control, which will slow down the rate at which it is able to estimate a_3 . Only our approach, which both plans and takes into account the cost while exploring, is able to reach a_3 enough times to efficiently estimate it, and learn a good controller.

This example illustrates that it is critical both to explore efficiently, and also to let the objective—learning a good controller—guide this exploration. We emphasize that the behavior in this example is only exhibited in nonlinear systems—though taking into account the task while exploring in linear systems is known to yield provable improvements ([Wagenmaker et al., 2021](#)), even playing random noise allows every direction to be learned in such systems. In nonlinear systems, however, this is not the case—one may fail to learn completely unless careful planning is performed.

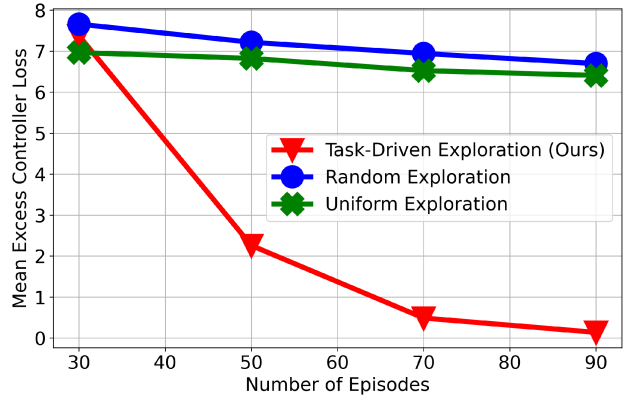


Figure 1: Performance on Motivating Example

2 Related Work

Online learning and control. Recently, there has been increased interest in studying online learning and control from a learning-theoretic perspective, largely for settings with linear systems such as online LQR or LQG with unknown dynamics (Abbasi-Yadkori et al., 2011; Simchowitz et al., 2018, 2019; Mania et al., 2019; Cohen et al., 2019; Dean et al., 2020; Yu et al., 2020a; Wagenmaker & Jamieson, 2020; Simchowitz & Foster, 2020; Simchowitz et al., 2020). In the nonlinear setting, (Foster et al., 2020; Oymak, 2019; Sattar & Oymak, 2022) provide formal guarantees on system identification in several different classes of nonlinear systems, yet they only consider noiseless systems, or systems that are significantly easier to excite than (1.1) (rendering the problem of exploration significantly easier). Kakade et al. (2020) study systems of the form (1.1), but consider only the regret minimization problem. While their bounds would yield a polynomial complexity via an online-to-batch conversion, our characterization is significantly tighter. The most relevant work, Mania et al. (2022), proposes an active learning approach to identify unknown parameters in (1.1), with the goal of minimizing the Euclidean distance in the parameter space. However, as we show, learning a uniformly good model could be significantly worse than learning a model with the goal task in mind. Also very related to our work is Wagenmaker et al. (2021), which seeks to answer a similar set of questions as what we consider: performing system identification in order to learn a good controller. This work is restricted to the setting of linear dynamics, however, and does not address the additional complexity of exploration in nonlinear systems.

System identification, dual control, and iterative learning control. There is a large body of classical work in system identification (Ljung, 1998), and our work can be seen as an instance of *active* system identification. While a variety of approaches have been proposed which study similar problems (Mehra, 1974; Gerencsér & Hjalmarsson, 2005; Katselis et al., 2012; Manchester, 2010; Rojas et al., 2007; Goodwin & Payne, 1977; Lindqvist & Hjalmarsson, 2001; Gerencsér et al., 2007), then tend to only consider linear systems, or lack rigorous theoretical guarantees. Recently deep learning approaches have also been applied in system identification (Shi et al., 2019; Nguyen-Tuong & Peters, 2011; Brunke et al., 2022; Williams et al., 2017; Shi et al., 2021b). In these works, the system identification phase is separate from the downstream controller design. Instead, in the control community, estimating parameters while simultaneously or iteratively optimizing for performance has been formulated as a dual control problem (Feldbaum, 1960; Mesbah, 2018) or an iterative learning control problem (Bristow et al., 2006). However, both settings focus on stability, robustness, or asymptotic convergence whereas our work quantifies the end-to-end suboptimality gap with a statistically optimal algorithm.

Model-based reinforcement learning. This paper falls into the broad category of model-based reinforcement learning (MBRL), where an agent explores the environment to learn a model and then computes an optimal policy using the learned model. On the empirical side, deep MBRL has made exciting progress in many domains (Kaiser et al., 2019; Yu et al., 2020b; Chua et al., 2018). Several task-aware methods have been designed to improve MBRL’s performance, such as uncertainty-aware policy optimization (Yu et al., 2020b; Chua et al., 2018) and active exploration to reduce model uncertainty (Nakka et al., 2020), yet these works lack formal guarantees. On the theoretical side, a variety of different model-based approaches exist (Osband & Van Roy, 2014; Sun et al., 2019; Agarwal et al., 2020; Zhou et al., 2021; Zanette & Brunskill, 2019; Azar et al., 2017; Song & Sun, 2021); however, the majority of these consider restricted settings such as tabular or linear MDPs. Of particular interest is the work of Song & Sun (2021) which presents a result in

systems of the form (1.1). While they show that polynomial sample complexity is possible, our results yield a significantly tighter characterization.

Adaptive nonlinear control. Adaptive nonlinear control also seeks to control an unknown nonlinear system with parametric uncertainties (Slotine et al., 1991; Åström & Wittenmark, 2013). In particular, the key idea of model-reference adaptive control (MRAC) bears affinity to this paper, in that the adaptation law in MRAC adapts unknown parameters in a task-aware manner, by relating the tracking error with the estimated parameter in a closed loop. In fact, the parameter estimation error in MRAC converges *only when necessary*, i.e., when the task is “rich” enough (the formal condition is called persistent excitation (Åström & Wittenmark, 2013; Slotine et al., 1991)). There are two main differences between MRAC and our work. First, adaptive control does not explicitly optimize a cost function. The objective of adaptive control is often tracking error convergence and Lyapunov stability, whereas our framework allows general cost functions. Moreover, adaptive control theory typically focuses on asymptotic convergence, but we give non-asymptotic optimality guarantees. Second, adaptive control has by and large been limited to specific system classes (e.g., fully-actuated systems (Åström & Wittenmark, 2013; Richards et al., 2021)) and policy classes (e.g., policy to directly cancel out the matched uncertainty (O’Connell et al., 2022; Boffi et al., 2021)), whereas our framework allows more general systems and policy classes.

3 Preliminaries

Notation. $\|\cdot\|_{\text{op}}$ denotes the operator norm (matrix 2-norm), $\|\cdot\|_F$ the Frobenius norm, and $\|\cdot\|_M$ the Mahalanobis norm, defined as $\|\mathbf{x}\|_M := \sqrt{\mathbf{x}^\top M \mathbf{x}}$ for $M \succeq 0$. $\text{vec}(A)$ denotes the vectorization of matrix A . $\mathcal{B}_p(A; r) := \{A' : \|A - A'\|_p \leq r\}$. $[H] = \{1, 2, \dots, H\}$. $\Delta_{\mathcal{X}}$ denotes the set of distributions over set $\mathcal{X}$. We let $\mathcal{S}^{d-1}$ refer to the unit ball in d dimensions and $\mathbb{S}_+^d$ (resp. $\mathbb{S}_{++}^d$) the set of positive semi-definite matrices (resp. positive definite matrices) in $\mathbb{R}^{d \times d}$. We let $\mathbb{E}_A[\cdot]$ denote the expectation over trajectories induced on system with parameter A , and $\mathbb{E}_{A,\pi}[\cdot]$ the expectation induced when policy π is played. Throughout, $\mathcal{O}(\cdot)$ denotes standard big-O notation, $\tilde{\mathcal{O}}(\cdot)$ hides additional logarithmic factors, and we use $\lesssim$ informally to highlight key parameters in an inequality.

Setting. In this work, we are interested in systems of the form (1.1). We consider the episodic setting, where episodes are of length H , and assume that each episode starts from a given state $\mathbf{x}_1$. We also assume $\|A_\star\|_{\text{op}} \leq B_A$ for some known $B_A > 0$. We note that the setting considered here encompasses many real-world systems of interest in robotics and control (e.g., (O’Connell et al., 2022; Song & Sun, 2021; Shi et al., 2021a; Richards et al., 2021) and Section 6).

The goal of the learner is to find a policy (controller) $\pi = (\pi_h)_{h=1}^H$ which achieves minimal cost on (1.1), for the cost defined by some (known) function $(\text{cost}_h(\cdot, \cdot))_{h=1}^H$, with $\text{cost}_h : \mathbb{R}^{d_x} \times \mathcal{U} \rightarrow \mathbb{R}_+$. For a given policy π , we define the expected cost on system A as

$$\mathcal{J}(\pi; A) := \mathbb{E}_{A,\pi} \left[\sum_{h=1}^H \text{cost}_h(\mathbf{x}_h, \mathbf{u}_h) \right].$$

We consider the following interaction protocol:

1. Learner interacts with system (1.1) for T episodes, at every episode playing an exploration policy $\pi_{\text{exp}} \in \Pi_{\text{exp}}$.

2. After T episodes, the learner proposes a policy $\hat{\pi}_T \in \Pi^*$.
3. The learner suffers cost $\mathcal{J}(\hat{\pi}_T; A_\star)$.

The goal of the learner is therefore first to explore and, after T episodes of exploration, to propose its best guess at the optimal controller for (1.1), $\hat{\pi}_T$. Here we take Π_{exp} to be a (known) set of admissible exploration policies (for example, policies with bounded input power), and Π^* a (known) set of admissible control policies. We assume that policies in Π^* are deterministic, but allow for randomized policies in Π_{exp} . Policies may be either open- or closed-loop. Note that we do not assume $\Pi^* = \Pi_{\text{exp}}$ —in general Π_{exp} need not be equal to Π^* .

System Notation. Before proceeding, we introduce several additional pieces of notation. First, we let $\mathcal{T}$ denote the space of all possible state-input trajectories, $\mathcal{T} \subseteq (\mathbb{R}^{d_x} \times \mathcal{U})^H \times \mathbb{R}^{d_x}$, and, for any $\tau \in \mathcal{T}$, let $\tau_{1:h}$ denote the first h states and inputs in τ . Second, for any policy π , we denote

$$\Lambda_{A,\pi} := \mathbb{E}_{A,\pi} \left[\sum_{h=1}^H \phi(\mathbf{x}_h, \mathbf{u}_h) \phi(\mathbf{x}_h, \mathbf{u}_h)^\top \right]$$

the expected covariance induced by playing π on system A . In particular, we set $\Lambda_\pi := \Lambda_{A_\star, \pi}$. We also denote $\tilde{\Lambda} := I_{d_x} \otimes \Lambda$ the Kronecker product of I_{d_x} and Λ . Finally, we let Ω denote the set of all possible covariance matrices induced by playing mixtures of policies in Π_{exp} :

$$\Omega := \{ \mathbb{E}_{\pi \sim \omega} [\Lambda_\pi] : \omega \in \Delta_{\Pi_{\text{exp}}} \},$$

where $\Delta_{\Pi_{\text{exp}}}$ denotes the set of distributions over Π_{exp} .

3.1 Regularity Assumptions

In order to make learning in (1.1) tractable, we need several regularity assumptions. We first introduce assumptions on the boundedness of the feature map ϕ , the boundedness of the cost, and the achievable minimum eigenvalue.

Assumption 1 (Bounded Features). *For all $\mathbf{x} \in \mathbb{R}^{d_x}$ and $\mathbf{u} \in \mathcal{U}$, we have $\|\phi(\mathbf{x}, \mathbf{u})\|_2 \leq B_\phi$.*

Assumption 2 (Bounded Cost). *There exists some $r_{\text{cost}}(A_\star) > 0$ such that, for all $A \in \mathcal{B}_F(A_\star; r_{\text{cost}}(A_\star))$ and all $\pi \in \Pi^*$, we have $\mathbb{E}_{A,\pi}[(\sum_{h=1}^H \text{cost}_h(\mathbf{x}_h, \mathbf{u}_h))^2] \leq L_{\text{cost}}$.*

Assumption 3 (Uniform Feature Excitation). *There exists $\omega \in \Delta_{\Pi_{\text{exp}}}$ such that $\lambda_{\min}(\mathbb{E}_{\pi \sim \omega} [\Lambda_{\pi_{\text{exp}}})) \geq \lambda_{\min}^*$ for some $\lambda_{\min}^* > 0$.*

We remark that these assumptions have appeared before in work on systems of the form (1.1) (Mania et al., 2022; Kakade et al., 2020). In order to precisely quantify the optimal rates of learning, we require that our system satisfy certain smoothness assumptions. First, we require that $\phi(\cdot, \cdot)$ is differentiable in its second argument.

Assumption 4 (Smooth Nonlinearity). *For all $\mathbf{x} \in \mathbb{R}^{d_x}$ and $\mathbf{u} \in \mathcal{U}$, $\phi(\mathbf{x}, \mathbf{u})$ is four-times differentiable in $\mathbf{u}$. Furthermore, $\|\nabla_{\mathbf{u}}^{(i)} \phi(\mathbf{x}, \mathbf{u})\|_{\text{op}} \leq L_\phi$, $\forall i \in \{1, 2, 3, 4\}$, $\mathbf{x} \in \mathbb{R}^{d_x}$, and $\mathbf{u} \in \mathcal{U}$.*

We also require that the class of admissible control policies, Π^* , has the following parametric form:

$$\Pi^* = \{ \pi^\theta : \theta \in \mathbb{R}^{d_\theta} \}$$

and that the parameterization is smooth in the following sense.

Assumption 5 (Smooth Controller Class). $\pi_h^\theta(\tau_{1:h})$ is four-times differentiable in θ for all $\tau \in \mathcal{T}$ and $h \in [H]$. Furthermore, $\|\nabla_\theta^{(i)} \pi_h^\theta(\tau_{1:h})\|_{\text{op}} \leq L_\theta$ for $\forall i \in \{1, 2, 3, 4\}$, $\theta \in \mathbb{R}^{d_\theta}$, and $\tau \in \mathcal{T}$.

Assumption 5 is satisfied for commonly considered classes of controllers, such as linear controllers, but is also satisfied by more complex classes such as neural network controllers. While the learner may propose any $\hat{\pi}_T \in \Pi^*$, we are particularly interested in the *certainty equivalence* decision rule (i.e., the learner decides $\hat{\pi}_T$ as if the estimated system is the actual one), defined as:

$$\pi_\star(A) := \pi^{\theta_\star(A)} \quad \text{for} \quad \theta_\star(A) := \arg \min_{\theta \in \mathbb{R}^{d_\theta}} \mathcal{J}(\pi^\theta; A). \quad (3.1)$$

To ensure that $\pi_\star(A)$ is well-defined and sufficiently regular, we make the following assumption.

Assumption 6 (Unique Optimal Controller). We assume that the global minimum of $\mathcal{J}(\pi^\theta; A_\star)$, $\theta_\star(A_\star)$, is unique, and that $\nabla_\theta^2 \mathcal{J}(\pi^\theta; A_\star)|_{\theta=\theta_\star(A_\star)} \succ 0$.

In general, the policy optimization problem in (3.1) may not be computationally tractable. As we show in Appendix D, the globally optimal decision rule of (3.1) can be replaced with a locally optimal decision rule (i.e. $\pi_\star(A)$ a local minimum of $\mathcal{J}(\pi; A)$). Furthermore, Assumption 6 can be replaced by assuming the differentiability of $\theta_\star(A)$ with respect to A for A near $A_\star$. For ease of exposition, in the main text we assume that Assumption 6 holds and that $\pi_\star(A)$ is defined as in (3.1). With these definitions and under Assumptions 1, 2, 4 and 5, we can show that $\mathcal{J}(\pi^\theta; A_\star)$ is differentiable in θ and, combined with Assumption 6, that $\theta_\star(A)$ is differentiable in A , for $A \in \mathcal{B}_F(A_\star; r_\theta(A_\star))$ and some $r_\theta(A_\star) > 0$. We let $L_{\pi_\star}$ denote an upper bound on the norm of the derivatives of $\theta_\star(A)$. We always take $B_\phi, L_{\text{cost}}, L_\phi, L_\theta, L_{\pi_\star} \geq 1$. Additional discussion on the setting of $\pi_\star(A)$ and the scaling of $r_\theta(A_\star)$ and $L_{\pi_\star}$ is given in Appendix D.

4 Optimal Exploration in Nonlinear Systems

In this work, we are interested in characterizing the instance-optimal rates of learning a controller $\pi \in \Pi^*$ which minimizes the loss $\mathcal{J}(\pi; A_\star)$. The following result, a generalization of Proposition 8.2 of Wagenmaker et al. (2021) to nonlinear systems, is the starting point of our analysis, and precisely quantifies how estimation error translates to controller loss.

Proposition 1 (Informal). Under Assumptions 1, 2 and 4 to 6 and on the system (1.1), we have

$$\mathcal{J}(\pi_\star(\hat{A}); A_\star) - \mathcal{J}(\pi_\star(A_\star); A_\star) = \|\text{vec}(A_\star - \hat{A})\|_{\mathcal{H}(A_\star)}^2 + \mathcal{O}^\star(\|A_\star - \hat{A}\|_F^3)$$

for

$$\mathcal{H}(A_\star) := \nabla_A^2 \mathcal{J}(\pi_\star(A); A_\star)|_{A=A_\star}$$

and where $\mathcal{O}^\star(\cdot)$ hides factors polynomial in the regularity parameters of Assumptions 1 to 6.

The quantity $\mathcal{H}(A_\star) := \nabla_A^2 \mathcal{J}(\pi_\star(A); A_\star)|_{A=A_\star}$, referred to as the *model-task Hessian* in Wagenmaker et al. (2021), corresponds to the *curvature* of the loss of the certainty-equivalence controller $\pi_\star(A)$ around $A \leftarrow A_\star$. It precisely quantifies how estimation error in each coordinate of $A_\star$ translates into suboptimality of the controller—providing an answer to our question of which parameters are most relevant to learning a good controller—and reduces the problem of minimizing the controller loss to estimating $A_\star$ in a particular norm. The following result gives a bound on this estimation error, $\|\text{vec}(A_\star - \hat{A})\|_{\mathcal{H}(A_\star)}^2$.

Proposition 2 (Informal). *Consider interacting with (1.1) for T episodes, and let*

$$\mathbf{\Lambda}_T = \sum_{t=1}^T \sum_{h=1}^H \phi(\mathbf{x}_h^t, \mathbf{u}_h^t) \phi(\mathbf{x}_h^t, \mathbf{u}_h^t)^\top$$

denote the observed covariates and

$$\hat{A} = \arg \min_A \sum_{t=1}^T \sum_{h=1}^H \|\mathbf{x}_{h+1}^t - A\phi(\mathbf{x}_h^t, \mathbf{u}_h^t)\|_2^2$$

the least-squares estimate of $A_\star$. Recalling that $\check{\mathbf{\Lambda}}_T = I_{d_x} \otimes \mathbf{\Lambda}_T$, we have, with high probability:

$$\|\text{vec}(A_\star - \hat{A})\|_{\mathcal{H}(A_\star)}^2 \lesssim \sigma_w^2 \cdot \text{tr}(\mathcal{H}(A_\star) \check{\mathbf{\Lambda}}_T^{-1}).$$

4.1 Algorithm and Upper Bound

Proposition 2 motivates our algorithmic approach: explore to collect covariates $\mathbf{\Lambda}_T$ minimizing $\text{tr}(\mathcal{H}(A_\star) \check{\mathbf{\Lambda}}_T^{-1})$. There are two primary challenges to achieving this: we do not know $\mathcal{H}(A_\star)$, as it depends on the (unknown) parameter $A_\star$ and, even if we did know $\mathcal{H}(A_\star)$, it is not clear how to explore so as to collect data minimizing $\text{tr}(\mathcal{H}(A_\star) \check{\mathbf{\Lambda}}_T^{-1})$. We address both of these challenges with our main algorithm, Algorithm 1.

Algorithm 1 Optimal Exploration in Nonlinear Systems (informal)

- 1: **inputs:** episodes T , $(\text{cost}_h)_{h=1}^H$, confidence δ , control policies $\Pi^\star$, exploration policies Π_{exp}
 - 2: $\hat{A}^1 \leftarrow 0$, $\ell_T \leftarrow \lceil \log_2 T/8 \rceil$, $T_\ell \leftarrow 2^\ell$
 - 3: **for** $\ell = 1, 2, 3, \dots, \ell_T$ **do**
 - 4: Compute estimate of model-task Hessian: $\mathcal{H}_\ell \leftarrow \mathcal{H}(\hat{A}^\ell)$
 - 5: Run DYNAMICOED on $\Phi_\ell(\mathbf{\Lambda}) \leftarrow \text{tr}(\mathcal{H}_\ell \cdot \mathbf{\Lambda}^{-1})$ to learn exploration policies $\Pi_\ell \subseteq \Pi_{\text{exp}}$
 - 6: Rerun each policy in Π_ℓ $N_\ell = \lceil T_\ell / |\Pi_\ell| \rceil$ times, denote collected data $\mathfrak{D}_\ell$
 - 7: Estimate $A_\star$: $\hat{A}^{\ell+1} = \arg \min_A \sum_{h=1}^H \sum_{(\mathbf{x}_{h+1}, \mathbf{u}_h, \mathbf{x}_h) \in \mathfrak{D}_\ell} \|\mathbf{x}_{h+1} - A\phi(\mathbf{x}_h, \mathbf{u}_h)\|_2^2$
 - 8: **return** $\hat{\pi}_T \leftarrow \pi_\star(\hat{A}^{\ell_T+1}) \in \Pi^\star$
-

Algorithm 1 proceeds in epochs of exponentially increasing length. At each epoch it first approximates $\mathcal{H}(A_\star)$ by computing the model-task Hessian of the estimated system, $\hat{A}^\ell$. Using this approximation of $\mathcal{H}(A_\star)$, it seeks to explore to minimize $\text{tr}(\mathcal{H}(\hat{A}^\ell) \check{\mathbf{\Lambda}}_T^{-1})$. This exploration routine is encapsulated in the DYNAMICOED (dynamic optimal experiment design) function, an adaptive experiment-design routine inspired by recent work in reinforcement learning (Wagenmaker & Jamieson, 2022) and described in more detail in Section 5. DYNAMICOED returns a set of exploration policies, Π_ℓ , which we run to collect data $\mathfrak{D}_\ell$. As we will show, the collected covariates, $\mathbf{\Lambda}_\ell := \sum_{h=1}^H \sum_{(\mathbf{u}_h, \mathbf{x}_h) \in \mathfrak{D}_\ell} \phi(\mathbf{x}_h, \mathbf{u}_h) \phi(\mathbf{x}_h, \mathbf{u}_h)^\top$, satisfy

$$\text{tr}(\mathcal{H}(\hat{A}^\ell) \check{\mathbf{\Lambda}}_\ell^{-1}) \lesssim T_\ell^{-1} \cdot \min_{\mathbf{\Lambda} \in \Omega} \text{tr}(\mathcal{H}(\hat{A}^\ell) \check{\mathbf{\Lambda}}^{-1}),$$

which implies that DYNAMICOED collects data minimizing $\text{tr}(\mathcal{H}(\hat{A}^\ell) \check{\mathbf{\Lambda}}_\ell^{-1})$ at a near-optimal rate. Given the data $\mathfrak{D}_\ell$, we form the least-squares estimate of $A_\star$, $\hat{A}^{\ell+1}$, and the process repeats. After running for T episodes, the certainty-equivalence controller on the last estimate obtained, $\hat{\pi}_T = \pi_\star(\hat{A}^{\ell_T+1})$, is returned. The following result bounds the suboptimality of $\hat{\pi}_T$ as compared to $\pi_\star(A_\star)$.

Theorem 1. Under Assumptions 1 to 6, if $T \geq C_{\text{poly}} \cdot \max\{1, r_{\text{cost}}(A_\star)^{-2}, r_\theta(A_\star)^{-2}\}$, then with probability at least $1 - \delta$, Algorithm 1 explores with policies in Π_{exp} at every episode, runs for at most T episodes, and returns $\hat{\pi}_T \in \Pi^\star$ satisfying:

$$\mathcal{J}(\hat{\pi}_T; A_\star) - \mathcal{J}(\pi_\star(A_\star); A_\star) \leq \frac{\sigma_w^2}{T} \cdot \min_{\Lambda \in \Omega} \text{tr}(\mathcal{H}(A_\star)\check{\Lambda}^{-1}) \cdot C \log \frac{6d_x d_\phi}{\delta} + \frac{C_{\text{poly}}}{T^{3/2}}$$

where we recall Ω is the set of possible expected covariates on (1.1), C is a universal constant, and

$$C_{\text{poly}} = \text{poly}(d_\phi, d_x, H, B_A, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, L_{\pi_\star}, \sigma_w, \sigma_w^{-1}, \frac{1}{\lambda_{\min}^\star}, \log \frac{T}{\delta}).$$

Theorem 1 shows that Algorithm 1 is able to explore so as to optimally minimize the exploration loss $\text{tr}(\mathcal{H}(A_\star)\check{\Lambda}_T^{-1})$, up to a lower-order term scaling as $T^{-3/2}$ and polynomially in system parameters. While Propositions 1 and 2 together show that collecting data which minimizes $\text{tr}(\mathcal{H}(A_\star)\check{\Lambda}_T^{-1})$ is in some sense fundamental to minimizing the cost of the certainty equivalent controller, it is not clear that this is necessary. In the following section, we show that this is indeed the case.

Remark 4.1 (Comparison to TOPLE Algorithm of Wagenmaker et al. (2021)). Algorithm 1 bears many similarities to the TOPLE algorithm of Wagenmaker et al. (2021), which performs an analogous task-driven exploration routine, but in the setting of linear dynamical systems. As noted in Section 1.1, the key challenge present in the nonlinear case compared to the linear is that, while in the linear case random noise will excite every direction, in the nonlinear case, the learner must actually traverse the system in order to reach the states that will excite the nonlinear modes. Though the overall structure of Algorithm 1 is similar to TOPLE, this added challenge requires a much more powerful exploration routine, encapsulated in the DYNAMICOED function and described in more detail in Section 5.

Remark 4.2 (Computational Efficiency of Algorithm 1). The primary computational burden of Algorithm 1 is in the computation of $\mathcal{H}(\hat{A}^\ell)$ —which involves differentiating $\pi_\star(\hat{A}^\ell)$ —the computation of $\pi_\star(\hat{A}^{\ell_T+1})$, and the DYNAMICOED subroutine. In general, if we define $\pi_\star(A)$ as in (3.1), it may not be efficiently computable, as it involves solving a possibly non-convex optimization problem. However, as we show in Appendix D, we can instead set $\pi_\star(A)$ to correspond to a local minimum rather than a global minimum of the loss, which will render it efficiently computable (though note that Theorem 1 will still in this case only bound the suboptimality of $\hat{\pi}_T$ as compared to $\pi_\star(A_\star)$). We discuss the computational efficiency of DYNAMICOED in more detail in Section 5, but note that in general it may not be computationally efficient as it relies on calls to the LC³ algorithm of Kakade et al. (2020), which requires access to a computational oracle. Despite these computational challenges, in Section 6 we demonstrate that in practice, by making several reasonable approximations, Algorithm 1 can be implemented efficiently, and that this efficient implementation performs very well on realistic systems.

4.2 Lower Bounds on Learning Controllers

Our goal is to show that, up to constants and lower-order terms, the bound given in Theorem 1 is not improvable, regardless of which controller estimate we use. To obtain such lower bounds, we need several additional assumptions. In particular, we require that the loss $\mathcal{J}(\pi^\theta; A)$ grows quadratically in the distance θ is from $\theta_\star(A)$, and strengthen Assumption 3 to ensure (1.1) is sufficiently easy to excite. Formal statements of these conditions are given in Appendix F. Our lower bound is as follows.

Theorem 2 (Informal). *Under Assumptions 1 to 6 and the additional regularity assumptions mentioned above, as long as $T \geq C_{\text{lb}}$, for any $\omega_{\text{exp}} \in \Delta_{\Pi_{\text{exp}}}$, we have*

$$\min_{\hat{\pi}} \max_{A \in \mathcal{B}_T} \mathbb{E}_{\mathcal{D}_T \sim A, \omega_{\text{exp}}} [\mathcal{J}(\hat{\pi}(\mathcal{D}_T); A) - \mathcal{J}(\pi_*(A); A)] \geq \frac{\sigma_w^2}{3T} \cdot \min_{\Lambda \in \Omega} \text{tr}(\mathcal{H}(A_*) \check{\Lambda}^{-1}) - \frac{C_{\text{lb}}}{T^{5/4}}$$

for $\mathcal{B}_T := \mathcal{B}_F(A_*; \mathcal{O}(T^{-5/6}))$, $\mathbb{E}_{\mathcal{D}_T \sim A, \omega_{\text{exp}}}[\cdot] = \mathbb{E}_{\pi_{\text{exp}} \sim \omega_{\text{exp}}}[\mathbb{E}_{\mathcal{D}_T \sim A, \pi_{\text{exp}}}[\cdot]]$ the expectation over trajectories generated by running policies $\pi \sim \omega_{\text{exp}}$ on system A for T episodes, $\hat{\pi}$ any mapping from observations to policies in Π^* , and C_{lb} some value scaling polynomially in problem parameters.

Note that this lower bound holds for *any* A_* and mapping ϕ , as long as our assumptions are met. Up to constants and lower-order terms, the scaling of Theorem 2 matches that of Theorem 1—both scale with $\min_{\Lambda \in \Omega} \text{tr}(\mathcal{H}(A_*) \check{\Lambda}^{-1})$ —which implies that Algorithm 1 is indeed optimal (under certain additional regularity conditions). To the best of our knowledge, this is the first result characterizing the optimal statistical rates for learning in nonlinear dynamical systems. We emphasize that Theorem 2 holds for *any* decision rule $\hat{\pi}$ —it does not require that we use the certainty equivalence decision rule. As Algorithm 1 does rely on certainty equivalence, this result also implies that the certainty equivalence decision rule is optimal for (certain classes of) nonlinear dynamical systems.

The proof of Theorem 2 builds on the work [Wagenmaker et al. \(2021\)](#), which shows a similar result for linear dynamical systems. It critically relies on our quadratic decomposition of the controller loss in Proposition 1, which reduces the problem of obtaining a lower bound on controller loss to a lower bound on estimating A_* in the $\mathcal{H}(A_*)$ norm. Given this, the result can be obtained by applying lower bounds on regression in general norms.

5 Optimal Experiment Design in Arbitrary Dynamical Systems

We turn now to the DYNAMICOED routine, which is the key algorithmic tool we use to prove Theorem 1. DYNAMICOED is a general reduction from policy optimization to optimal experiment design in arbitrary dynamical systems, and is an extension of a recently proposed approach for experiment design in linear MDPs ([Wagenmaker & Jamieson, 2022](#)). This section may be of independent interest.

To illustrate the generality of this reduction, in this section we consider the following system:

$$\mathbf{x}_{h+1} = f_h(\mathbf{x}_h, \mathbf{u}_h, \mathbf{w}_h), \quad h = 1, 2, \dots, H, \quad (5.1)$$

where $\mathbf{x}_h \in \mathcal{X} \subseteq \mathbb{R}^{d_x}$ denotes the state, $\mathbf{u}_h \in \mathcal{U} \subseteq \mathbb{R}^{d_u}$ the input, and $\mathbf{w}_h \in \mathbb{R}^{d_w}$ the noise. We take the dynamics $(f_h)_{h=1}^H$ to be unknown and arbitrary. We assume there is some known featurization of our system that is of interest, $\phi(\mathbf{x}, \mathbf{u}) \rightarrow \mathbb{R}^{d_\phi}$, and an experiment design object on this featurization, $\Phi : \mathbb{R}^{d_\phi \times d_\phi} \rightarrow \mathbb{R}$. Our goal is to collect some set of trajectories $\{\boldsymbol{\tau}_t\}_{t=1}^T$ which minimizes Φ :

$$\Phi\left(\frac{1}{TH} \cdot \sum_{t=1}^T \sum_{h=1}^H \phi(\mathbf{x}_h^t, \mathbf{u}_h^t) \phi(\mathbf{x}_h^t, \mathbf{u}_h^t)^\top\right).$$

As an example, if $\Phi(\Lambda) = \log \det(\Lambda)$, this reduces to D -optimal design, and if $\Phi(\Lambda) = \text{tr}(\mathcal{H} \cdot \Lambda^{-1})$, the setting considered in Section 4, this reduces to weighted A -optimal design. As before, we assume we have access to some set of exploration policies Π_{exp} , and define Λ_π and Ω as in Section 3, but with respect to this new feature map ϕ and system (5.1). We also define $\hat{\Omega}$ to be the space of all possible covariance matrices:

$$\hat{\Omega} := \left\{ \sum_{h=1}^H \phi(\mathbf{x}_h, \mathbf{u}_h) \phi(\mathbf{x}_h, \mathbf{u}_h)^\top : \mathbf{x}_h \in \mathcal{X}, \mathbf{u}_h \in \mathcal{U}, \forall h \in [H] \right\}.$$

To facilitate efficient experiment design in this setting, we will make the following assumption on Φ .

Assumption 7 (Regularity of Φ). Φ is regular in the following sense:

1. Φ is convex, differentiable, and β -smooth in the norm $\|\cdot\|$ (with dual-norm $\|\cdot\|_*$):

$$\|\nabla_{\Lambda}\Phi(\Lambda) - \nabla_{\Lambda'}\Phi(\Lambda')\|_* \leq \beta \cdot \|\Lambda - \Lambda'\|, \quad \forall \Lambda, \Lambda' \in \widehat{\Omega}.$$

2. There exists some $M < \infty$ satisfying $\sup_{\Lambda \in \widehat{\Omega}} \sup_{\mathbf{x} \in \mathcal{X}, \mathbf{u} \in \mathcal{U}} |\phi(\mathbf{x}, \mathbf{u})^\top \nabla_{\Lambda}\Phi(\Lambda) \phi(\mathbf{x}, \mathbf{u})| \leq M$.

The key algorithmic assumption we make is access to a regret minimization oracle on (5.1).

Assumption 8 (Regret Minimization Oracle). Let $\text{cost}_h(\mathbf{x}, \mathbf{u}) = \phi(\mathbf{x}, \mathbf{u})^\top Q_h \phi(\mathbf{x}, \mathbf{u})$ for some $Q_h \in \mathbb{R}^{d_\phi \times d_\phi}$ such that $|\sum_h \text{cost}_h(\mathbf{x}_h, \mathbf{u}_h)| \leq 1$ for all $\mathbf{x}_h \in \mathcal{X}, \mathbf{u}_h \in \mathcal{U}$. We assume we have access to some learner $\mathbb{A}_{\mathcal{R}}$ which is able to achieve low regret on costs $\{\text{cost}_h(\cdot, \cdot)\}_{h=1}^H$ with respect to policy class Π_{exp} . That is, with probability at least $1 - \delta$:

$$\sum_{t=1}^T \mathbb{E}_{f, \pi_t} \left[\sum_{h=1}^H \text{cost}_h(\mathbf{x}_h^t, \mathbf{u}_h^t) \right] - T \cdot \min_{\pi \in \Pi_{\text{exp}}} \mathbb{E}_{f, \pi} \left[\sum_{h=1}^H \text{cost}_h(\mathbf{x}_h, \mathbf{u}_h) \right] \leq C_{\mathcal{R}} \cdot \log^{p_{\mathcal{R}}} \frac{T}{\delta} \cdot T^\alpha$$

for some $C_{\mathcal{R}} > 0$, $p_{\mathcal{R}} > 0$, and $\alpha \in (0, 1)$, and where π_t is the policy $\mathbb{A}_{\mathcal{R}}$ plays at episode t .

Note that the regret minimization algorithm satisfying Assumption 8 may be arbitrary. For example, for linear systems, we could apply provably efficient algorithms for the Linear Quadratic Regulator (Simchowitz & Foster, 2020; Mania et al., 2019); for nonlinear systems of the form (1.1) we could apply the LC³ algorithm of (Kakade et al., 2020); for more general settings of reinforcement learning with function approximation, algorithms such as BILIN-UCB (Du et al., 2021) or E2D (Foster et al., 2021) could be applied. In practice, though they may not formally satisfy the guarantee of Assumption 8, deep RL approaches could be used. We have the following result.

Theorem 3. Fix $T > 0$ and denote $R := \sup_{\Lambda, \Lambda' \in \widehat{\Omega}} \|\Lambda - \Lambda'\|$. Under Assumption 7, and assuming we have access to a learner $\mathbb{A}_{\mathcal{R}}$ satisfying Assumption 8 with $\alpha = 1/2$, DYNAMICOED runs for T episodes on (5.1), and with probability at least $1 - \delta$ collects data $\{(\mathbf{x}_h^t, \mathbf{u}_h^t)\}_{h \in [H], t \in [T]}$ satisfying

$$\Phi \left(\frac{1}{T} \cdot \sum_{t=1}^T \sum_{h=1}^H \phi(\mathbf{x}_h^t, \mathbf{u}_h^t) \phi(\mathbf{x}_h^t, \mathbf{u}_h^t)^\top \right) - \min_{\Lambda \in \Omega} \Phi(\Lambda) \leq \frac{\beta R^2 \log T + HM(C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2T}{\delta} + 3 \log^{1/2} \frac{4T}{\delta})}{T^{1/3}}$$

where $R = \sup_{\Lambda, \Lambda' \in \widehat{\Omega}} \|\Lambda - \Lambda'\|$.

Theorem 3 shows that, given access only to a regret minimization oracle, it is possible to solve experiment design problems on arbitrary dynamical systems. The requirement that $\alpha = 1/2$ is for expositional purposes only—we generalize this result to arbitrary α (and more general feature maps) in Appendix C. Under certain conditions, it can be shown that, if the exploration policies DYNAMICOED runs to collect $\mathfrak{D}$ are rerun, the newly collected data satisfies a similar guarantee as Theorem 3. This lets us run DYNAMICOED to learn an approximate solution of $\min_{\Lambda \in \Omega} \Phi(\Lambda)$, and then rerun the learned policies as many times as desired to collect additional data approximately minimizing Φ .

5.1 Overview of DynamicOED Algorithm

DYNAMICOED is inspired by recent work on experiment design in reinforcement learning (Hazan et al., 2019; Zahavy et al., 2021; Wagenmaker & Jamieson, 2022; Wagenmaker & Pacchiano, 2022),

and can be seen as an extension of the FWREGRET algorithm of Wagenmaker & Jamieson (2022) to arbitrary systems. We refer the reader to Wagenmaker & Jamieson (2022) for a more in-depth discussion of the FWREGRET algorithm, and briefly sketch its extension to arbitrary systems here (see Appendix C and Algorithm 4 for precise definitions).

Algorithm 2 Dynamic Optimal Experiment Design (DYNAMICOED, Informal)

- 1: **input:** objective Φ , episodes T , confidence δ , regret algorithm $\mathbb{A}_{\mathcal{R}}$, exploration policies Π_{exp}
 - 2: Set $K \leftarrow \mathcal{O}(T^{2/3})$, $N \leftarrow \mathcal{O}(T^{1/3})$, $\gamma_n \leftarrow \frac{1}{n+1}$
// $\phi_h^{k,n} := \phi(\mathbf{x}_h^{k,n}, \mathbf{u}_h^{k,n})$ for $(\mathbf{x}_h^{k,n}, \mathbf{u}_h^{k,n})$ the state-input at step h of episode k of iteration n
 - 3: Play any $\pi_{\text{exp}} \in \Pi_{\text{exp}}$ for K episodes, set $\mathbf{\Lambda}_0 \leftarrow \frac{1}{K} \sum_{k=1}^K \sum_{h=1}^H \phi_h^{k,0} (\phi_h^{k,0})^\top$
 - 4: **for** $n = 1, 2, \dots, N$ **do**
 - 5: Compute derivative of $\Phi(\mathbf{\Lambda}_{n-1})$, $\Xi_n \leftarrow \nabla_{\mathbf{\Lambda}} \Phi(\mathbf{\Lambda})|_{\mathbf{\Lambda}=\mathbf{\Lambda}_{n-1}}$
 - 6: Run $\mathbb{A}_{\mathcal{R}}$ on cost $\text{cost}_h^n(\mathbf{x}, \mathbf{u}) \leftarrow \frac{1}{M} \cdot \phi(\mathbf{x}, \mathbf{u})^\top (\Xi_n) \phi(\mathbf{x}, \mathbf{u})$ for K episodes
 - 7: $\mathbf{\Lambda}_n \leftarrow (1 - \gamma_n) \mathbf{\Lambda}_{n-1} + \frac{\gamma_n}{K} \cdot \sum_{k=1}^K \sum_{h=1}^H \phi_h^{k,n} (\phi_h^{k,n})^\top$
 - 8: **return** $\frac{1}{T} \sum_{n=0}^N \sum_{k=1}^K \sum_{h=1}^H \phi_h^{k,n} (\phi_h^{k,n})^\top$
-

Conceptually, DYNAMICOED runs a variant of conditional gradient descent on the objective $\Phi(\mathbf{\Lambda})$. At each iteration, n , it computes the gradient of the loss at the current iterate, $\Xi_n \leftarrow \nabla_{\mathbf{\Lambda}} \Phi(\mathbf{\Lambda})|_{\mathbf{\Lambda}=\mathbf{\Lambda}_{n-1}}$. To run a standard gradient descent algorithm on this objective, we would simply update $\mathbf{\Lambda}_{n-1}$ by taking a step in the direction Ξ_n . However, our objective is to minimize Φ over the constraint set, $\mathbf{\Omega}$. Thus, rather than taking a step in the direction Ξ_n , we wish to take a step in the direction of steepest descent *within the constraint set*.

The challenge is that the constraint set in our setting, $\mathbf{\Omega}$, is *unknown*, as it depends on the expectation over trajectories induced on the unknown dynamics $(f_h)_{h=1}^H$, and therefore we cannot directly compute this steepest descent direction. The key observation is that the computation of this steepest descent direction is equivalent to solving:

$$\arg \min_{\pi_{\text{exp}} \in \Pi_{\text{exp}}} \mathbb{E}_{f, \pi_{\text{exp}}} \left[\sum_{h=1}^H \phi(\mathbf{x}_h, \mathbf{u}_h)^\top (\Xi_n) \phi(\mathbf{x}_h, \mathbf{u}_h) \right].$$

This is simply a policy optimization problem, however, and can be solved approximately by $\mathbb{A}_{\mathcal{R}}$ under Assumption 8. Thus, in the call to $\mathbb{A}_{\mathcal{R}}$ on Line 6, we approximate the steepest descent direction, and on Line 7 update $\mathbf{\Lambda}_{n-1}$ in this direction. Convergence of this procedure to the optimal value, $\min_{\mathbf{\Lambda} \in \mathbf{\Omega}} \Phi(\mathbf{\Lambda})$, can then be shown by the standard analysis of conditional gradient descent. We remark that, under Assumption 7 and Assumption 8, this argument is completely generic and does not require that our system, (5.1), exhibit any additional properties.

5.2 From Theorem 3 to Theorem 1

In Algorithm 1, our goal is to collect covariates, $\check{\mathbf{\Lambda}}_{T_\ell}^{-1}$, such that $\text{tr}(\mathcal{H}(\hat{A}^\ell) \check{\mathbf{\Lambda}}_{T_\ell}^{-1})$ is as small as possible. To achieve this, we apply DYNAMICOED to the objective $\Phi_\ell(\mathbf{\Lambda}) = \text{tr}(\mathcal{H}(\hat{A}^\ell) \mathbf{\Lambda}^{-1})$, with Assumption 8 instantiated by the LC³ algorithm of Kakade et al. (2020). By the guarantee given in Theorem 3, after running for a number of episodes N which scales polynomially in problem parameters, DYNAMICOED will collect covariates $\mathbf{\Lambda}_N$ such that $\Phi_\ell(\frac{1}{N} \mathbf{\Lambda}_N) \leq 2 \cdot \min_{\mathbf{\Lambda} \in \mathbf{\Omega}} \Phi_\ell(\mathbf{\Lambda})$, which implies $\text{tr}(\mathcal{H}(\hat{A}^\ell) \check{\mathbf{\Lambda}}_N^{-1}) \leq \frac{2}{N} \cdot \min_{\mathbf{\Lambda} \in \mathbf{\Omega}} \text{tr}(\mathcal{H}(\hat{A}^\ell) \check{\mathbf{\Lambda}}^{-1})$. By rerunning the policies DYNAMICOED used to

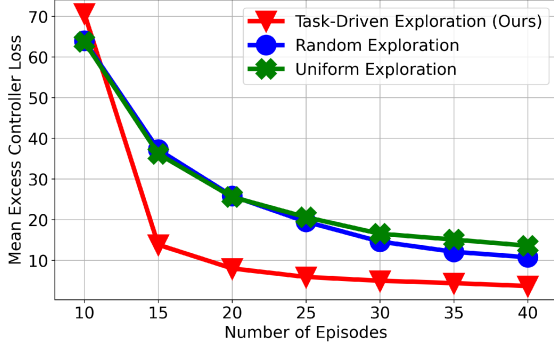


Figure 2: Performance on Drone

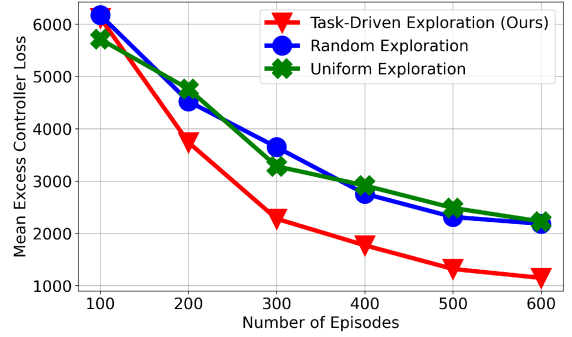


Figure 3: Performance on Car

collect this $\mathbf{A}_N$ for T_ℓ/N additional times, we can ensure $\text{tr}(\mathcal{H}(\hat{A}^\ell)\check{\mathbf{A}}_{T_\ell}^{-1}) \lesssim \frac{1}{T_\ell} \cdot \min_{\mathbf{A} \in \Omega} \text{tr}(\mathcal{H}(\hat{A}^\ell)\check{\mathbf{A}}^{-1})$ as desired.

We remark that the LC^3 algorithm requires access to a computation oracle. As the focus of this work is primarily statistical, we leave addressing this computational challenge for future work. Furthermore, as we show in the following section, computationally efficient, sampling-based implementations of our approach are very effective in practice. We remark as well that the objective we ultimately care about minimizing is $\text{tr}(\mathcal{H}(A_\star)\check{\mathbf{A}}^{-1})$. As we show, by including a small amount of uniform exploration, we can ensure that $\mathcal{H}(\hat{A}^\ell)$ is not too far from $\mathcal{H}(A_\star)$, and so the suboptimality incurred optimizing $\text{tr}(\mathcal{H}(\hat{A}^\ell)\check{\mathbf{A}}^{-1})$ instead of $\text{tr}(\mathcal{H}(A_\star)\check{\mathbf{A}}^{-1})$ only contributes to the lower-order terms of the final guarantee in Theorem 1.

6 Experimental Results

Finally, we demonstrate the effectiveness of our proposed approach (Algorithm 1, the **Task-Driven Exploration** method in Figures 1 to 3) on several systems motivated by robotic applications. We compare Algorithm 1 with an approach that plays $\mathbf{u}_h \sim \mathcal{N}(0, \sigma_u^2 \cdot I)$ (**Gaussian Exploration**), and an approach inspired by Mania et al. (2022) (**Uniform Exploration**), which seeks to estimate $A_\star$ uniformly well, playing inputs that reduce $\|\hat{A} - A_\star\|_{\text{op}}$.

To benchmark the performance of these approaches, we consider an affine system with dynamics corresponding to that of a simplified 3-D drone (i.e., 3-D double integrator with a gravity term), and a nonlinear system with dynamics corresponding to that of a 2-D car. For both systems, we choose $H = 50$, and plot the value of $\mathcal{J}(\hat{\pi}_t; A_\star) - \mathcal{J}(\pi_\star(A_\star); A_\star)$ for $\hat{\pi}_t$ the certainty-equivalence controller computed on the estimate of the system obtained at time t . For the drone, we let $\Pi^\star$ be the class of linear-affine feedback controllers, and for the car, $\Pi^\star$ is a set of nonlinear controllers with dimension 4. While the optimal controller for the drone can be computed in closed-form, for the car we rely on a sampling-based routine to find an approximately optimal controller. The model-task hessian $\mathcal{H}(\hat{A}^\ell)$ is computed via automatic differentiation. For the exploration policies of **Task-Driven Exploration** and **Uniform Exploration**, Π_{exp} , we rely on MPC-style sampling based methods. For all approaches, we require that $\mathbb{E}_{A, \pi_{\text{exp}}} [\sum_{h=1}^H \|\mathbf{u}_h\|_2^2] \leq \gamma^2$ for some $\gamma^2 > 0$ and all $\pi_{\text{exp}} \in \Pi_{\text{exp}}$. On all examples, we implement DYNAMICOED with $\mathbb{A}_{\mathcal{R}}$ a posterior sampling-inspired version of the LC^3 of Kakade et al. (2020). Figures 1 and 3 shows performance averaged over 100 trials, and Figure 2 over 200 trials. Additional experimental details can be found in Appendix G.

As illustrated in Figures 1 to 3, our approach yields a non-trivial gain over existing approaches on all systems. In particular, in Figures 2 and 3 it improves on the sample complexity of existing

approaches by roughly a factor of 2—for example, in the drone system, reaching excess controller cost of 10 after less than 20 episodes, as compared to over 40 episodes for existing approaches.

Our implementation is very modular, and any piece (for example, the parameterization of Π_{exp} and Π^* , the policy optimizer, or the exploration routine) can be easily replaced with other procedures. Our results therefore highlight that, even when using, for example, a possibly suboptimal policy optimizer, exploring so as to minimize uncertainty in the model-task hessian yields a non-trivial gain. We expect that this would hold true regardless of the policy optimizer used—the model-task hessian will adapt to the structure of the policy optimizer, inducing the exploration that will minimize parameter uncertainty most relevant to the given optimizer. Integration of our approach with deep model-based RL approaches is an interesting direction for future work, but we believe the approach will scale to these settings as well.

Acknowledgements

AW would like to thank Kevin Tully for helpful discussions. The work of AW is supported by NSF HDR 62-0221. The work of KJ is supported in part by NSF TRIPODS 2023166 and CIF 2007036.

References

- Abbasi-Yadkori, Y., Pál, D., and Szepesvári, C. Improved algorithms for linear stochastic bandits. *Advances in neural information processing systems*, 24:2312–2320, 2011.
- Agarwal, A., Kakade, S., and Yang, L. F. Model-based reinforcement learning with a generative model is minimax optimal. In *Conference on Learning Theory*, pp. 67–83. PMLR, 2020.
- Åström, K. J. and Wittenmark, B. *Adaptive control*. Courier Corporation, 2013.
- Azar, M. G., Osband, I., and Munos, R. Minimax regret bounds for reinforcement learning. In *International Conference on Machine Learning*, pp. 263–272. PMLR, 2017.
- Boffi, N. M., Tu, S., and Slotine, J.-J. E. Regret bounds for adaptive nonlinear control. In *Learning for Dynamics and Control*, pp. 471–483. PMLR, 2021.
- Bristow, D. A., Tharayil, M., and Alleyne, A. G. A survey of iterative learning control. *IEEE control systems magazine*, 26(3):96–114, 2006.
- Brunke, L., Greeff, M., Hall, A. W., Yuan, Z., Zhou, S., Panerati, J., and Schoellig, A. P. Safe learning in robotics: From learning-based control to safe reinforcement learning. *Annual Review of Control, Robotics, and Autonomous Systems*, 5:411–444, 2022.
- Chua, K., Calandra, R., McAllister, R., and Levine, S. Deep reinforcement learning in a handful of trials using probabilistic dynamics models. *Advances in neural information processing systems*, 31, 2018.
- Cohen, A., Koren, T., and Mansour, Y. Learning linear-quadratic regulators efficiently with only $\sqrt{T}$ regret. *arXiv preprint arXiv:1902.06223*, 2019.
- Dean, S., Mania, H., Matni, N., Recht, B., and Tu, S. On the sample complexity of the linear quadratic regulator. *Foundations of Computational Mathematics*, 20(4):633–679, 2020.
- Dieudonné, J. *Foundations of modern analysis*. Read Books Ltd, 2011.
- Du, S., Kakade, S., Lee, J., Lovett, S., Mahajan, G., Sun, W., and Wang, R. Bilinear classes: A structural framework for provable generalization in rl. In *International Conference on Machine Learning*, pp. 2826–2836. PMLR, 2021.
- Feldbaum, A. A. Dual control theory. i. *Avtomatika i Telemekhanika*, 21(9):1240–1249, 1960.
- Foster, D., Sarkar, T., and Rakhlin, A. Learning nonlinear dynamical systems from a single trajectory. In *Learning for Dynamics and Control*, pp. 851–861. PMLR, 2020.
- Foster, D. J., Kakade, S. M., Qian, J., and Rakhlin, A. The statistical complexity of interactive decision making. *arXiv preprint arXiv:2112.13487*, 2021.
- Gerencsér, L. and Hjalmarsson, H. Adaptive input design in system identification. In *Proceedings of the 44th IEEE Conference on Decision and Control*, pp. 4988–4993. IEEE, 2005.
- Gerencsér, L., Mårtensson, J., and Hjalmarsson, H. Adaptive input design for arx systems. In *2007 European Control Conference (ECC)*, pp. 5707–5714. IEEE, 2007.

- Goodwin, G. C. and Payne, R. L. *Dynamic system identification: experiment design and data analysis*. Academic press, 1977.
- Hazan, E., Kakade, S., Singh, K., and Van Soest, A. Provably efficient maximum entropy exploration. In *International Conference on Machine Learning*, pp. 2681–2691. PMLR, 2019.
- Kaiser, L., Babaeizadeh, M., Milos, P., Osinski, B., Campbell, R. H., Czechowski, K., Erhan, D., Finn, C., Kozakowski, P., Levine, S., et al. Model-based reinforcement learning for atari. *arXiv preprint arXiv:1903.00374*, 2019.
- Kakade, S., Krishnamurthy, A., Lowrey, K., Ohnishi, M., and Sun, W. Information theoretic regret bounds for online nonlinear control. *Advances in Neural Information Processing Systems*, 33: 15312–15325, 2020.
- Katselis, D., Rojas, C. R., Hjalmarsson, H., and Bengtsson, M. Application-oriented finite sample experiment design: A semidefinite relaxation approach. *IFAC Proceedings Volumes*, 45(16): 1635–1640, 2012.
- Lindqvist, K. and Hjalmarsson, H. Identification for control: Adaptive input design using convex optimization. In *Proceedings of the 40th IEEE Conference on Decision and Control (Cat. No. 01CH37228)*, volume 5, pp. 4326–4331. IEEE, 2001.
- Ljung, L. *System identification*. Springer, 1998.
- Manchester, I. R. Input design for system identification via convex relaxation. In *49th IEEE Conference on Decision and Control (CDC)*, pp. 2041–2046. IEEE, 2010.
- Mania, H., Tu, S., and Recht, B. Certainty equivalence is efficient for linear quadratic control. *Advances in Neural Information Processing Systems*, 32, 2019.
- Mania, H., Jordan, M. I., and Recht, B. Active learning for nonlinear system identification with guarantees. *J. Mach. Learn. Res.*, 23:32–1, 2022.
- Mehra, R. Optimal input signals for parameter estimation in dynamic systems—survey and new results. *IEEE Transactions on Automatic Control*, 19(6):753–768, 1974.
- Mesbah, A. Stochastic model predictive control with active uncertainty learning: A survey on dual control. *Annual Reviews in Control*, 45:107–117, 2018.
- Nakka, Y. K., Liu, A., Shi, G., Anandkumar, A., Yue, Y., and Chung, S.-J. Chance-constrained trajectory optimization for safe exploration and learning of nonlinear systems. *IEEE Robotics and Automation Letters*, 6(2):389–396, 2020.
- Nguyen-Tuong, D. and Peters, J. Model learning for robot control: a survey. *Cognitive processing*, 12:319–340, 2011.
- Osband, I. and Van Roy, B. Model-based reinforcement learning and the eluder dimension. *Advances in Neural Information Processing Systems*, 27, 2014.
- Oymak, S. Stochastic gradient descent learns state equations with nonlinear activations. In *conference on Learning Theory*, pp. 2551–2579. PMLR, 2019.

- O’Connell, M., Shi, G., Shi, X., Azizzadenesheli, K., Anandkumar, A., Yue, Y., and Chung, S.-J. Neural-fly enables rapid learning for agile flight in strong winds. *Science Robotics*, 7(66):eabm6597, 2022.
- Rahimi, A. and Recht, B. Uniform approximation of functions with random bases. In *2008 46th annual allerton conference on communication, control, and computing*, pp. 555–561. IEEE, 2008.
- Richards, S. M., Azizan, N., Slotine, J.-J., and Pavone, M. Adaptive-control-oriented meta-learning for nonlinear systems. *arXiv preprint arXiv:2103.04490*, 2021.
- Rojas, C. R., Welsh, J. S., Goodwin, G. C., and Feuer, A. Robust optimal experiment design for system identification. *Automatica*, 43(6):993–1008, 2007.
- Sattar, Y. and Oymak, S. Non-asymptotic and accurate learning of nonlinear dynamical systems. *The Journal of Machine Learning Research*, 23(1):6248–6296, 2022.
- Schulman, J., Levine, S., Abbeel, P., Jordan, M., and Moritz, P. Trust region policy optimization. In *International conference on machine learning*, pp. 1889–1897. PMLR, 2015.
- Schulman, J., Wolski, F., Dhariwal, P., Radford, A., and Klimov, O. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- Shi, G., Shi, X., O’Connell, M., Yu, R., Azizzadenesheli, K., Anandkumar, A., Yue, Y., and Chung, S.-J. Neural lander: Stable drone landing control using learned dynamics. In *2019 International Conference on Robotics and Automation (ICRA)*, pp. 9784–9790. IEEE, 2019.
- Shi, G., Azizzadenesheli, K., O’Connell, M., Chung, S.-J., and Yue, Y. Meta-adaptive nonlinear control: Theory and algorithms. *Advances in Neural Information Processing Systems*, 34:10013–10025, 2021a.
- Shi, G., Hönig, W., Shi, X., Yue, Y., and Chung, S.-J. Neural-swarm2: Planning and control of heterogeneous multirotor swarms using learned interactions. *IEEE Transactions on Robotics*, 38(2):1063–1079, 2021b.
- Simchowitz, M. and Foster, D. Naive exploration is optimal for online lqr. In *International Conference on Machine Learning*, pp. 8937–8948. PMLR, 2020.
- Simchowitz, M., Mania, H., Tu, S., Jordan, M. I., and Recht, B. Learning without mixing: Towards a sharp analysis of linear system identification. *arXiv preprint arXiv:1802.08334*, 2018.
- Simchowitz, M., Boczar, R., and Recht, B. Learning linear dynamical systems with semi-parametric least squares. *arXiv preprint arXiv:1902.00768*, 2019.
- Simchowitz, M., Singh, K., and Hazan, E. Improper learning for non-stochastic control. In *Conference on Learning Theory*, pp. 3320–3436. PMLR, 2020.
- Slotine, J.-J. E., Li, W., et al. *Applied nonlinear control*, volume 199. Prentice hall Englewood Cliffs, NJ, 1991.
- Song, Y. and Sun, W. Pc-mlp: Model-based reinforcement learning with policy cover guided exploration. In *International Conference on Machine Learning*, pp. 9801–9811. PMLR, 2021.

- Sun, W., Jiang, N., Krishnamurthy, A., Agarwal, A., and Langford, J. Model-based rl in contextual decision processes: Pac bounds and exponential improvements over model-free approaches. In *Conference on learning theory*, pp. 2898–2933. PMLR, 2019.
- Sutton, R. S. and Barto, A. G. *Reinforcement learning: An introduction*. MIT press, 2018.
- Wagenmaker, A. and Jamieson, K. Active learning for identification of linear dynamical systems. In *Conference on Learning Theory*, pp. 3487–3582. PMLR, 2020.
- Wagenmaker, A. and Jamieson, K. Instance-dependent near-optimal policy identification in linear mdps via online experiment design. *arXiv preprint arXiv:2207.02575*, 2022.
- Wagenmaker, A. and Pacchiano, A. Leveraging offline data in online reinforcement learning. *arXiv preprint arXiv:2211.04974*, 2022.
- Wagenmaker, A. J., Simchowitz, M., and Jamieson, K. Task-optimal exploration in linear dynamical systems. In *International Conference on Machine Learning*, pp. 10641–10652. PMLR, 2021.
- Williams, G., Wagener, N., Goldfain, B., Drews, P., Reh, J. M., Boots, B., and Theodorou, E. A. Information theoretic mpc for model-based reinforcement learning. In *2017 IEEE International Conference on Robotics and Automation (ICRA)*, pp. 1714–1721. IEEE, 2017.
- Yu, C., Shi, G., Chung, S.-J., Yue, Y., and Wierman, A. The power of predictions in online control. *Advances in Neural Information Processing Systems*, 33:1994–2004, 2020a.
- Yu, T., Thomas, G., Yu, L., Ermon, S., Zou, J. Y., Levine, S., Finn, C., and Ma, T. Mopo: Model-based offline policy optimization. *Advances in Neural Information Processing Systems*, 33: 14129–14142, 2020b.
- Zahavy, T., O’Donoghue, B., Desjardins, G., and Singh, S. Reward is enough for convex mdps. *Advances in Neural Information Processing Systems*, 34:25746–25759, 2021.
- Zanette, A. and Brunskill, E. Tighter problem-dependent regret bounds in reinforcement learning without domain knowledge using value function bounds. In *International Conference on Machine Learning*, pp. 7304–7312. PMLR, 2019.
- Zhou, D., Gu, Q., and Szepesvari, C. Nearly minimax optimal reinforcement learning for linear mixture markov decision processes. In *Conference on Learning Theory*, pp. 4532–4576. PMLR, 2021.

A Technical Tools

Lemma A.1. *Let $\mathbf{w}_i \sim \mathcal{N}(0, I_{d_x})$ for all $i \in \{1, 2, \dots, n\}$. Then,*

$$\mathbb{E} \left[\left(\sum_{i=1}^n \|\mathbf{w}_i\|_2 \right)^c \right] \leq n^2 \cdot \text{poly}(d_x).$$

for c an absolute constant.

Proof. We first bound

$$\left(\sum_{i=1}^n \|\mathbf{w}_i\|_2 \right)^c \leq n \cdot \max_i \|\mathbf{w}_i\|_2^c \leq n \cdot \sum_i \|\mathbf{w}_i\|_2^c.$$

The result then follows since we can bound the $\mathbb{E}[\|\mathbf{w}_i\|_2^c] \leq \text{poly}(d)$ for $\mathbf{w}_i \sim \mathcal{N}(0, I)$ and c an absolute constant. $\square$

Lemma A.2 (Lemma I.4 of [Wagenmaker et al. \(2021\)](#)). *Assume $A, B \in \mathbb{S}_{++}^d$, $\|A - B\|_{\text{op}} \leq \epsilon$, and $\epsilon < \lambda_{\min}(B)$. Then*

$$\|A^{-1} - B^{-1}\|_{\text{op}} \leq \frac{\epsilon}{\lambda_{\min}(B)(\lambda_{\min}(B) - \epsilon)}.$$

A.1 Martingale Regression in General Norms

For the following two results, we consider the martingale regression setting of [Wagenmaker et al. \(2021\)](#) (referred to as the MDM setting). In particular, we consider observations of the form

$$y_t = \langle \boldsymbol{\mu}^*, \mathbf{z}_t \rangle + w_t, \tag{A.1}$$

for $y_t \in \mathbb{R}$, unknown parameter $\boldsymbol{\mu}^* \in \mathbb{R}^{d_\mu}$, $w_t \mid \mathcal{F}_{t-1} \sim \mathcal{N}(0, \sigma_w^2)$, and $\mathbf{z}_t$ $\mathcal{F}_{t-1}$ -measurable, for a filtration $(\mathcal{F}_t)_{t \geq 1}$. This setting therefore encompasses general stochastic processes where the observations are linear—the evolution of $\mathbf{z}_t$ could be arbitrary.

We consider the setting where we interact with (A.1) for T steps, collecting observations $\{(\mathbf{z}_t, y_t)\}_{t=1}^T$, and then form the least-squares estimate of $\boldsymbol{\mu}^*$:

$$\hat{\boldsymbol{\mu}} = \left(\sum_{t=1}^T \mathbf{z}_t \mathbf{z}_t^\top \right)^{-1} \sum_{t=1}^T \mathbf{z}_t y_t.$$

We also denote $\boldsymbol{\Sigma}_T := \sum_{t=1}^T \mathbf{z}_t \mathbf{z}_t^\top$. The following results characterize the estimation error of $\hat{\boldsymbol{\mu}}$ in the M -norm and 2-norm.

Proposition 3 (Theorem 7.2 of [Wagenmaker et al. \(2021\)](#)). *Fix any matrices $\boldsymbol{\Gamma} \in \mathbb{S}_{++}^{d_\mu}$, $M \in \mathbb{S}_+^{d_\mu}$, with $M \neq 0$. Given a parameter $\beta \in (0, 1/4)$, define the event*

$$\mathcal{E} := \{\|\boldsymbol{\Sigma}_T - \boldsymbol{\Gamma}\|_{\text{op}} \leq \beta \lambda_{\min}(\boldsymbol{\Gamma})\}.$$

Then, if $\mathcal{E}$ holds, the following holds with probability at least $1 - \delta$:

$$\|\hat{\boldsymbol{\mu}} - \boldsymbol{\mu}^*\|_M^2 \leq 5(1 + \zeta) \cdot \sigma_w^2 \log \frac{6d_\mu}{\delta} \cdot \text{tr}(M\boldsymbol{\Gamma}^{-1})$$

where $\zeta = 26\beta^2 \lambda_{\max}(\boldsymbol{\Gamma}) \text{tr}(\boldsymbol{\Gamma}^{-1})$.

Proposition 4 (Lemma E.1 of Wagenmaker et al. (2021)). *On the event*

$$\mathcal{E}_{\text{op}} := \{\lambda_{\min}(\Sigma_{\mathbf{T}}) \geq \underline{\lambda}T, \Sigma_{\mathbf{T}} \preceq T\bar{\Gamma}_{\mathbf{T}}\},$$

then we have that with probability at least $1 - \delta$:

$$\|\hat{\boldsymbol{\mu}} - \boldsymbol{\mu}^{\star}\|_2 \leq C \cdot \sigma_{\mathbf{w}} \sqrt{\frac{\log 1/\delta + d_{\boldsymbol{\mu}} + \log \det(\bar{\Gamma}_{\mathbf{T}}/\underline{\lambda} + I)}{\underline{\lambda}T}}.$$

A.1.1 Connection Between (1.1) and (A.1)

We will apply the results Proposition 3 and Proposition 4 in the setting of (1.1) in order to obtain estimation bounds on $A_{\star}$. As the setting of (1.1) has vector observations, we briefly describe here how it can be mapped into the setting described above.

Recall that (1.1) evolves as

$$\mathbf{x}_{h+1} = A_{\star}\phi(\mathbf{x}_h, \mathbf{u}_h) + \mathbf{w}_h, \quad h = 1, \dots, H,$$

for $\mathbf{x}_h \in \mathbb{R}^{d_{\mathbf{x}}}$, $\phi(\mathbf{x}, \mathbf{u}) \in \mathbb{R}^{d_{\phi}}$, and $A_{\star} \in \mathbb{R}^{d_{\mathbf{x}} \times d_{\phi}}$. We assume that $\mathbf{x}_1$ is some fixed starting state. Assume that we have run for T episodes, and collected observations $\{(\mathbf{x}_1^t, \mathbf{u}_1^t, \mathbf{x}_2^t, \dots, \mathbf{x}_h^t, \mathbf{u}_h^t, \mathbf{x}_{h+1}^t)\}_{t=1}^T$. Now let $\boldsymbol{\mu}^{\star} := \text{vec}(A_{\star})$. Furthermore, for any t, h , and $i \in [d_{\mathbf{x}}]$, let $\mathbf{t} = (t, h, i)$ and $\mathbf{z}_{\mathbf{t}} = [\mathbf{0}_{d_{\phi}(i-1)}, \phi(\mathbf{x}_h^t, \mathbf{u}_h^t), \mathbf{0}_{d_{\phi}(d_{\mathbf{x}}-i)}] \in \mathbb{R}^{d_{\mathbf{x}}d_{\phi}}$ where $\mathbf{0}_d$ denotes the zero vector of length d . Then we see that

$$[\mathbf{x}_{h+1}^t]_i = \langle \boldsymbol{\mu}^{\star}, \mathbf{z}_{\mathbf{t}} \rangle + [\mathbf{w}_h^t]_i.$$

Setting $y_{\mathbf{t}} = [\mathbf{x}_{h+1}^t]_i$ and $w_{\mathbf{t}} = [\mathbf{w}_h^t]_i$, it is clear that this follows the observation model of (A.1) with $d_{\boldsymbol{\mu}} = d_{\phi}d_{\mathbf{x}}$ and $T = d_{\mathbf{x}}TH$. It is also straightforward to see that the measurability assumptions of the setting of (A.1) are satisfied by this.

B Proof of Main Result

Theorem 4 (Full Version of Theorem 1). *Assume Assumptions 1 to 5 and 13 hold. Then if*

$$T \geq \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, H, B_A, B_{\phi}, L_{\phi}, L_{\theta}, L_{\text{cost}}, L_{\pi_{\star}}, \sigma_{\mathbf{w}}, \sigma_{\mathbf{w}}^{-1}, \frac{1}{\lambda_{\min}^{\star}}, \log \frac{T}{\delta} \right) \cdot \max \left\{ 1, \frac{1}{r_{\text{cost}}(A_{\star})^2}, \frac{1}{r_{\theta}(A_{\star})^2} \right\}, \quad (\text{B.1})$$

with probability at least $1 - \delta$, Algorithm 3 plays exploration policies $\pi_{\text{exp}} \in \Pi_{\text{exp}}$ at every episode, runs for at most T episodes, and the controller $\hat{\pi}_T$ returned Algorithm 3 satisfies, with probability at least $1 - \delta$:

$$\mathcal{J}(\hat{\pi}_T; A_{\star}) - \mathcal{J}(\pi_{\star}(A_{\star}); A_{\star}) \leq \frac{\sigma_{\mathbf{w}}^2}{T} \cdot \min_{\Lambda \in \Omega} \text{tr}(\mathcal{H}(A_{\star})\check{\Lambda}^{-1}) \cdot C \log \frac{6d_{\mathbf{x}}d_{\phi}}{\delta} + \frac{C_{\text{lot}}}{T^{3/2}}$$

for C a universal constant and

$$C_{\text{poly}} := \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, H, B_A, B_{\phi}, L_{\phi}, L_{\theta}, L_{\text{cost}}, L_{\pi_{\star}}, \sigma_{\mathbf{w}}, \sigma_{\mathbf{w}}^{-1}, \frac{1}{\lambda_{\min}^{\star}}, \log \frac{T}{\delta} \right).$$

Algorithm 3 Optimal Exploration in Nonlinear Systems (Full Version of Algorithm 1)

- 1: **inputs:** number of episodes to run T , cost function $(\text{cost}_h)_{h=1}^H$, confidence δ , control policies Π^* , exploration policies Π_{exp}
- 2: $\hat{A}^1 \leftarrow$ anything, $\ell_T \leftarrow \lceil \log_2 T/8 \rceil$
- 3: **for** $\ell = 1, 2, 3, \dots, \ell_T$ **do**
- 4: $T_\ell \leftarrow 2^\ell, \delta_\ell \leftarrow \delta/8\ell^2$
- 5: Compute estimate of cost matrix: $\mathcal{H}_\ell \leftarrow \mathcal{H}(\hat{A}^\ell)$
- 6: $\Pi_\ell \leftarrow \text{LEARNEXP}\Pi(\mathcal{H}_\ell, T_\ell, \delta_\ell, \mathbb{A}_\mathcal{R}, \Pi_{\text{exp}})$ (Algorithm 7), with $\mathbb{A}_\mathcal{R}$ the LC³ algorithm (Kakade et al., 2020)
- 7: Rerun each policy in Π_ℓ $N_\ell = \lceil T_\ell/|\Pi_\ell| \rceil$ times, denote collected data $\mathfrak{D}_\ell$
- 8: Estimate system parameters
- 9:

$$\hat{A}^{\ell+1} = \arg \min_A \sum_{h=1}^H \sum_{(\mathbf{x}_{h+1}, \mathbf{u}_h, \mathbf{x}_h) \in \mathfrak{D}_\ell} \|\mathbf{x}_{h+1} - A\phi(\mathbf{x}_h, \mathbf{u}_h)\|_2^2$$

10: **return** $\hat{\pi}_T \leftarrow \pi_*(\hat{A}^{\ell_T+1})$

Proof. Let $\mathcal{E}_\ell$ denote that the good event of Lemma B.3 holds at round ℓ which, by Lemma B.3, occurs with probability at least $1 - 6\delta_\ell$. By our setting of $\delta_\ell = \delta/12\ell^2$, we have that the total failure probability of $\mathcal{E}_\ell$ for all ℓ is bounded as

$$\sum_{\ell=1}^{\infty} 6 \cdot \frac{\delta}{12\ell^2} \leq \delta.$$

Henceforth we assume that $\mathcal{E} := \cap_\ell \mathcal{E}_\ell$ holds. Let $\hat{A} := \hat{A}^{\ell_T+1}$, and $\hat{A}^- := \hat{A}^{\ell_T}$.

Before proceeding to the main proof, we note that the conclusion that Algorithm 3 only explores with policies in Π_{exp} follows from the definition of LEARNEXP Π and LC³. Note that LEARNEXP Π only interacts with (1.1) through calls to DYNAMICOED, which itself only interacts with (1.1) through calls to $\mathbb{A}_\mathcal{R}$, instantiated in Algorithm 3 by LC³. Inspection of the LC³ algorithm in Kakade et al. (2020) reveals that LC³ only interacts with (1.1) by playing policies in Π_{exp} , from which the conclusion follows.

Bounding the Number of Episodes. Denote $T_\ell^{\text{oad}} = |\Pi_\ell|$. Note that by construction we always have $N_\ell T_\ell^{\text{oad}} \geq T_\ell$. By Lemma B.2, as long as (B.1) is met, we can bound the total number of episodes collected up to and including round ℓ by $4T_\ell$ for $\ell \in \{\ell_T, \ell_T - 1\}$. We therefore, in the following, will make use of the fact that

$$c \cdot T \leq N_\ell K_\ell \leq c' \cdot T$$

for $\ell \in \{\ell_T, \ell_T - 1\}$ and absolute constants c, c' . Furthermore, it also follows from this that the total number of episodes run by Algorithm 3 is bounded by

$$4T_{\ell_T} = 4 \cdot 2^{\lceil \log T/8 \rceil} \leq 8 \cdot \frac{T}{8} = T,$$

so Algorithm 3 runs for at most T episodes.

Approximating the Controller Loss. Let $r_{\text{est}}(A_\star) := \min\{1, r_{\text{cost}}(A_\star), r_\theta(A_\star)\}$, for $r_{\text{cost}}(A_\star)$ as in Assumption 2 and $r_\theta(A_\star)$ as in Assumption 13. By Lemma D.2, under Assumptions 1, 2, 4, 5 and 13, as long as $\hat{A} \in \mathcal{B}_F(A_\star; r_{\text{est}}(A_\star))$, we have

$$\begin{aligned} \mathcal{J}(\hat{\pi}_T; A_\star) - \mathcal{J}(\pi_\star(A_\star); A_\star) &\leq \|\text{vec}(\hat{A} - A_\star)\|_{\mathcal{H}(A_\star)}^2 \\ &\quad + \text{poly}(L_{\pi_\star}, \|A_\star\|_{\text{op}}, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x) \cdot \|\hat{A} - A_\star\|_{\text{op}}^3. \end{aligned}$$

Furthermore, we can bound

$$\begin{aligned} \|\text{vec}(\hat{A} - A_\star)\|_{\mathcal{H}(A_\star)}^2 &= \|\text{vec}(\hat{A} - A_\star)\|_{\mathcal{H}(\hat{A}^-)}^2 + \text{vec}(\hat{A} - A_\star)^\top (\mathcal{H}(A_\star) - \mathcal{H}(\hat{A}^-)) \text{vec}(\hat{A} - A_\star) \\ &\leq (\hat{A} - A_\star)^\top \mathcal{H}(\hat{A}^-) (\hat{A} - A_\star) + \|\hat{A}^- - A_\star\|_{\text{op}}^2 \|\mathcal{H}(A_\star) - \mathcal{H}(\hat{A}^-)\|_{\text{op}}. \end{aligned}$$

Bounding the Hessian Estimation Error. On $\mathcal{E}$, by Lemma B.3 and as long as (B.1) is met, we have (note that at the final epoch, the plug-in estimator $\mathcal{H}(\hat{A}^-)$ is given as input to DYNAMICOED):

$$\begin{aligned} \|\text{vec}(\hat{A} - A_\star)\|_{\mathcal{H}(\hat{A}^-)}^2 &\leq \frac{60}{N_{\ell_T} T_{\ell_T}^{\text{oad}}} \cdot \min_{\mathbf{A} \in \Omega} \text{tr} \left(\mathcal{H}(\hat{A}^-) \check{\mathbf{A}}^{-1} \right) \cdot \sigma_w^2 \log \frac{6d_x d_\phi}{\delta} \\ &\quad + \text{poly} \left(d_\phi, d_x, \frac{1}{\lambda_{\min}^\star}, B_A, B_\phi, \log \frac{1}{\sigma_w}, H, \|\mathcal{H}(\hat{A}^-)\|_{\text{op}}, \log \frac{T_{\ell_T}}{\delta} \right) \cdot \frac{1}{N_{\ell_T}^2} \\ &\leq \frac{C}{T} \cdot \min_{\mathbf{A} \in \Omega} \text{tr} \left(\mathcal{H}(\hat{A}^-) \check{\mathbf{A}}^{-1} \right) \cdot \sigma_w^2 \log \frac{6d_x d_\phi}{\delta} \\ &\quad + \text{poly} \left(d_\phi, d_x, \frac{1}{\lambda_{\min}^\star}, B_A, B_\phi, \log \frac{1}{\sigma_w}, H, \|\mathcal{H}(\hat{A}^-)\|_{\text{op}}, \log \frac{T}{\delta} \right) \cdot \frac{1}{T^2} \end{aligned}$$

where the last line uses that $N_{\ell_T} T_{\ell_T}^{\text{oad}}$ is within a constant of T , and that $T_{\ell_T}^{\text{oad}}$ can be bounded by

$$\text{poly} \left(d_\phi, d_x, \frac{1}{\lambda_{\min}^\star}, B_A, B_\phi, \log \frac{1}{\sigma_w}, H, \log \frac{T}{\delta} \right)$$

by Lemma B.3. By Lemma B.4 we can bound

$$\min_{\mathbf{A} \in \Omega} \text{tr} \left(\mathcal{H}(\hat{A}^-) \check{\mathbf{A}}^{-1} \right) \leq \min_{\mathbf{A} \in \Omega} 2\text{tr} \left(\mathcal{H}(A_\star) \check{\mathbf{A}}^{-1} \right) + \frac{2d_x d_\phi}{\lambda_{\min}^\star} \cdot \|\mathcal{H}(A_\star) - \mathcal{H}(\hat{A}^-)\|_{\text{op}}$$

and by Lemma D.3, under Assumptions 1, 2, 4, 5 and 13, and as long as $\hat{A}^- \in \mathcal{B}_F(A_\star; r_{\text{est}}(A_\star))$, we can bound

$$\|\mathcal{H}(A_\star) - \mathcal{H}(\hat{A}^-)\|_{\text{op}} \leq \text{poly}(L_{\pi_\star}, \|A_\star\|_{\text{op}}, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x) \cdot \|\hat{A}^- - A_\star\|_{\text{op}}.$$

Let

$$C_{\text{lot}} := \text{poly} \left(L_{\pi_\star}, B_A, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, d_x, d_\phi, \frac{1}{\lambda_{\min}^\star}, \sigma_w, \sigma_w^{-1}, H, \|\mathcal{H}(A_\star)\|_{\text{op}}, \log \frac{T}{\delta} \right)$$

denote some lower-order constant, whose precise polynomial dependence may change from line to line. On $\mathcal{E}$, by Lemma B.3 we can bound

$$\|\hat{A} - A_\star\|_{\text{op}}, \|\hat{A}^- - A_\star\|_{\text{op}} \leq C_{\text{lot}} \cdot \frac{1}{\sqrt{T}}, \quad (\text{B.2})$$

and so assuming the burn-in (B.1) is met, we can bound $\|\hat{A}^- - A_\star\|_{\text{op}} \leq \frac{1}{2}r_{\text{est}}(A_\star)$ and $\|\hat{A} - A_\star\|_{\text{op}} \leq \frac{1}{2}r_{\text{est}}(A_\star)$. This then implies that

$$\|\mathcal{H}(A_\star) - \mathcal{H}(\hat{A}^-)\|_{\text{op}} \leq C_{\text{lot}} \cdot \frac{1}{\sqrt{T}},$$

so in particular we can bound

$$\begin{aligned} & \text{poly} \left(d_\phi, d_x, \frac{1}{\lambda_{\min}^\star}, B_A, B_\phi, H, \log \frac{1}{\sigma_w}, \|\mathcal{H}(\hat{A}^-)\|_{\text{op}}, \log \frac{T}{\delta} \right) \\ & \leq \text{poly} \left(d_\phi, d_x, \frac{1}{\lambda_{\min}^\star}, B_A, B_\phi, H, \log \frac{1}{\sigma_w}, \|\mathcal{H}(A_\star)\|_{\text{op}}, \log \frac{T}{\delta} \right). \end{aligned}$$

We have therefore shown that

$$\begin{aligned} \mathcal{J}(\hat{\pi}_T; A_\star) - \mathcal{J}(\pi_\star(A_\star); A_\star) & \leq \frac{C}{T} \cdot \min_{\mathbf{A} \in \Omega} \text{tr}(\mathcal{H}(A_\star) \check{\mathbf{A}}^{-1}) \cdot \sigma_w^2 \log \frac{6d_x d_\phi}{\delta} \\ & \quad + C_{\text{lot}} \cdot \left(\|\hat{A} - A_\star\|_{\text{op}}^3 + \|\hat{A}^- - A_\star\|_{\text{op}}^3 + \frac{1}{T^2} + \frac{1}{T} \cdot \|\hat{A}^- - A_\star\|_{\text{op}} \right) \\ & \leq \frac{C}{T} \cdot \min_{\mathbf{A} \in \Omega} \text{tr}(\mathcal{H}(A_\star) \check{\mathbf{A}}^{-1}) \cdot \sigma_w^2 \log \frac{6d_x d_\phi}{\delta} + \frac{C_{\text{lot}}}{T^{3/2}}, \end{aligned}$$

The final result follows from using Lemma D.4 to bound

$$\|\mathcal{H}(A_\star)\|_{\text{op}} \leq \text{poly}(\|A_\star\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, L_{\pi_\star}, \sigma_w^{-1}, H, d_x).$$

□

Proof of Theorem 1. The proof of Theorem 1 is identical to that of Theorem 4, the only difference being that we replace Assumption 13 with Assumption 6. However, by Proposition 6, the conditions of Assumption 13 are met when Assumption 6 holds. □

B.1 Supporting Lemmas

Lemma B.1. *Under Assumptions 1 and 3, the system (1.1) satisfies Assumptions 11 and 12 with*

$$\psi(\boldsymbol{\tau}) \leftarrow I_{d_x} \otimes \sum_{h=1}^H \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau) \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau)^\top, \quad D = d_x H B_\phi^2,$$

$d_\psi \leftarrow d_\phi d_x$, and where $\mathbf{x}_h^\tau$ (resp. $\mathbf{u}_h^\tau$) denotes the state (resp. input) at step h of trajectory $\boldsymbol{\tau}$. Furthermore, it satisfies Assumption 10 with $\mathbb{A}_\mathcal{R}$ instantiated with the LC³ algorithm of Kakade et al. (2020) and

$$C_\mathcal{R} = C \cdot H \sqrt{d_\phi(d_\phi + d_x + B_A) \cdot \log(1 + B_\phi H / \sigma_w)}, \quad p_\mathcal{R} = 3/2, \quad \alpha = 1/2$$

for a universal constant C .

Proof. That Assumption 12 is satisfied is immediate under Assumption 3. It is clear that $\psi(\boldsymbol{\tau}) \in \mathcal{S}_+^{d_\phi d_x}$. To obtain a bound on D , we only need to bound the trace of $\psi(\boldsymbol{\tau})$:

$$\text{tr}(\psi(\boldsymbol{\tau})) = \sum_{h=1}^H \text{tr}(I_{d_x}) \cdot \text{tr}(\phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau) \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau)^\top)$$

$$\begin{aligned}
&= d_{\mathbf{x}} \sum_{h=1}^H \|\phi(\mathbf{x}_h^{\tau}, \mathbf{u}_h^{\tau})\|_2^2 \\
&\leq d_{\mathbf{x}} H B_{\phi}^2
\end{aligned}$$

where the inequality holds under Assumption 1.

To show that Assumption 10 is satisfied in this setting, we have by Theorem 6 that with probability at least $1 - \delta$, LC^3 has regret bounded as (using that $c_{\max} \leq 1$ in the setting of Assumption 10):

$$\mathcal{R}_T \leq C \cdot H \sqrt{d_{\phi} \cdot (d_{\phi} + d_{\mathbf{x}} + B_A + \log \frac{1}{\delta}) \cdot T \cdot \log(1 + B_{\phi} H T / \sigma_{\mathbf{w}})}$$

for C a universal constant. We can therefore take $\alpha = 1/2$, $p_{\mathcal{R}} = 3/2$, and

$$C_{\mathcal{R}} = C' \cdot H \sqrt{d_{\phi}(d_{\phi} + d_{\mathbf{x}} + B_A) \cdot \log(1 + B_{\phi} H / \sigma_{\mathbf{w}})}.$$

□

Lemma B.2. *Let $\bar{T}_{\ell}$ denote the total number of episodes collected by Algorithm 3 at round ℓ . For*

$$T_{\ell} \geq \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_{\phi}, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{T_{\ell}}{\delta} \right), \quad (\text{B.3})$$

on the success event of Lemma B.3, we have $2T_{\ell} \geq \bar{T}_{\ell}$ and

$$2T_{\ell} \geq \sum_{i=1}^{\ell-1} \bar{T}_i.$$

Proof. Recall that $T_{\ell} = 2^{\ell}$ and $N_{\ell} = \lceil T_{\ell} / T_{\ell}^{\text{od}} \rceil$ for $T_{\ell}^{\text{od}} = |\Pi_{\ell}|$. By Lemma B.3, $\bar{T}_{\ell}$ can be bounded as

$$\bar{T}_{\ell} \leq N_{\ell} T_{\ell}^{\text{od}} + (16 + 2 \log T_{\ell}^{\text{od}}) T_{\ell}^{\text{od}}$$

and T_{ℓ}^{od} can be bounded as

$$T_{\ell}^{\text{od}} \leq \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_{\phi}, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{T_{\ell}}{\delta} \right). \quad (\text{B.4})$$

Note that we can bound

$$\begin{aligned}
\bar{T}_i &\leq N_i T_i^{\text{od}} + (16 + 2 \log T_i^{\text{od}}) T_i^{\text{od}} \\
&\leq T_i + (17 + 2 \log T_i^{\text{od}}) T_i^{\text{od}} \\
&\leq T_i + \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_{\phi}, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{T_i}{\delta} \right).
\end{aligned}$$

From this it is immediately obvious that $2T_{\ell} \geq \bar{T}_{\ell}$ as long as (B.3) is satisfied.

To show the second conclusion note that, by our choice of $T_i = 2^i$, we have that $T_{\ell} \geq \sum_{i=1}^{\ell-1} T_i$, so it therefore remains to show that

$$T_{\ell} \geq \sum_{i=1}^{\ell-1} \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_{\phi}, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{T_i}{\delta} \right).$$

However, we can bound

$$\sum_{i=1}^{\ell-1} \text{poly} \left(d_\phi, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_\phi, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{T_i}{\delta} \right) \leq \text{poly} \left(d_\phi, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_\phi, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{T_\ell}{\delta} \right) \cdot \log T_\ell,$$

so a sufficient condition is

$$T_\ell \geq \text{poly} \left(d_\phi, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_\phi, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{T_\ell}{\delta} \right) \cdot \log T_\ell$$

which we see is met when (B.3) holds. $\square$

Lemma B.3. *Consider running Algorithm 7 with weight matrix $\mathcal{H}$, parameter $\tilde{N}$, and confidence δ , and rerunning each policy in Π_{out} $N \leq \tilde{N}$ times. Then, under Assumptions 1 and 3, with probability at least $1 - 6\delta$:*

$$\begin{aligned} \|\text{vec}(\hat{A} - A_\star)\|_{\mathcal{H}}^2 &\leq \frac{60}{NT_{\text{out}}} \cdot \min_{\mathbf{\Lambda} \in \Omega} \text{tr}(\mathcal{H}\check{\mathbf{\Lambda}}^{-1}) \cdot \sigma_{\mathbf{w}}^2 \log \frac{6d_{\mathbf{x}}d_\phi}{\delta} \\ &\quad + \text{poly} \left(d_\phi, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_\phi, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \|\mathcal{H}\|_{\text{op}}, \log \frac{\tilde{N}}{\delta} \right) \cdot \frac{1}{N^2} \end{aligned}$$

where $\hat{A}$ denotes the least-squares estimate of $A_\star$ obtained on the data generated by rerunning Π_{out} . In addition, we have

$$\|\hat{A} - A_\star\|_{\text{F}} \leq \text{poly} \left(d_\phi, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_\phi, \sigma_{\mathbf{w}}, H, \log \frac{\tilde{N}}{\delta} \right) \cdot \frac{1}{\sqrt{N}}.$$

Furthermore, we have

$$T_{\text{out}} \leq \text{poly} \left(d_\phi, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^*}, B_A, B_\phi, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{\tilde{N}}{\delta} \right)$$

and the total number of episodes collected by this procedure is bounded by $NT_{\text{out}} + (16 + 2\log(T_{\text{out}})) \cdot T_{\text{out}}$, for $T_{\text{out}} = |\Pi_{\text{out}}|$.

Proof. By Lemma B.1, the assumptions of Lemma C.6 and Lemma C.7 are met, so we can therefore apply these results in our setting. By Lemma C.6, the event $\mathcal{E}_{\text{exp}}$ occurs with probability at least $1 - \delta$. Throughout the remainder of the proof we union bound over the success event of Lemma C.7 and $\mathcal{E}_{\text{exp}}$, which together occur with probability at least $1 - 4\delta$.

Let

$$\tilde{\mathbf{\Lambda}} := \sum_{t=1}^{NT_{\text{out}}} \psi(\boldsymbol{\tau}^t) = I_{d_{\mathbf{x}}} \otimes \sum_{t=1}^{NT_{\text{out}}} \sum_{h=1}^H \phi(\mathbf{x}_h^t, \mathbf{u}_h^t) \phi(\mathbf{x}_h^t, \mathbf{u}_h^t)^\top$$

denote the features returned by rerunning every policy in Π_{out} N times. By Lemma C.7, we then have that:

$$\left\| \tilde{\mathbf{\Lambda}} - N \cdot \sum_{\pi \in \Pi_{\text{out}}} \check{\mathbf{\Lambda}}_\pi \right\|_{\text{op}} \leq \underbrace{\frac{\sqrt{T_{\text{out}}} \cdot \sqrt{8d_\phi d_{\mathbf{x}} \log(1 + 8\sqrt{NT_{\text{out}}})} + 8\log 1/\delta}{\sqrt{N} \cdot 6272d_\phi d_{\mathbf{x}} \log \frac{68\tilde{N}}{\delta}}}_{=:\beta} \cdot \lambda_{\min} \left(N \cdot \sum_{\pi \in \Pi_{\text{out}}} \check{\mathbf{\Lambda}}_\pi \right).$$

Applying Proposition 3 with $\mathcal{E}$ the event that the above conclusion holds and $\mathbf{\Gamma} := N \cdot \sum_{\pi \in \Pi_{\text{out}}} \check{\mathbf{\Lambda}}_{\pi}$, we obtain that, with probability at least $1 - \delta$ (using the mapping to the martingale regression setting described in Appendix A.1.1):

$$\|\text{vec}(\hat{A} - A_{\star})\|_{\mathcal{H}}^2 \leq 5(1 + \zeta) \cdot \sigma_{\mathbf{w}}^2 \log \frac{6d_{\mathbf{x}}d_{\phi}}{\delta} \cdot \text{tr}(\mathcal{H}\mathbf{\Gamma}^{-1})$$

for $\zeta = 26\beta^2\lambda_{\max}(\mathbf{\Gamma})\text{tr}(\mathbf{\Gamma}^{-1})$ and β as defined above. Since $\|\phi(\mathbf{x}, \mathbf{u})\|_2 \leq B_{\phi}$ under Assumption 1, we have $\|\check{\mathbf{\Lambda}}_{\pi}\|_2 \leq B_{\phi}^2$, so we can upper bound $\lambda_{\max}(\mathbf{\Gamma}) \leq NT_{\text{out}}B_{\phi}^2$. By Lemma C.7, we can also bound (using that $D = d_{\mathbf{x}}HB_{\phi}^2$ by Lemma B.1):

$$\text{tr}(\mathbf{\Gamma}^{-1}) \leq \frac{1}{N \cdot 6272d_{\mathbf{x}}HB_{\phi}^2 \log \frac{68\tilde{N}}{\delta}}.$$

Combining these and using that

$$T_{\text{out}} \leq \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^{\star}}, B_A, B_{\phi}, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{\tilde{N}}{\delta} \right) \quad (\text{B.5})$$

as shown in Lemma C.8 (and using our bounds on $C_{\mathcal{R}}$ and $p_{\mathcal{R}}$ in Lemma B.1), we can therefore bound

$$\zeta \leq \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^{\star}}, B_A, B_{\phi}, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \log \frac{\tilde{N}}{\delta} \right) \cdot \frac{1}{N}.$$

Using that $\text{tr}(\mathcal{H}\mathbf{\Gamma}^{-1}) \leq \|\mathcal{H}\|_{\text{op}} \cdot \text{tr}(\mathbf{\Gamma}^{-1})$, and the bound on $\text{tr}(\mathbf{\Gamma}^{-1})$ given above, it follows that

$$\|\text{vec}(\hat{A} - A_{\star})\|_{\mathcal{H}}^2 \leq 5\sigma_{\mathbf{w}}^2 \log \frac{6d_{\mathbf{x}}d_{\phi}}{\delta} \cdot \text{tr}(\mathcal{H}\mathbf{\Gamma}^{-1}) + \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^{\star}}, B_A, B_{\phi}, \log \frac{1}{\sigma_{\mathbf{w}}}, H, \|\mathcal{H}\|_{\text{op}}, \log \frac{\tilde{N}}{\delta} \right) \cdot \frac{1}{N^2}.$$

Finally, by Lemma C.7, we can bound

$$\text{tr}(\mathcal{H}\mathbf{\Gamma}^{-1}) \leq \frac{12}{NT_{\text{out}}} \cdot \min_{\mathbf{\Lambda} \in \Omega} \text{tr}(\mathcal{H}\check{\mathbf{\Lambda}}^{-1}).$$

By Lemma C.8, we can bound the total number of episodes collected by Algorithm 7 by $(16 + 2\log(T_{\text{out}})) \cdot T_{\text{out}}$.

Bound on Frobenius Norm Error. By Lemma C.7, we can lower bound

$$\lambda_{\min}(\tilde{\mathbf{\Lambda}}) \geq N \cdot 6272d_{\mathbf{x}}d_{\phi}HB_{\phi}^2 \log \frac{68\tilde{N}}{\delta}.$$

Furthermore, since $\|\phi(\mathbf{x}, \mathbf{u})\|_2 \leq B_{\phi}$, we always have $\|\tilde{\mathbf{\Lambda}}\|_{\text{op}} \leq NT_{\text{out}}B_{\phi}^2$, which implies $\tilde{\mathbf{\Lambda}} \preceq NT_{\text{out}}B_{\phi}^2 \cdot I$. By Proposition 4, we then have that with probability at least $1 - \delta$ (again using the mapping to the martingale regression setting described in Appendix A.1.1):

$$\begin{aligned} \|\hat{A} - A_{\star}\|_{\text{F}} &\leq C \cdot \sigma_{\mathbf{w}} \sqrt{\frac{\log 1/\delta + d_{\mathbf{x}}d_{\phi} + \log \det\left(\frac{T_{\text{out}}}{6272d_{\mathbf{x}}d_{\phi}H \log \frac{68\tilde{N}}{\delta}} \cdot I + I\right)}{N \cdot 6272d_{\mathbf{x}}d_{\phi}HB_{\phi}^2 \log \frac{68\tilde{N}}{\delta}}} \\ &\leq \text{poly} \left(d_{\phi}, d_{\mathbf{x}}, \frac{1}{\lambda_{\min}^{\star}}, B_A, B_{\phi}, \sigma_{\mathbf{w}}, H, \log \frac{\tilde{N}}{\delta} \right) \cdot \frac{1}{\sqrt{N}}. \end{aligned}$$

□

Lemma B.4. Under Assumption 3, for any $\mathcal{H}, \mathcal{H}'$, we can bound

$$\min_{\check{\mathbf{\Lambda}} \in \Omega} \text{tr}(\mathcal{H}\check{\mathbf{\Lambda}}^{-1}) \leq \min_{\check{\mathbf{\Lambda}} \in \Omega} 2\text{tr}(\mathcal{H}'\check{\mathbf{\Lambda}}^{-1}) + \frac{2d_{\mathbf{x}}d_{\phi}}{\lambda_{\min}^*} \cdot \|\mathcal{H} - \mathcal{H}'\|_{\text{op}}.$$

Proof. We have

$$\begin{aligned} \min_{\check{\mathbf{\Lambda}} \in \Omega} \text{tr}(\mathcal{H}\check{\mathbf{\Lambda}}^{-1}) &= \min_{\check{\mathbf{\Lambda}} \in \Omega} \text{tr}(\mathcal{H}'\check{\mathbf{\Lambda}}^{-1}) + \text{tr}((\mathcal{H} - \mathcal{H}')\check{\mathbf{\Lambda}}^{-1}) \\ &\leq \min_{\check{\mathbf{\Lambda}} \in \Omega} \text{tr}(\mathcal{H}'\check{\mathbf{\Lambda}}^{-1}) + \|\mathcal{H} - \mathcal{H}'\|_{\text{op}} \cdot \text{tr}(\check{\mathbf{\Lambda}}^{-1}) \end{aligned}$$

Under Assumption 3, we know that there exists some $\check{\mathbf{\Lambda}}' \in \Omega$ such that $\lambda_{\min}(\check{\mathbf{\Lambda}}') \geq \lambda_{\min}^*$. We can then bound

$$\begin{aligned} &\min_{\check{\mathbf{\Lambda}} \in \Omega} \text{tr}(\mathcal{H}'\check{\mathbf{\Lambda}}^{-1}) + \|\mathcal{H} - \mathcal{H}'\|_{\text{op}} \cdot \text{tr}(\check{\mathbf{\Lambda}}^{-1}) \\ &\leq \min_{\check{\mathbf{\Lambda}} \in \Omega} \text{tr}(\mathcal{H}'(\frac{1}{2}\check{\mathbf{\Lambda}} + \frac{1}{2}\check{\mathbf{\Lambda}}')^{-1}) + \|\mathcal{H} - \mathcal{H}'\|_{\text{op}} \cdot \text{tr}((\frac{1}{2}\check{\mathbf{\Lambda}} + \frac{1}{2}\check{\mathbf{\Lambda}}')^{-1}) \\ &\leq \min_{\check{\mathbf{\Lambda}} \in \Omega} 2\text{tr}(\mathcal{H}'\check{\mathbf{\Lambda}}^{-1}) + 2\|\mathcal{H} - \mathcal{H}'\|_{\text{op}} \cdot \frac{d_{\mathbf{x}}d_{\phi}}{\lambda_{\min}^*} \end{aligned}$$

which proves the result. $\square$

C Experiment Design in Arbitrary Dynamical Systems

In this section we generalize somewhat the setting of Section 5. In particular, our goal will now be to collect some set of trajectories $\mathfrak{D} = \{\boldsymbol{\tau}_t\}_{t=1}^T$, which minimize

$$\Phi \left(\frac{1}{T} \sum_{t=1}^T \psi(\boldsymbol{\tau}_t) \right)$$

for some general feature mapping $\psi : \mathcal{T} \rightarrow \mathbb{R}^d$, $\Phi : \mathbb{R}^d \rightarrow \mathbb{R}$, and $\mathcal{T} = (\mathcal{X} \times \mathcal{U})^H \times \mathcal{X}$ the space of possible state-input trajectories, $\boldsymbol{\tau} = (\mathbf{x}_0, \mathbf{u}_0, \mathbf{x}_1, \dots, \mathbf{x}_{H-1}, \mathbf{u}_{H-1}, \mathbf{x}_H) \in \mathcal{T}$. We will assume that ψ can be decomposed additively as

$$\psi(\boldsymbol{\tau}) = \sum_{h=1}^H \psi_h(\mathbf{x}_h, \mathbf{u}_h).$$

In Section 5 we considered the special case where $\psi_h(\mathbf{x}, \mathbf{u}) = \phi(\mathbf{x}, \mathbf{u})\phi(\mathbf{x}, \mathbf{u})^\top$; in this section ψ could instead be any arbitrary mapping.

As before, we will be interested in defining optimal exploration with respect to some set of exploration policies, Π_{exp} . Let

$$\Omega_{\psi} := \{\mathbb{E}_{\pi \sim \omega}[\mathbb{E}_{\pi}[\psi(\boldsymbol{\tau})]] : \omega \in \Delta_{\Pi}\}$$

denote the space of expected value of $\psi(\boldsymbol{\tau})$ for mixtures of policies in Π_{exp} . To distinguish elements $\mathbf{\Lambda} \in \Omega$ from elements in Ω_{ψ} , we will let $\mathbf{\Gamma} = \mathbb{E}_{\pi \sim \omega}[\mathbb{E}_{\pi}[\psi(\boldsymbol{\tau})]]$ refer to elements of Ω_{ψ} , and in particular define $\mathbf{\Gamma}_{\pi} := \mathbb{E}_{\pi}[\psi(\boldsymbol{\tau})]$ (where it is assumed that the expectation is collected over trajectories on

(5.1)). We will usually denote unnormalized sums of features, e.g. $\sum_{t=1}^T \psi(\tau_t)$, with Σ . We also define $\hat{\Omega}_\psi$ to be the space of all possible combinations of $\psi(\tau)$:

$$\hat{\Omega}_\psi := \{\mathbb{E}_{\tau \sim \omega}[\psi(\tau)] : \omega \in \Delta_{\mathcal{T}}\}.$$

We generalize Assumption 7 and Assumption 8 as follows.

Assumption 9 (Regularity of Φ). *We make the following assumptions:*

1. Φ is convex, differentiable, and β -smooth in the norm $\|\cdot\|$:

$$\|\nabla_{\Gamma}\Phi(\Gamma) - \nabla_{\Gamma'}\Phi(\Gamma')\|_* \leq \beta \cdot \|\Gamma - \Gamma'\|, \quad \forall \Gamma, \Gamma' \in \hat{\Omega}_\psi$$

for $\|\cdot\|_*$ the dual norm of $\|\cdot\|$.

2. There exists some $M < \infty$ satisfying

$$\sup_{\Gamma \in \hat{\Omega}_\psi} \sup_{\tau \in \mathcal{T}} |\langle \nabla_{\Gamma}\Phi(\Gamma), \psi(\tau) \rangle| \leq M.$$

Assumption 10 (Regret Minimization Oracle). *Let $\text{cost}_h(\tau) = \langle Q_h, \psi_h(\tau) \rangle$ for some $Q_h \in \mathbb{R}^d$, and $\text{cost}(\tau) = \sum_{h=1}^H \text{cost}_h(\tau)$ the total cost of trajectory τ . We assume we have access to some learner $\mathbb{A}_{\mathcal{R}}$ which, in the setting when $|\text{cost}(\tau)| \leq 1$ for all $\tau \in \mathcal{T}$, is able to achieve low regret on $\{\text{cost}_h(\cdot, \cdot)\}_{h=1}^H$ with respect to policy class Π_{exp} . That is, with probability at least $1 - \delta$:*

$$\sum_{t=1}^T \mathbb{E}_{f, \pi_t}[\text{cost}(\tau_t)] - T \cdot \inf_{\pi \in \Pi_{\text{exp}}} \mathbb{E}_{f, \pi}[\text{cost}(\tau)] \leq C_{\mathcal{R}} \cdot \log^{p_{\mathcal{R}}} \frac{T}{\delta} \cdot T^{\alpha}$$

for some $C_{\mathcal{R}} > 0$, $p_{\mathcal{R}} > 0$, and $\alpha \in (0, 1)$, and where π_t is the policy $\mathbb{A}_{\mathcal{R}}$ plays at episode t .

Algorithm 4 Dynamic Optimal Experiment Design (DYNAMICOEED)

- 1: **input:** objective Φ , number of episodes T (OR number of iterates N , episodes per iterate K), confidence δ , regret minimization algorithm $\mathbb{A}_{\mathcal{R}}$, exploration policies Π_{exp}
- 2: Play any policy $\pi_{\text{exp}} \in \Pi_{\text{exp}}$ for K episodes, collect trajectories $\mathfrak{D}_0 = \{\tau_k^0\}_{k=1}^K$, set $\Gamma_0 \leftarrow K^{-1} \sum_{k=1}^K \psi(\tau_k^0)$
- 3: **for** $n = 1, 2, \dots, N$ **do**
- 4: Set $\gamma_n \leftarrow \frac{1}{n+1}$
- 5: Run $\mathbb{A}_{\mathcal{R}}$ on cost

$$\text{cost}_h^n(\tau) \leftarrow \frac{1}{M} \langle \Xi_n, \psi_h(x_h, u_h) \rangle \quad \text{for } \Xi_n \leftarrow \nabla_{\Gamma}\Phi(\Gamma)|_{\Gamma=\Gamma_{n-1}}$$

for K episodes, collect trajectories $\mathfrak{D}_n = \{\tau_k^n\}_{k=1}^K$, denote policies run as Π_n

- 6: $\Gamma_n \leftarrow (1 - \gamma_n)\Gamma_{n-1} + \gamma_n K^{-1} \sum_{k=1}^K \psi(\tau_k^n)$
 - 7: **return** $(N+1)K\Gamma_N, \cup_{n=0}^N \Pi_n, \cup_{n=0}^N \mathfrak{D}_n$
-

We define DYNAMICOEED as in Algorithm 4. We then have the following generalization of Theorem 3.

Theorem 5 (Full Version of Theorem 3). *Let Assumption 7 hold, and assume that we have access to a learner $\mathbb{A}_{\mathcal{R}}$ satisfying Assumption 8. Fix $N, K > 0$. Then, with probability at least $1 - \delta$, DYNAMICOED runs for at most $(N + 1)K$ episodes, and collects a dataset satisfying $\mathfrak{D} = \{\{\boldsymbol{\tau}_k^n\}_{k=1}^K\}_{n=0}^N$ satisfying*

$$\Phi\left(\frac{1}{K(N+1)} \sum_{n=0}^N \sum_{k=1}^K \boldsymbol{\psi}(\boldsymbol{\tau}_k^n)\right) - \min_{\boldsymbol{\Gamma} \in \boldsymbol{\Omega}_{\boldsymbol{\psi}}} \Phi(\boldsymbol{\Gamma}) \leq \frac{\beta R^2 (\log N + 1)}{2(N+1)} + M \cdot \left(C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2NK}{\delta} \cdot K^{\alpha-1} + \sqrt{\frac{8 \log(4N/\delta)}{K}} \right)$$

where $R = \sup_{\boldsymbol{\Gamma}, \boldsymbol{\Gamma}' \in \hat{\boldsymbol{\Omega}}_{\boldsymbol{\psi}}} \|\boldsymbol{\Gamma} - \boldsymbol{\Gamma}'\|$.

In this work we are particularly interested in the case where $\boldsymbol{\psi}(\boldsymbol{\tau}) \in \mathcal{S}_+^{d_{\boldsymbol{\psi}}}$. We encapsulate this in the following assumption.

Assumption 11 (Matrix Experiment Design). *We assume that $\boldsymbol{\psi}(\boldsymbol{\tau}) \in \mathcal{S}_+^{d_{\boldsymbol{\psi}}}$ and that, for all $\boldsymbol{\tau} \in \mathcal{T}$, $\text{tr}(\boldsymbol{\psi}(\boldsymbol{\tau})) \leq D$ for some $D > 0$.*

The following corollary instantiates Theorem 3 under Assumption 11 with objective $\Phi(\boldsymbol{\Gamma}) = \text{tr}(\mathcal{H}(\boldsymbol{\Gamma} + \boldsymbol{\Gamma}_0)^{-1})$, the objective considered in Algorithm 1.

Corollary 1. *Consider the objective*

$$\Phi(\boldsymbol{\Gamma}) = \text{tr}(\mathcal{H} \cdot (\boldsymbol{\Gamma} + \boldsymbol{\Gamma}_0)^{-1})$$

and assume that $\mathcal{H} \succeq 0$ and Assumption 10 holds with $\alpha = 1/2$ and Assumption 11 holds. Fix N, K , let $T := (N + 1)K$, and consider running Algorithm 4 on this objective and with these choices of N and K . Then Algorithm 4 will run for at most T episodes, and, with probability at least $1 - \delta$, will return data satisfying

$$\begin{aligned} \text{tr}\left(\mathcal{H}\left(\sum_{t=1}^T \boldsymbol{\psi}(\boldsymbol{\tau}_t) + T\boldsymbol{\Gamma}_0\right)^{-1}\right) &\leq \frac{1}{T} \cdot \min_{\boldsymbol{\Gamma} \in \boldsymbol{\Omega}_{\boldsymbol{\psi}}} \text{tr}(\mathcal{H}(\boldsymbol{\Gamma} + \boldsymbol{\Gamma}_0)^{-1}) + \frac{8D^4 \|\mathcal{H}\|_{\text{op}} \|\boldsymbol{\Gamma}_0^{-1}\|_{\text{op}}^3}{T(N+1)} \\ &\quad + \frac{8D \|\mathcal{H}\|_{\text{op}} \|\boldsymbol{\Gamma}_0^{-1}\|_{\text{op}}^2 (\log^{1/2} \frac{4T}{\delta} + C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2T}{\delta})}{T\sqrt{K}} \end{aligned}$$

C.1 Proof of Theorem 3 and Theorem 5

Lemma C.1 (Lemma C.1 of Wagenmaker & Jamieson (2022)). *Consider running Algorithm 5 with some convex function f that is β -smooth with respect to some norm $\|\cdot\|$, assume that $\mathbf{y}_n \in \mathcal{Y}$ for some $\mathcal{Y}$ and all n , and let $R := \sup_{\mathbf{z}, \mathbf{y} \in \mathcal{Z} \cup \mathcal{Y}} \|\mathbf{z} - \mathbf{y}\|$. Then for $N \geq 2$, we have*

$$f(\mathbf{z}_{N+1}) - \min_{\mathbf{z} \in \mathcal{Z}} f(\mathbf{z}) \leq \frac{\beta R^2 (\log N + 1)}{2(N+1)} + \frac{1}{N+1} \sum_{n=1}^N \epsilon_n.$$

Lemma C.2 (Lemma C.2 of Wagenmaker & Jamieson (2022)). *When running Algorithm 5, we have*

$$\mathbf{z}_{N+1} = \frac{1}{N+1} \left(\sum_{n=1}^N \mathbf{y}_n + \mathbf{z}_1 \right).$$

Algorithm 5 Approximate Frank-Wolfe

- 1: **input:** function to optimize f , number of iterations to run N , starting iterate $\mathbf{x}_1$
- 2: **for** $t = 1, 2, \dots, N$ **do**
- 3: Set $\gamma_n \leftarrow \frac{1}{n+1}$
- 4: Choose $\mathbf{y}_n$ to be any point such that

$$\nabla f(\mathbf{z}_n)^\top \mathbf{y}_n \leq \min_{\mathbf{y} \in \mathcal{Z}} \nabla f(\mathbf{z}_n)^\top \mathbf{y} + \epsilon_n$$

- 5: $\mathbf{z}_{n+1} \leftarrow (1 - \gamma_n)\mathbf{z}_n + \gamma_n \mathbf{y}_n$
 - 6: **return** $\mathbf{x}_{N+1}$
-

Proof of Theorem 3. By our assumption on $\mathbb{A}_{\mathcal{R}}$, Assumption 8, we have that, at round n , with probability at least $1 - \delta/2N$,

$$\sum_{k=1}^K \mathbb{E}_{\pi_k} [\text{cost}^n(\boldsymbol{\tau}_k)] - K \cdot \inf_{\pi \in \Pi_{\text{exp}}} \mathbb{E}_{\pi} [\text{cost}^n(\boldsymbol{\tau})] \leq C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2NK}{\delta} \cdot K^{\alpha}$$

where we have used that, under Assumption 9 and by the definition of $\text{cost}_h^n(\boldsymbol{\tau})$, $|\text{cost}^n(\boldsymbol{\tau})| \leq 1$ for all $\boldsymbol{\tau} \in \mathcal{T}$. This implies that

$$\frac{1}{K} \sum_{k=1}^K \mathbb{E}_{\pi_k} [\langle \Xi_n, \boldsymbol{\psi}(\boldsymbol{\tau}_k) \rangle] \leq M \cdot \inf_{\pi \in \Pi_{\text{exp}}} \mathbb{E}_{\pi} [\text{cost}^n(\boldsymbol{\tau})] + MC_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2NK}{\delta} \cdot K^{\alpha-1}.$$

Furthermore, by Azuma-Hoeffding and under Assumption 9, we have that, with probability at least $1 - \delta/2N$,

$$\left| \frac{1}{K} \sum_{k=1}^K \langle \Xi_n, \boldsymbol{\psi}(\boldsymbol{\tau}_k) \rangle - \frac{1}{K} \sum_{k=1}^K \mathbb{E}_{\pi_k} [\langle \Xi_n, \boldsymbol{\psi}(\boldsymbol{\tau}_k) \rangle] \right| \leq \sqrt{\frac{8M^2 \log(4N/\delta)}{K}}.$$

This implies that

$$\frac{1}{K} \sum_{k=1}^K \langle \Xi_n, \boldsymbol{\psi}(\boldsymbol{\tau}_k) \rangle \leq M \cdot \inf_{\pi \in \Pi_{\text{exp}}} \mathbb{E}_{\pi} [\text{cost}^n(\boldsymbol{\tau})] + M \cdot \left(C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{K}{\delta} \cdot K^{\alpha-1} + \sqrt{\frac{8 \log(4N/\delta)}{K}} \right). \quad (\text{C.1})$$

Note that

$$\langle \Xi_n, \boldsymbol{\psi}(\boldsymbol{\tau}_k) \rangle = \langle \nabla_{\mathbf{\Gamma}} \Phi(\mathbf{\Gamma})|_{\mathbf{\Gamma}=\mathbf{\Gamma}_n}, \boldsymbol{\psi}(\boldsymbol{\tau}) \rangle,$$

and that for any $\mathbf{\Gamma} \in \boldsymbol{\Omega}_{\psi}$, we have

$$\langle \nabla_{\mathbf{\Gamma}} \Phi(\mathbf{\Gamma})|_{\mathbf{\Gamma}=\mathbf{\Gamma}_n}, \mathbf{\Gamma} \rangle = \mathbb{E}_{\pi \sim \omega} [\mathbb{E}_{\pi} [\langle \nabla_{\mathbf{\Gamma}} \Phi(\mathbf{\Gamma})|_{\mathbf{\Gamma}=\mathbf{\Gamma}_n}, \boldsymbol{\psi}(\boldsymbol{\tau}) \rangle]]$$

for some ω . This implies that

$$\begin{aligned} \inf_{\mathbf{\Gamma} \in \boldsymbol{\Omega}_{\psi}} \langle \nabla_{\mathbf{\Gamma}} \Phi(\mathbf{\Gamma})|_{\mathbf{\Gamma}=\mathbf{\Gamma}_n}, \mathbf{\Gamma} \rangle &= \inf_{\omega \in \Delta_{\Pi}} \mathbb{E}_{\pi \sim \omega} [\mathbb{E}_{\pi} [\langle \nabla_{\mathbf{\Gamma}} \Phi(\mathbf{\Gamma})|_{\mathbf{\Gamma}=\mathbf{\Gamma}_n}, \boldsymbol{\psi}(\boldsymbol{\tau}) \rangle]] \\ &= \inf_{\pi \in \Pi_{\text{exp}}} \mathbb{E}_{\pi} [\langle \Xi_n, \boldsymbol{\psi}(\boldsymbol{\tau}) \rangle] \end{aligned}$$

$$= M \cdot \inf_{\pi \in \Pi_{\text{exp}}} \mathbb{E}_{\pi}[\text{cost}^n(\boldsymbol{\tau})].$$

By (C.1) above, we have that

$$\frac{1}{K} \sum_{k=1}^K \boldsymbol{\psi}(\boldsymbol{\tau}_k)$$

is an approximate minimizer of $M \cdot \sup_{\pi \in \Pi_{\text{exp}}} \mathbb{E}_{\pi}[\text{cost}^n(\boldsymbol{\tau})]$, with approximation tolerance $M(C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2NK}{\delta} \cdot K^{\alpha-1} + \sqrt{\frac{8 \log(4N/\delta)}{K}})$. We can therefore apply Lemma C.1 with

$$\epsilon_n = M \cdot \left(C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{K}{\delta} \cdot K^{\alpha-1} + \sqrt{\frac{8 \log(4N/\delta)}{K}} \right)$$

to get that

$$\Phi(\mathbf{\Gamma}_{N+1}) - \min_{\mathbf{\Gamma} \in \mathbf{\Omega}_{\psi}} \Phi(\mathbf{\Gamma}) \leq \frac{\beta R^2 (\log N + 1)}{2(N+1)} + M \cdot \left(C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2NK}{\delta} \cdot K^{\alpha-1} + \sqrt{\frac{8 \log(4N/\delta)}{K}} \right).$$

The result then follows since $\mathbf{\Gamma}_{N+1} = \frac{1}{K(N+1)} \sum_{n=0}^N \sum_{k=1}^K \boldsymbol{\psi}(\boldsymbol{\tau}_k^n)$ by Lemma C.2. $\square$

Proof of Corollary 1. By Theorem 5, for any setting of N and K , we have that with probability at least $1 - \delta$:

$$\begin{aligned} & \text{tr} \left(\mathcal{H} \left(\frac{1}{K(N+1)} \sum_{n=0}^N \sum_{k=1}^K \boldsymbol{\psi}(\boldsymbol{\tau}_k^n) + \mathbf{\Gamma}_0 \right)^{-1} \right) - \min_{\mathbf{\Gamma} \in \mathbf{\Omega}} \text{tr} (\mathcal{H}(\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1}) \\ & \leq \frac{\beta R^2 \log N}{N+1} + \frac{MC_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2NK}{\delta}}{K^{1-\alpha}} + M \sqrt{\frac{8 \log \frac{4N}{\delta}}{K}} \end{aligned}$$

which implies

$$\begin{aligned} & \text{tr} \left(\mathcal{H} \left(\sum_{n=0}^N \sum_{k=1}^K \boldsymbol{\psi}(\boldsymbol{\tau}_k^n) + T\mathbf{\Gamma}_0 \right)^{-1} \right) - \frac{\min_{\mathbf{\Gamma} \in \mathbf{\Omega}} \text{tr} (\mathcal{H}(\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1})}{T} \\ & \leq \frac{\beta R^2 \log N}{T(N+1)} + \frac{MC_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2NK}{\delta}}{T\sqrt{K}} + M \frac{1}{T} \sqrt{\frac{8 \log \frac{4N}{\delta}}{K}} \end{aligned}$$

This gives

$$\begin{aligned} \text{tr} \left(\mathcal{H} \left(\sum_{n=0}^N \sum_{k=1}^K \boldsymbol{\psi}(\boldsymbol{\tau}_k^n) + T\mathbf{\Gamma}_0 \right)^{-1} \right) & \leq \frac{\min_{\mathbf{\Gamma} \in \mathbf{\Omega}} \text{tr} (\mathcal{H}(\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1})}{T} \\ & \quad + \frac{\beta R^2 \log T}{T(N+1)} + \frac{M(3 \log^{1/2} \frac{4T}{\delta} + C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2T}{\delta})}{T\sqrt{K}}. \end{aligned}$$

It then remains to bound R, β , and M . By Lemma D.6 of [Wagenmaker & Jamieson \(2022\)](#), we have that

$$\nabla_{\mathbf{\Gamma}} \Phi(\mathbf{\Gamma})[\tilde{\mathbf{\Gamma}}] = -\text{tr} \left(\mathcal{H}(\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \tilde{\mathbf{\Gamma}} (\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \right).$$

We can then compute the second derivative as, using Lemma D.6 of [Wagenmaker & Jamieson \(2022\)](#):

$$\begin{aligned} \nabla_{\mathbf{\Gamma}}^2 \Phi(\mathbf{\Gamma})[\tilde{\mathbf{\Gamma}}, \bar{\mathbf{\Gamma}}] &= \frac{d}{dt} \left[-\text{tr} \left(\mathcal{H}(\mathbf{\Gamma} + \mathbf{\Gamma}_0 + t\bar{\mathbf{\Gamma}})^{-1} \tilde{\mathbf{\Gamma}} (\mathbf{\Gamma} + \mathbf{\Gamma}_0 + t\bar{\mathbf{\Gamma}})^{-1} \right) \right] \\ &= \text{tr} \left(\mathcal{H}(\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \bar{\mathbf{\Gamma}} (\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \tilde{\mathbf{\Gamma}} (\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \right) \\ &\quad + \text{tr} \left(\mathcal{H}(\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \tilde{\mathbf{\Gamma}} (\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \bar{\mathbf{\Gamma}} (\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \right). \end{aligned}$$

Recall that M is any bound on

$$\sup_{\mathbf{\Gamma} \in \hat{\Omega}_{\psi}} \sup_{\boldsymbol{\tau} \in \mathcal{T}} |\langle \nabla_{\mathbf{\Gamma}} \Phi(\mathbf{\Gamma}), \boldsymbol{\psi}(\boldsymbol{\tau}) \rangle|.$$

By the above computation of the gradient, we can bound this as

$$\begin{aligned} \sup_{\mathbf{\Gamma} \in \hat{\Omega}_{\psi}} \sup_{\boldsymbol{\tau} \in \mathcal{T}} |\langle \nabla_{\mathbf{\Gamma}} \Phi(\mathbf{\Gamma}), \boldsymbol{\psi}(\boldsymbol{\tau}) \rangle| &\leq \sup_{\mathbf{\Gamma} \in \hat{\Omega}_{\psi}} \sup_{\boldsymbol{\tau} \in \mathcal{T}} \left| \text{tr} \left(\mathcal{H}(\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \boldsymbol{\psi}(\boldsymbol{\tau}) (\mathbf{\Gamma} + \mathbf{\Gamma}_0)^{-1} \right) \right| \\ &\leq \|\mathcal{H}\|_{\text{op}} \|\mathbf{\Gamma}_0^{-1}\|_{\text{op}}^2 \cdot \sup_{\boldsymbol{\tau} \in \mathcal{T}} \text{tr}(\boldsymbol{\psi}(\boldsymbol{\tau})) \\ &\leq D \|\mathcal{H}\|_{\text{op}} \|\mathbf{\Gamma}_0^{-1}\|_{\text{op}}^2. \end{aligned} \tag{C.2}$$

To bound β , by the Mean Value Theorem it suffices to bound the operator norm of $\nabla_{\mathbf{\Gamma}}^2 \Phi(\mathbf{\Gamma})$. Using the expression above, we can bound this as

$$\begin{aligned} \sup_{\tilde{\mathbf{\Gamma}}, \bar{\mathbf{\Gamma}} \in \hat{\Omega}_{\psi}} |\nabla_{\mathbf{\Gamma}}^2 \Phi(\mathbf{\Gamma})[\boldsymbol{\psi}(\boldsymbol{\tau}_1), \boldsymbol{\psi}(\boldsymbol{\tau}_2)]| &\leq 2 \|\mathcal{H}\|_{\text{op}} \|\mathbf{\Gamma}_0^{-1}\|_{\text{op}}^3 \cdot \sup_{\tilde{\mathbf{\Gamma}}, \bar{\mathbf{\Gamma}} \in \hat{\Omega}_{\psi}} \text{tr}(\tilde{\mathbf{\Gamma}} \bar{\mathbf{\Gamma}}) \\ &\leq 2D^2 \|\mathcal{H}\|_{\text{op}} \|\mathbf{\Gamma}_0^{-1}\|_{\text{op}}^3. \end{aligned}$$

Finally, it's straightforward to bound $R \leq 2D$. Putting all of this together gives the result. $\square$

C.2 Collecting Full-Rank Data

Algorithm 6 Minimum Eigenvalue Maximization (MINEIG)

- 1: **input:** scale N , confidence δ , regret minimization algorithm $\mathbb{A}_{\mathcal{R}}$, exploration policies Π_{exp}
 - 2: **for** $j = 1, 2, 3, \dots$ **do**
 - 3: $N_j \leftarrow \lceil 2^{j/3} \rceil - 1, K_j \leftarrow \lceil 2^{2j/3} \rceil, T_j \leftarrow (N_j + 1)K_j, \lambda_j \leftarrow T_j^{-1/18}, \delta_j \leftarrow \frac{\delta}{4j^2}$
 - 4: $\Sigma_j, \Pi_j \leftarrow \text{DYNAMICOEED}(\Phi, N_j, K_j, \delta_j, \mathbb{A}_{\mathcal{R}}, \Pi_{\text{exp}})$ for $\Phi(\mathbf{\Gamma}) = \text{tr}((\mathbf{\Gamma} + \lambda_j \cdot I)^{-1})$
 - 5: **if** $\lambda_{\min}(\Sigma_j) \geq 12544Dd_{\psi} \log \frac{2N(2+32T_j)}{\delta}$ **then**
 - 6: **break**
 - 7: **return** Π_j
-

In this section, we consider the setting where $\boldsymbol{\psi}(\boldsymbol{\tau}) \in \mathcal{S}_+^{d_{\psi}}$, and our goal is to collect $\{\boldsymbol{\tau}_t\}_{t=1}^T$ such that $\boldsymbol{\psi}(\frac{1}{T} \sum_{t=1}^T \boldsymbol{\psi}(\boldsymbol{\tau}_t)) > 0$. For this to be achievable, we need the following assumption, a generalization of Assumption 3.

Assumption 12 (Full-Rank Data). Consider $\psi(\tau)$ such that $\psi(\tau) \in \mathbb{S}_+^{d_\psi}$. Then we have $\sup_{\Gamma \in \Omega_\psi} \lambda_{\min}(\Gamma) \geq \lambda_{\min}^*$ for some $\lambda_{\min}^* > 0$.

Throughout this section we also assume that Assumption 10 is satisfied with $\alpha = 1/2$ (though all results generalize in a straightforward way for $\alpha \neq 1/2$). We have the following result.

Lemma C.3. Under Assumptions 10 to 12, running Algorithm 6 we have that with probability at least $1 - \delta$, it will terminate after collecting at most

$$\text{poly} \left(d_\psi, \frac{1}{\lambda_{\min}^*}, D, C_{\mathcal{R}}, \log^{p_{\mathcal{R}}} \frac{N}{\delta} \right)$$

episodes, and return policy set Π such that

$$\lambda_{\min} \left(\sum_{\pi \in \Pi} \Gamma_\pi \right) \geq 6272 D d_\psi \log \frac{68N}{\delta}.$$

Furthermore, if we rerun each policy in Π once, the resulting features Σ will satisfy, with probability at least $1 - \delta/N$:

$$\lambda_{\min}(\Sigma) \geq 6272 D d_\psi \log \frac{68N}{\delta}.$$

Proof. By Lemma C.4 and our choice of N_j and K_j in Algorithm 6, we have that if $\lambda_j \leq \frac{\lambda_{\min}^*}{4d_\psi}$ and

$$T_j^{1/3} \geq \tilde{\Omega} \left(\left(D \lambda_j^{-2} (D^3 \lambda_j^{-1} + C_{\mathcal{R}} \cdot \log^{p_{\mathcal{R}}} \frac{1}{\delta_j}) \right) \cdot \frac{\lambda_{\min}^*}{d_\psi} \right), \quad (\text{C.3})$$

then $\lambda_{\min}(\Sigma_j) \geq \frac{\lambda_{\min}^*}{4d_\psi} \cdot T_j$ with probability at least $1 - \delta_j$. It follows that, with probability at least $1 - \delta_j$, the if statement on Line 5 will be true once $\lambda_j \leq \frac{\lambda_{\min}^*}{4d_\psi}$, (C.3) holds, and

$$\frac{\lambda_{\min}^*}{4d_\psi} \cdot T_j \geq 12544 D d_\psi \log \frac{2N(2 + 32T_j)}{\delta}. \quad (\text{C.4})$$

By our choice of $\lambda_j = T_j^{-1/18}$, a sufficient condition to ensure $\lambda_j \leq \frac{\lambda_{\min}^*}{4d_\psi}$, (C.3), and (C.4) is

$$T_j \geq \tilde{\Omega} \left(\max \left\{ \left(\frac{d_\psi}{\lambda_{\min}^*} \right)^{18}, \left(\frac{D^4 \lambda_{\min}^*}{d_\psi} \right)^6, \left(D C_{\mathcal{R}} \cdot \log^{p_{\mathcal{R}}} \frac{1}{\delta_j} \cdot \frac{\lambda_{\min}^*}{d_\psi} \right)^{9/2}, \frac{D d_\psi^2}{\lambda_{\min}^*} \cdot \log \frac{N T_j}{\delta} \right\} \right).$$

Since $T_j = \lceil 2^{j/3} \rceil \lceil 2^{2j/3} \rceil \in [2^j, 4 \cdot 2^j]$, it follows that the if statement on Line 5 will be met after running for at most

$$\tilde{\mathcal{O}} \left(\max \left\{ \left(\frac{d_\psi}{\lambda_{\min}^*} \right)^{18}, \left(\frac{D^4 \lambda_{\min}^*}{d_\psi} \right)^6, \left(D C_{\mathcal{R}} \cdot \log^{p_{\mathcal{R}}} \frac{1}{\delta_j} \cdot \frac{\lambda_{\min}^*}{d_\psi} \right)^{9/2}, \frac{D d_\psi^2}{\lambda_{\min}^*} \cdot \log \frac{N}{\delta} \right\} \right) \quad (\text{C.5})$$

episodes.

By Lemma C.5, if $\lambda_{\min}(\Sigma_j) \geq 12544 D d_\psi \log \frac{2N(2+32T_j)}{\delta}$ and we rerun all policies in Π_j , then we will collect data $\tilde{\Sigma}$ such that $\lambda_{\min}(\tilde{\Sigma}) \geq \frac{1}{2} \lambda_{\min}(\Sigma_j)$, with probability at least $1 - \delta/2N$. As the if

statement on Line 5 will only be true once this is met, it follows that, with probability at least $1 - \delta/2N$, rerunning all policies in Π_j once, we will collect data Σ which satisfies

$$\lambda_{\min}(\Sigma) \geq \frac{1}{2} \lambda_{\min}(\Sigma_j) \geq 6272Dd_\psi \log \frac{2N(2 + 32T_j)}{\delta} \geq 6272Dd_\psi \log \frac{68N}{\delta}.$$

The lower bound on $\lambda_{\min}(\sum_{\pi \in \Pi} \Gamma_\pi)$ follows analogously from Lemma C.5.

The result then follows noting that the failure probability of running DYNAMICOED is at most

$$\sum_{j=1}^{\infty} \frac{\delta}{4j^2} \leq \delta/2.$$

□

C.2.1 Supporting Lemmas

Lemma C.4. *Under Assumptions 10 to 12, consider running DYNAMICOED on the objective*

$$\Phi(\Gamma) = \text{tr}((\Gamma + \lambda \cdot I)^{-1})$$

with $N = \lceil 2^{i/3} \rceil - 1$ and $K = \lceil 2^{2i/3} \rceil$, for some $\lambda > 0$ and i . Let $T := (N + 1)K$. Then if $\lambda \leq \frac{\lambda_{\min}^*}{4d_\psi}$ and

$$T^{1/3} \geq \tilde{\Omega} \left(\left(D\lambda^{-2}(D^3\lambda^{-1} + C_{\mathcal{R}} \cdot \log^{p_{\mathcal{R}}} \frac{1}{\delta}) \right) \cdot \frac{\lambda_{\min}^*}{d_\psi} \right), \quad (\text{C.6})$$

with probability at least $1 - \delta$,

$$\lambda_{\min} \left(\sum_{t=1}^T \psi(\tau_t) \right) \geq \frac{\lambda_{\min}^*}{4d_\psi} \cdot T.$$

Proof. Applying Corollary 1 with $\mathcal{H} = I$ and $\Gamma_0 = \lambda \cdot I$, we have that, with probability at least $1 - \delta$:

$$\begin{aligned} \text{tr} \left(\left(\sum_{t=1}^T \psi(\tau_t) + T\lambda \cdot I \right)^{-1} \right) &\leq \frac{1}{T} \cdot \min_{\Gamma \in \Omega_\psi} \text{tr}((\Gamma + \lambda \cdot I)^{-1}) + \frac{8D^4\lambda^{-3}}{T(N+1)} + \frac{8D\lambda^{-2}(\log^{1/2} \frac{T}{\delta} + C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{T}{\delta})}{T\sqrt{K}} \\ &\leq \frac{1}{T} \cdot \min_{\Gamma \in \Omega_\psi} \text{tr}((\Gamma + \lambda \cdot I)^{-1}) + \frac{24D^4\lambda^{-3}}{T^{4/3}} + \frac{24D\lambda^{-2}(\log^{1/2} \frac{T}{\delta} + C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{T}{\delta})}{T^{4/3}} \end{aligned}$$

where the second inequality follows since

$$T = \lceil 2^{i/3} \rceil \lceil 2^{2i/3} \rceil \leq 4 \cdot 2^i$$

which implies $N + 1 = \lceil 2^{i/3} \rceil \geq T^{1/3}/4^{1/3}$ and $K = \lceil 2^{2i/3} \rceil \geq T^{2/3}/4^{2/3}$. If T satisfies (C.6), then we can bound

$$\frac{24D^4\lambda^{-3}}{T^{4/3}} + \frac{24D\lambda^{-2}(\log^{1/2} \frac{T}{\delta} + C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{T}{\delta})}{T^{4/3}} \leq \frac{1}{T} \cdot \frac{d_\psi}{\lambda_{\min}^*}.$$

Furthermore, under Assumption 3 there exists some $\mathbf{\Gamma} \in \mathbf{\Omega}_\psi$ such that $\mathbf{\Gamma} \succeq \lambda_{\min}^* \cdot I$, so we can upper bound

$$\min_{\mathbf{\Gamma} \in \mathbf{\Omega}_\psi} \text{tr}((\mathbf{\Gamma} + \lambda \cdot I)^{-1}) \leq \frac{d_\psi}{\lambda_{\min}^*}$$

and we can lower bound

$$\text{tr} \left(\left(\sum_{t=1}^T \psi(\tau_t) + T\lambda \cdot I \right)^{-1} \right) \geq \frac{1}{\lambda_{\min}(\sum_{t=1}^T \psi(\tau_t)) + T\lambda}.$$

Thus,

$$\frac{1}{\lambda_{\min}(\sum_{t=1}^T \psi(\tau_t)) + T\lambda} \leq \frac{1}{T} \cdot \frac{2d_\psi}{\lambda_{\min}^*} \implies \lambda_{\min} \left(\sum_{t=1}^T \psi(\tau_t) \right) \geq \frac{T\lambda_{\min}^*}{2d_\psi} - T\lambda.$$

It follows that if $\lambda \leq \frac{\lambda_{\min}^*}{4d_\psi}$, then we have

$$\lambda_{\min} \left(\sum_{t=1}^T \psi(\tau_t) \right) \geq T \cdot \frac{\lambda_{\min}^*}{4d_\psi}$$

which proves the result. $\square$

Lemma C.5. *Consider running some policies $(\pi_\tau)_{\tau=1}^T$, for π_τ $\mathcal{F}_{\tau-1}$ -measurable, and collecting covariance $\mathbf{\Sigma}_T = \sum_{t=1}^T \psi(\tau_t)$. Then under Assumption 11, as long as*

$$\lambda_{\min}(\mathbf{\Sigma}_T) \geq 12544Dd_\psi \log \frac{2 + 32T}{\delta}$$

with probability at least $1 - \delta$, if we rerun each $(\pi_\tau)_{\tau=1}^T$, we will collect features $\tilde{\mathbf{\Sigma}}_T$ such that

$$\lambda_{\min}(\tilde{\mathbf{\Sigma}}_T) \geq \frac{1}{2} \lambda_{\min}(\mathbf{\Sigma}_T).$$

Furthermore,

$$\lambda_{\min} \left(\sum_{\tau=1}^T \mathbf{\Gamma}_{\pi_\tau} \right) \geq \frac{1}{2} \lambda_{\min}(\mathbf{\Sigma}_T).$$

Proof. This follows from applying Lemma D.7 of Wagenmaker & Jamieson (2022) to the matrix $\frac{1}{D}\mathbf{\Sigma}_T$. Note that while Wagenmaker & Jamieson (2022) considers the setting of linear MDPs, the proof of Lemma D.7 of Wagenmaker & Jamieson (2022) does not make use of the linear MDP assumption, and the proof therefore extends immediately to our setting. Furthermore, though it is not explicitly stated, the lower bound on $\lambda_{\min}(\sum_{\tau=1}^T \mathbf{\Gamma}_{\pi_\tau})$ is also proved in Lemma D.7 of Wagenmaker & Jamieson (2022). $\square$

Algorithm 7 Learn Minimizing Exploration Policies (LEARNEXP Π)

```

1: input:  $\mathcal{H}$ , iterates bound  $\tilde{N}$ , confidence  $\delta$ , regret minimization algorithm  $\mathbb{A}_{\mathcal{R}}$ , exploration
   policies  $\Pi_{\text{exp}}$ 
2: for  $i = 1, 2, 3, \dots$  do
3:    $N_i \leftarrow \lceil 2^{i/3} \rceil - 1, K_i \leftarrow \lceil 2^{2i/3} \rceil, T_i \leftarrow (N_i + 1)K_i, \delta_i \leftarrow \delta/4i^2$ 
4:    $\Pi_{\text{MINEIG}}^i \leftarrow \text{MINEIG}(\tilde{N}T_i, \delta_i, \mathbb{A}_{\mathcal{R}}, \Pi_{\text{exp}})$ 
5:   Run policies in  $\Pi_{\text{MINEIG}}^i$   $\lceil T_i/|\Pi_{\text{MINEIG}}^i| \rceil$  times, set  $\mathbf{\Gamma}_0^i$  to collected features
6:    $\Phi(\mathbf{\Gamma}) \leftarrow \text{tr}(\mathcal{H}(\mathbf{\Gamma} + T_i^{-1}\mathbf{\Gamma}_0^i)^{-1})$ 
7:    $\mathbf{\Gamma}_{\text{fw}}^i, \Pi_{\text{fw}}^i \leftarrow \text{DYNAMICOED}(\Phi, N_i, K_i, \delta, \mathbb{A}_{\mathcal{R}}, \Pi_{\text{exp}})$ 
8:   if

```

$$\max_{j=1, \dots, i} |\Pi_{\text{MINEIG}}^j| \leq T_i \quad (\text{C.7})$$

$$\frac{16D^4 \|\mathcal{H}\|_{\text{op}} \|(T_i^{-1}\mathbf{\Gamma}_0^i)^{-1}\|_{\text{op}}^3}{T_i(N_i + 1)} + \frac{16D \|\mathcal{H}\|_{\text{op}} \|(T_i^{-1}\mathbf{\Gamma}_0^i)^{-1}\|_{\text{op}}^2 (\log^{1/2} \frac{4T_i}{\delta} + C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2T_i}{\delta})}{T_i \sqrt{K_i}} \leq \text{tr} \left(\mathcal{H} (\mathbf{\Gamma}_{\text{fw}}^i + \mathbf{\Gamma}_0^i)^{-1} \right) \quad (\text{C.8})$$

$$\text{tr}(\mathcal{H}) \cdot D \sqrt{2T_i} \sqrt{8d_{\psi} \log(1 + 8\sqrt{2T_i}) + 8 \log 1/\delta} \cdot \frac{2}{\lambda_{\min} (\mathbf{\Gamma}_{\text{fw}}^i + \mathbf{\Gamma}_0^i)^2} \leq \text{tr} \left(\mathcal{H} (\mathbf{\Gamma}_{\text{fw}}^i + \mathbf{\Gamma}_0^i)^{-1} \right) \quad (\text{C.9})$$

$$D \sqrt{2T_i} \sqrt{8d_{\psi} \log(1 + 8\sqrt{2T_i}) + 8 \log 1/\delta} \leq \frac{1}{2} \lambda_{\min} (\mathbf{\Gamma}_{\text{fw}}^i + \mathbf{\Gamma}_0^i) \quad (\text{C.10})$$

```

9:   then
10:      $\mathbf{\Gamma}_{\text{out}} \leftarrow \mathbf{\Gamma}_{\text{fw}}^i + \mathbf{\Gamma}_0^i, \Pi_{\text{out}} \leftarrow \Pi_{\text{MINEIG}}^i \cup (\cup_{j=1}^{\lceil T_i/|\Pi_{\text{MINEIG}}^i| \rceil} \Pi_{\text{fw}}^j)$ 
11:   return  $\mathbf{\Gamma}_{\text{out}}, \Pi_{\text{out}}$ 

```

C.3 Rerunning Policies

In this section, we build on the analysis of the DYNAMICOED algorithm to show that, not only do the features collected by DYNAMICOED approximately minimize Φ , but that, under certain conditions, if we rerun the policies that DYNAMICOED ran to collect this data, we will collect a new set of features which also approximately minimizes Φ .

In particular, we specialize this argument to objectives of the form $\Phi(\mathbf{\Gamma}) = \text{tr}(\mathcal{H}\mathbf{\Gamma}^{-1})$. LEARN-EXP Π (Algorithm 7) proceeds by first calling MINEIG to collect full-rank data, using this data as a regularizer of $\Phi(\mathbf{\Gamma})$, and the running DYNAMICOED on this objective. After meeting a certain termination criteria, it terminates, and returns the policies it has run over its operation.

Lemma C.6. *Let $\mathcal{E}_{\text{exp}}$ denote the event that, for all $i = 1, 2, 3, \dots$, the success event of MINEIG and DYNAMICOED occur, and*

$$\lambda_{\min}(\mathbf{\Gamma}_0^i) \geq \lceil T_i/|\Pi_{\text{MINEIG}}^i| \rceil \cdot 6272Dd_{\psi} \log \frac{68\tilde{N}}{\delta}.$$

Then if Assumptions 10 to 12 hold, $\mathbb{P}[\mathcal{E}_{\text{exp}}] \geq 1 - \delta$.

Lemma C.7. *Consider rerunning each policy in Π_{out} $N \leq \tilde{N}$ times, and let $\tilde{\mathbf{\Gamma}}$ denote the obtained*

features. Then, if Assumptions 10 to 12 hold, with probability at least $1 - 3\delta$, on the event $\mathcal{E}_{\text{exp}}$:

$$\left\| \tilde{\mathbf{\Gamma}} - N \cdot \sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right\|_{\text{op}} \leq \frac{\sqrt{T_{\text{out}}} \cdot \sqrt{8d_{\psi} \log(1 + 8\sqrt{NT_{\text{out}}}) + 8 \log 1/\delta}}{\sqrt{N} \cdot 6272d_{\psi} \log \frac{68\tilde{N}}{\delta}} \cdot \lambda_{\min} \left(N \cdot \sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right), \quad (\text{C.11})$$

$$N \cdot 6272Dd_{\psi} \log \frac{68\tilde{N}}{\delta} \leq \min \left\{ \lambda_{\min} \left(N \cdot \sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right), \lambda_{\min}(\tilde{\mathbf{\Gamma}}) \right\}, \quad (\text{C.12})$$

and

$$\text{tr} \left(\mathcal{H} \left(\sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right)^{-1} \right) \leq \frac{12}{T_{\text{out}}} \cdot \min_{\mathbf{\Gamma} \in \mathbf{\Omega}} \text{tr}(\mathcal{H}\mathbf{\Gamma}^{-1}) \quad (\text{C.13})$$

for $T_{\text{out}} := |\Pi_{\text{out}}|$.

Lemma C.8. On the event $\mathcal{E}_{\text{exp}}$, under Assumptions 10 to 12, we can bound

$$T_{\text{out}} \leq \text{poly} \left(d_{\psi}, \frac{1}{\lambda_{\min}^*}, D, C_{\mathcal{R}}, \log^{p_{\mathcal{R}}} \frac{\tilde{N}}{\delta} \right).$$

Furthermore, the total number of episodes collected by Algorithm 7 is bounded by $(16 + 2 \log(T_{\text{out}})) \cdot T_{\text{out}}$.

C.3.1 Supporting Lemmas and Proofs

Lemma C.9. Under Assumption 11, for any $\mathbf{\Gamma} = \mathbb{E}_{\boldsymbol{\tau} \sim \omega}[\boldsymbol{\psi}(\boldsymbol{\tau})]$ and $\mathcal{H} \succeq 0$ we can bound

$$\text{tr}(\mathcal{H}\mathbf{\Gamma}^{-1}) \geq D^{-1} \cdot \text{tr}(\mathcal{H}).$$

Proof. By Von Neumann's Trace Inequality we can lower bound

$$\text{tr}(\mathcal{H}\mathbf{\Gamma}^{-1}) \geq \lambda_{\min}(\mathbf{\Gamma}^{-1}) \cdot \text{tr}(\mathcal{H}) = \|\mathbf{\Gamma}\|_{\text{op}}^{-1} \cdot \text{tr}(\mathcal{H}).$$

By our assumption that $\text{tr}(\boldsymbol{\psi}(\boldsymbol{\tau})) \leq D$, we can bound $\|\mathbf{\Gamma}\|_{\text{op}} \leq D$, which proves the result. $\square$

Lemma C.10. Assume $\text{tr}(\boldsymbol{\psi}(\boldsymbol{\tau})) \leq D$ for all $\boldsymbol{\tau}$. Let $\mathbf{\Gamma}_K$ denote the time-normalized features obtained by playing policies $\{\pi_k\}_{k=1}^K$, where π_k is $\mathcal{F}_{k-1}$ -measurable. Then, with probability at least $1 - \delta$,

$$\left\| \frac{1}{K} \sum_{k=1}^K \mathbf{\Gamma}_{\pi_k} - \mathbf{\Gamma}_K \right\|_{\text{op}} \leq D \sqrt{\frac{8d_{\psi} \log(1 + 8\sqrt{K}) + 8 \log 1/\delta}{K}}.$$

Proof. This follows from an argument identical to the proof of Lemma C.4 of Wagenmaker & Jamieson (2022). While Wagenmaker & Jamieson (2022) considers the setting of linear MDPs, we note that the proof of Lemma C.4 of Wagenmaker & Jamieson (2022) nowhere relies on the linear MDP assumption. The result stated here then follows identically as Lemma C.4 of Wagenmaker & Jamieson (2022), after normalizing $\boldsymbol{\psi}(\boldsymbol{\tau})$ by D . $\square$

Proof of Lemma C.6. By Lemma C.3, the failure probability of running MINEIG at round i is $\delta_i = \delta/8i^2$, and by Corollary 1 the failure probability of DYNAMICOED at round i is also bounded by $\delta_i = \delta/8i^2$. It follows that the total failure probability of running MINEIG and DYNAMICOED is bounded by

$$\sum_{i=1}^{\infty} \frac{2 \cdot \delta}{8i^2} \leq \frac{\delta}{2}.$$

Furthermore, by Lemma C.3, we have that rerunning all policies in Π_{MINEIG}^i , we will obtain features $\mathbf{\Gamma}$ satisfying, with probability at least $1 - \delta_i/\tilde{N}T_i$:

$$\lambda_{\min}(\mathbf{\Gamma}) \geq 6272Dd_{\psi} \log \frac{68\tilde{N}}{\delta_i}.$$

Repeating this $\lceil T_i/|\Pi_{\text{MINEIG}}^i| \rceil$ times and union bounding, we have that

$$\lambda_{\min}(\mathbf{\Gamma}_0^i) \geq \lceil T_i/|\Pi_{\text{MINEIG}}^i| \rceil \cdot 6272Dd_{\psi} \log \frac{68\tilde{N}}{\delta_i}$$

with probability at least $1 - \delta/\tilde{N}T_i \cdot \lceil T_i/|\Pi_{\text{MINEIG}}^i| \rceil \geq 1 - \delta/\tilde{N}$. $\square$

Proof of Lemma C.7. Let $\Pi_{\text{MINEIG}}, \mathbf{\Gamma}_0, \Pi_{\text{fw}}, \mathbf{\Gamma}_{\text{fw}}$ denote the policies and features obtained on the round at which Algorithm 7 terminates. Let $T_{\text{fw}} = |\Pi_{\text{fw}}|$ denote the number of episodes of DYNAMICOED on the terminating round, and $N_{\text{fw}}, K_{\text{fw}}$ the corresponding values of N_i and K_i . Throughout the proof we make use of the fact that at termination of Algorithm 7, all of (C.7)-(C.10) are met.

Proof of (C.11) and (C.12). By Lemma C.10 we have that, with probability at least $1 - \delta$:

$$\left\| \tilde{\mathbf{\Gamma}} - N \cdot \sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right\|_{\text{op}} \leq D\sqrt{NT_{\text{out}}} \cdot \sqrt{8d_{\psi} \log(1 + 8\sqrt{NT_{\text{out}}})} + 8 \log 1/\delta.$$

On $\mathcal{E}_{\text{exp}}$, by Lemma C.3, we can bound

$$|\Pi_{\text{MINEIG}}| \leq \text{poly} \left(d_{\psi}, \frac{1}{\lambda_{\min}^*}, D, C_{\mathcal{R}}, \log^{p_{\mathcal{R}}} \frac{\tilde{N}T_{\text{out}}}{\delta} \right)$$

and, furthermore, we can lower bound

$$\lambda_{\min} \left(\sum_{\pi \in \Pi_{\text{MINEIG}}} \mathbf{\Gamma}_{\pi} \right) \geq 6272Dd_{\psi} \log \frac{68\tilde{N}}{\delta}.$$

Since $\Pi_{\text{MINEIG}} \subseteq \Pi_{\text{out}}$, it follows that

$$\lambda_{\min} \left(N \cdot \sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right) \geq N \cdot 6272Dd_{\psi} \log \frac{68\tilde{N}}{\delta}.$$

Combining these, we therefore have that, with probability at least $1 - \delta$:

$$\left\| \tilde{\mathbf{\Gamma}} - N \cdot \sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right\|_{\text{op}} \leq \frac{\sqrt{NT_{\text{out}}} \cdot \sqrt{8d_{\psi} \log(1 + 8\sqrt{NT_{\text{out}}})} + 8 \log 1/\delta}{N \cdot 6272Dd_{\psi} \log \frac{68\tilde{N}}{\delta}} \cdot \lambda_{\min} \left(N \cdot \sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right).$$

In addition, also by Lemma C.3, we have that with probability at least $1 - \delta/\tilde{N}T_{\text{out}} \cdot N \geq 1 - \delta$, that

$$\lambda_{\min}(\tilde{\mathbf{\Gamma}}) \geq N \cdot 6272Dd_{\psi} \log \frac{68\tilde{N}}{\delta}.$$

Proof of (C.13). By Corollary 1, on $\mathcal{E}_{\text{exp}}$ we have that:

$$\begin{aligned} \Phi(\mathbf{\Gamma}_{\text{fw}}) = \text{tr} \left(\mathcal{H} (\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0)^{-1} \right) &\leq \frac{\min_{\mathbf{\Gamma} \in \Omega} \text{tr} \left(\mathcal{H} (\mathbf{\Gamma} + T_{\text{fw}}^{-1} \mathbf{\Gamma}_0)^{-1} \right)}{T_{\text{fw}}} \\ &+ \frac{8D^4 \|\mathcal{H}\|_{\text{op}} \|(T_{\text{fw}}^{-1} \mathbf{\Gamma}_0)^{-1}\|_{\text{op}}^3}{T_{\text{fw}}(N_{\text{fw}} + 1)} + \frac{8D \|\mathcal{H}\|_{\text{op}} \|(T_{\text{fw}}^{-1} \mathbf{\Gamma}_0)^{-1}\|_{\text{op}}^2 (\log^{1/2} \frac{4T_{\text{fw}}}{\delta} + C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2T_{\text{fw}}}{\delta})}{T_{\text{fw}} \sqrt{K_{\text{fw}}}}. \end{aligned}$$

Since T_{fw} satisfies (C.8), we can bound

$$\frac{8D^4 \|\mathcal{H}\|_{\text{op}} \|(T_{\text{fw}}^{-1} \mathbf{\Gamma}_0)^{-1}\|_{\text{op}}^3}{T_{\text{fw}}(N_{\text{fw}} + 1)} + \frac{8D \|\mathcal{H}\|_{\text{op}} \|(T_{\text{fw}}^{-1} \mathbf{\Gamma}_0)^{-1}\|_{\text{op}}^2 (\log^{1/2} \frac{4T_{\text{fw}}}{\delta} + C_{\mathcal{R}} \log^{p_{\mathcal{R}}} \frac{2T_{\text{fw}}}{\delta})}{T_{\text{fw}} \sqrt{K_{\text{fw}}}} \leq \frac{1}{2} \Phi(\mathbf{\Gamma}_{\text{fw}}).$$

It follows that

$$\begin{aligned} \Phi(\mathbf{\Gamma}_{\text{fw}}) &\leq \frac{\min_{\mathbf{\Gamma} \in \Omega} \text{tr} \left(\mathcal{H} (\mathbf{\Gamma} + T_{\text{fw}}^{-1} \mathbf{\Gamma}_0)^{-1} \right)}{T_{\text{fw}}} + \frac{1}{2} \Phi(\mathbf{\Gamma}_{\text{fw}}) \\ \implies \Phi(\mathbf{\Gamma}_{\text{fw}}) &\leq 2 \cdot \frac{\min_{\mathbf{\Gamma} \in \Omega} \text{tr} \left(\mathcal{H} (\mathbf{\Gamma} + T_{\text{fw}}^{-1} \mathbf{\Gamma}_0)^{-1} \right)}{T_{\text{fw}}}. \end{aligned} \quad (\text{C.14})$$

By Lemma C.10, we have that, with probability at least $1 - \delta$:

$$\left\| \sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} - (\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0) \right\|_{\text{op}} \leq D \sqrt{T_{\text{out}}} \sqrt{8d_{\psi} \log(1 + 8\sqrt{T_{\text{out}}}) + 8 \log 1/\delta}.$$

Since (C.10) is satisfied and $|\Pi_{\text{MINEIG}}^i| \leq T_i$, we have

$$D \sqrt{T_{\text{out}}} \sqrt{8d_{\psi} \log(1 + 8\sqrt{T_{\text{out}}}) + 8 \log 1/\delta} \leq \frac{1}{2} \lambda_{\min}(\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0).$$

By Lemma A.2 it follows that

$$\left\| \left(\sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right)^{-1} - (\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0)^{-1} \right\|_{\text{op}} \leq D \sqrt{T_{\text{out}}} \sqrt{8d_{\psi} \log(1 + 8\sqrt{T_{\text{out}}}) + 8 \log 1/\delta} \cdot \frac{2}{\lambda_{\min}(\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0)^2}.$$

This implies that

$$\begin{aligned} \text{tr} \left(\mathcal{H} \left(\sum_{\pi \in \Pi_{\text{out}}} \mathbf{\Gamma}_{\pi} \right)^{-1} \right) &\leq \text{tr} \left(\mathcal{H} (\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0)^{-1} \right) \\ &+ \text{tr}(\mathcal{H}) \cdot D \sqrt{T_{\text{out}}} \sqrt{8d_{\psi} \log(1 + 8\sqrt{T_{\text{out}}}) + 8 \log 1/\delta} \cdot \frac{2}{\lambda_{\min}(\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0)^2}. \end{aligned}$$

Now if

$$\text{tr}(\mathcal{H}) \cdot D \sqrt{T_{\text{out}}} \sqrt{8d_{\psi} \log(1 + 8\sqrt{T_{\text{out}}}) + 8 \log 1/\delta} \cdot \frac{2}{\lambda_{\min}(\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0)^2} \leq \text{tr} \left(\mathcal{H} (\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0)^{-1} \right), \quad (\text{C.15})$$

we can bound this all by

$$\leq 2 \text{tr} \left(\mathcal{H} (\mathbf{\Gamma}_{\text{fw}} + \mathbf{\Gamma}_0)^{-1} \right) \leq \frac{4}{T_{\text{fw}}} \cdot \min_{\mathbf{\Gamma} \in \Omega} \text{tr} \left(\mathcal{H} (\mathbf{\Gamma} + T_{\text{fw}}^{-1} \mathbf{\Gamma}_0)^{-1} \right)$$

where the last inequality follows from (C.14). However, note that (C.15) since (C.9) holds. Finally, note that

$$T_{\text{out}} = T_{\text{fw}} + \lceil T_{\text{fw}} / |\Pi_{\text{MINEIG}}| \rceil |\Pi_{\text{MINEIG}}| \leq 2T_{\text{fw}} + |\Pi_{\text{MINEIG}}| \leq 3T_{\text{fw}}$$

where the last inequality follows since (C.7) holds. We can therefore upper bound $\frac{4}{T_{\text{fw}}} \leq \frac{12}{T_{\text{out}}}$. Putting this together proves the result. $\square$

Proof of Lemma C.8. To bound T_{out} , it suffices to show that (C.7)-(C.10) are satisfied for sufficiently large T_i .

On $\mathcal{E}_{\text{exp}}$, by Lemma C.3, we can bound

$$|\Pi_{\text{MINEIG}}^i| \leq \text{poly} \left(d_\psi, \frac{1}{\lambda_{\min}^*}, D, C_{\mathcal{R}}, \log^{p_{\mathcal{R}}} \frac{\tilde{N}T_i}{\delta} \right),$$

so to ensure (C.7) is met it suffices that

$$T_i \geq \text{poly} \left(d_\psi, \frac{1}{\lambda_{\min}^*}, D, C_{\mathcal{R}}, \log^{p_{\mathcal{R}}} \frac{\tilde{N}}{\delta} \right).$$

On $\mathcal{E}_{\text{exp}}$, we have

$$\lambda_{\min}(\mathbf{\Gamma}_{\text{fw}}^i + \mathbf{\Gamma}_0^i) \geq \lambda_{\min}(\mathbf{\Gamma}_0^i) \geq \lceil T_i / |\Pi_{\text{MINEIG}}^i| \rceil \cdot 6272 D d_\psi \log \frac{68\tilde{N}}{\delta}.$$

Which also implies

$$\|(T_i^{-1} \mathbf{\Gamma}_0^i)^{-1}\|_{\text{op}} = \frac{T_i}{\lambda_{\min}(\mathbf{\Gamma}_0^i)} \leq \frac{T_i}{\lceil T_i / |\Pi_{\text{MINEIG}}^i| \rceil} \cdot \frac{1}{6272 D d_\psi \log \frac{68\tilde{N}}{\delta}} \leq \frac{|\Pi_{\text{MINEIG}}^i|}{6272 D d_\psi \log \frac{68\tilde{N}}{\delta}}$$

Furthermore, by Lemma C.9 we can lower bound

$$\text{tr} \left(\mathcal{H} (\mathbf{\Gamma}_{\text{fw}}^i + \mathbf{\Gamma}_0^i)^{-1} \right) \geq \frac{\text{tr}(\mathcal{H})}{D(T_i + \lceil T_i / |\Pi_{\text{MINEIG}}^i| \rceil |\Pi_{\text{MINEIG}}^i|)} \geq \frac{\text{tr}(\mathcal{H})}{3DT_i}.$$

Combining these and using that $N_i = \mathcal{O}(T_i^{1/3})$ and $K_i = \mathcal{O}(T_i^{2/3})$, it is easy to see that (C.8)-(C.10) will be met once

$$T_i \geq \text{poly} \left(d_\psi, \frac{1}{\lambda_{\min}^*}, D, C_{\mathcal{R}}, \log^{p_{\mathcal{R}}} \frac{\tilde{N}}{\delta} \right).$$

The bound on T_{out} then follows since $T_i = \lceil 2^{i/3} \rceil \lceil 2^{2i/3} \rceil \in [2^i, 4 \cdot 2^i]$, so it can be at most a constant larger than the sufficient condition before terminating.

Let i^* denote the round that Algorithm 7 terminates on. Note that at round i , MINEIG runs for at most $|\Pi_{\text{MINEIG}}^i|$, DYNAMICOED runs for at most T_i episodes, and we run for an additional $\lceil T_i / |\Pi_{\text{MINEIG}}^i| \rceil \cdot |\Pi_{\text{MINEIG}}^i|$ episodes on Line 5. In total, then, the number of episodes Algorithm 7 runs for is bounded by

$$\begin{aligned} \sum_{i=1}^{i^*} (T_i + |\Pi_{\text{MINEIG}}^i| + \lceil T_i / |\Pi_{\text{MINEIG}}^i| \rceil \cdot |\Pi_{\text{MINEIG}}^i|) &\leq 2 \sum_{i=1}^{i^*} (T_i + |\Pi_{\text{MINEIG}}^i|) \\ &\leq 16T_{i^*} + 2 \sum_{i=1}^{i^*} |\Pi_{\text{MINEIG}}^i| \end{aligned}$$

where the last inequality follows since $T_i \in [2^i, 4 \cdot 2^i]$. Now note that, since Algorithm 7 only terminates once (C.7) is met, we will have $\max_{j=1, \dots, i^*} |\Pi_{\text{MINEIG}}^j| \leq T_{i^*}$. This implies that $2 \sum_{i=1}^{i^*} |\Pi_{\text{MINEIG}}^i| \leq 2i^* T_{i^*} \leq 2 \log(T_{i^*}) \cdot T_{i^*}$. Bounding $T_{i^*} \leq T_{\text{out}}$ gives the result. $\square$

D Smooth Nonlinear Systems

In this section we restrict to the nonlinear regulator system of (1.1). Our goal will be to show that, under our assumptions, the nonlinear regulator system exhibits certain smooth behavior. As we have assumed

$$\Pi^\star = \{\pi^\theta : \theta \in \mathbb{R}^{d_\theta}\},$$

it will be convenient to define $\mathcal{J}(\theta; A) := \mathcal{J}(\pi^\theta; A)$ and $\theta_\star(A) = \theta_\star$. For the remainder of this section, we will typically use θ in place of π^θ . In addition, when considering radius terms such as $r_\theta(A_\star)$ and $r_{\text{cost}}(A_\star)$, to simplify results we assume that $r_\theta(A_\star) \leq 1$ and $r_{\text{cost}}(A_\star) \leq 1$. Note that this does not change the validity of the result since, for example, if a result holds with $A \in \mathcal{B}_F(A_\star; r)$ for some $r > r_\theta(A_\star)$, it also holds for $A \in \mathcal{B}_F(A_\star; r_\theta(A_\star))$. Throughout this section, we let $\nabla_x f(x)[\Delta]$ refer to the directional gradient of $f(x)$ in direction Δ .

We first have the following result, which shows that under our assumptions, the controller loss is differentiable.

Lemma D.1. *Under Assumptions 1, 2, 4 and 5, for any A satisfying $A \in \mathcal{B}_F(A_\star; r_\theta(A_\star))$, the controller loss $\mathcal{J}(\theta; A)$ is four-times differentiable in θ and A . Furthermore, we can bound*

$$\|\nabla_A^{(i)} \nabla_\theta^{(j)} \mathcal{J}(\theta; A)\|_{\text{op}} \leq \text{poly}(\|A\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x)$$

for $i, j \in \{0, 1, 2, 3, 4\}$ satisfying $1 \leq i + j \leq 3$.

In this section, we generalize Assumption 6 to the following.

Assumption 13. *We assume there exists some $r_\theta(A_\star) > 0$ such that, for all $A \in \mathcal{B}_F(A_\star; r_\theta(A_\star))$, $\theta_\star(A)$ satisfies:*

- $\nabla_\theta \mathcal{J}(\theta; A)|_{\theta=\theta_\star(A)} = 0$,
- $\theta_\star(A)$ is three-times differentiable in A , and we can bound $\|\nabla_A^{(i)} \theta_\star(A)\|_{\text{op}} \leq L_{\pi_\star}$ for some $L_{\pi_\star} > 0$ and $i \in \{1, 2, 3\}$.

The first condition requires that $\theta_\star(A)$ corresponds to a stationary point of the loss. This will be met, for example, by choosing $\theta_\star(A)$ to be a minima (local or global) of $\mathcal{J}(\theta; A)$. It is not obvious, however, that the first and second condition can be simultaneously satisfied. In the following we show that, assuming $\nabla_\theta^2 \mathcal{J}(\theta; A_\star)|_{\theta=\theta_\star(A_\star)}$ is full-rank (which will be the case, for example, when $\theta_\star(A_\star)$ is a strict local minimum of $\mathcal{J}(\pi^\theta; A_\star)$), there always exists some $\theta_\star(A)$ satisfying both conditions of Assumption 13, with $L_{\pi_\star}$ scaling polynomially in problem parameters, and $r_\theta(A_\star)$ scaling inverse polynomially in problem parameters. Note that this definition of $\pi_\star(A)$ is general enough to capture settings where the global minimum of $\mathcal{J}(\pi; A)$ cannot be efficiently computed—it suffices to take $\pi_\star(A)$ a local minimum of the loss.

Proposition 5. *Assume that Assumptions 1, 2, 4 and 5 hold and that $\lambda_{\min}(\nabla_\theta^2 \mathcal{J}(\theta; A_\star)|_{\theta=\theta_\star(A_\star)}) > 0$. Let $r_\theta(A_\star) > 0$ be some value satisfying*

$$r_\theta(A_\star) = \min \left\{ r_{\text{cost}}(A_\star), \text{poly} \left(\frac{1}{\lambda_{\min}(\nabla_\theta^2 \mathcal{J}(\theta; A_\star)|_{\theta=\theta_\star(A_\star)})}, \|A_\star\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x \right)^{-1} \right\}.$$

Then there exists some function $\theta_\star(A)$ such that, for all $A \in \mathcal{B}_F(A_\star; r_\theta(A_\star))$:

- $\nabla_{\boldsymbol{\theta}} \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_*(A)} = 0$,
- $\boldsymbol{\theta}_*(A)$ is three-times differentiable in A ,

and it suffices that we take

$$L_{\pi_*} = \text{poly} \left(\frac{1}{\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_*)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_*(A_*)})}, \|A_*\|_{\text{op}}, B_{\phi}, L_{\phi}, L_{\boldsymbol{\theta}}, L_{\text{cost}}, \sigma_{\mathbf{w}}^{-1}, H, d_{\mathbf{x}} \right).$$

While Proposition 5 shows that there exists some $\boldsymbol{\theta}_*(A)$ satisfying Assumption 13, it does not directly give a recipe for constructing such a map. The following result shows that under a mild additional assumption, the minimizer of the loss satisfies Assumption 13.

Proposition 6. *Let*

$$\boldsymbol{\theta}_*(A) := \arg \min_{\boldsymbol{\theta} \in \mathbb{R}^{d_{\boldsymbol{\theta}}}} \mathcal{J}(\boldsymbol{\theta}; A).$$

Then under Assumptions 1, 2 and 4 to 6, there exists some $r_{\boldsymbol{\theta}}(A_) > 0$ and $L_{\pi_*} < \infty$ such that $\boldsymbol{\theta}_*(A)$ satisfies Assumption 13.*

The scaling of L_{π_*} in Proposition 6 can be shown to match that of Proposition 5, but in general $r_{\boldsymbol{\theta}}(A_*)$ could be smaller than the value of $r_{\boldsymbol{\theta}}(A_*)$ given in Proposition 5. In particular, in the setting of Proposition 6, we can only show that $r_{\boldsymbol{\theta}}(A_*)$ scales with $\min_{\boldsymbol{\theta} \notin \mathcal{B}_2(\boldsymbol{\theta}_*(A_*); r)} \mathcal{J}(\boldsymbol{\theta}; A_*) - \mathcal{J}(\boldsymbol{\theta}_*(A_*); A_*)$ for some $r > 0$ which scales inverse polynomially in problem parameters. While we can show that $\mathcal{J}(\boldsymbol{\theta}; A_*) - \mathcal{J}(\boldsymbol{\theta}_*(A_*); A_*)$ scales inverse polynomially in problem parameters, including in $\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_*)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_*(A_*)})$, for $\boldsymbol{\theta}$ approximately a distance of r from $\boldsymbol{\theta}_*(A_*)$, it is possible $\mathcal{J}(\boldsymbol{\theta}; A_*)$ has some local minimizer $\boldsymbol{\theta}'$ arbitrarily far away from $\boldsymbol{\theta}_*(A_*)$, such that $\mathcal{J}(\boldsymbol{\theta}'; A_*)$ and $\mathcal{J}(\boldsymbol{\theta}_*(A_*); A_*)$ are arbitrarily close, in which case Δ^* , and therefore $r_{\boldsymbol{\theta}}(A_*)$, could be arbitrarily small. The failure mode here is that, while $\boldsymbol{\theta}_*(A_*)$ may be the global minimum of $\mathcal{J}(\boldsymbol{\theta}; A_*)$, for A arbitrarily close to A_* , the global minimum of $\mathcal{J}(\boldsymbol{\theta}; A)$ could instead be near $\boldsymbol{\theta}'$, which would render the map $\boldsymbol{\theta}_*(A)$ discontinuous.

By making further assumptions on $\mathcal{J}(\boldsymbol{\theta}; A_*)$ which exclude this case, we can obtain a value of $r_{\boldsymbol{\theta}}(A_*)$ scaling similarly to in Proposition 5. For example, in the following, we show that under the assumption that $\mathcal{J}(\boldsymbol{\theta}; A)$ is convex, this holds.

Proposition 7. *Assume that there exists some $r_{\text{conv}}(A_*) > 0$ such that, for all $A \in \mathcal{B}_{\text{F}}(A_*; r_{\text{conv}}(A_*))$, $\mathcal{J}(\boldsymbol{\theta}; A)$ is convex in $\boldsymbol{\theta}$, and set*

$$\boldsymbol{\theta}_*(A) = \arg \min_{\boldsymbol{\theta} \in \mathbb{R}^{d_{\boldsymbol{\theta}}}} \mathcal{J}(\boldsymbol{\theta}; A).$$

Then we have that $\boldsymbol{\theta}_$ satisfies Assumption 13 with*

$$r_{\boldsymbol{\theta}}(A_*) = \min \left\{ r_{\text{conv}}(A_*), r_{\text{cost}}(A_*), \text{poly} \left(\frac{1}{\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_*)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_*(A_*)})}, \|A_*\|_{\text{op}}, B_{\phi}, L_{\phi}, L_{\boldsymbol{\theta}}, L_{\text{cost}}, \sigma_{\mathbf{w}}^{-1}, H, d_{\mathbf{x}} \right)^{-1} \right\}$$

and it suffices that we take

$$L_{\pi_*} = \text{poly} \left(\frac{1}{\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_*)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_*(A_*)})}, \|A_*\|_{\text{op}}, B_{\phi}, L_{\phi}, L_{\boldsymbol{\theta}}, L_{\text{cost}}, \sigma_{\mathbf{w}}^{-1}, H, d_{\mathbf{x}} \right).$$

Note that, if $\mathcal{J}(\boldsymbol{\theta}; A)$ is μ -strongly convex in $\boldsymbol{\theta}$ for all A near A_* , we can lower bound $\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_*)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_*(A_*)}) \geq \mu$.

Approximating the Controller Loss. In order to efficiently direct our exploration, it is convenient to derive a quadratic approximation to the controller loss. The following result shows that, under our assumptions, this is indeed possible.

Lemma D.2 (Formal Version of Proposition 1). *Under Assumptions 1, 2, 4 to 5 and 13, for $\hat{A} \in \mathcal{B}_F(A_\star; \min\{r_{\text{cost}}(A_\star), r_\theta(A_\star)\})$, we have*

$$\mathcal{J}(\theta_\star(\hat{A}); A_\star) - \mathcal{J}(\theta_\star(A_\star); A_\star) \leq \text{vec}(\hat{A} - A_\star)^\top \mathcal{H}(A_\star) \text{vec}(\hat{A} - A_\star) + M[\hat{A} - A_\star, \hat{A} - A_\star, \hat{A} - A_\star],$$

for some tensor M such that

$$\|M[\hat{A} - A_\star, \hat{A} - A_\star, \hat{A} - A_\star]\|_{\text{op}} \leq \text{poly}(L_{\pi_\star}, \|A_\star\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x) \cdot \|\hat{A} - A_\star\|_{\text{op}}^3.$$

In practice we do not know $\mathcal{H}(A_\star)$ and must estimate it. The following result shows that the distance between $\mathcal{H}(A_\star)$ and $\mathcal{H}(\hat{A})$ can be bounded.

Lemma D.3. *Under Assumptions 1, 2, 4, 5 and 13, and if $\hat{A} \in \mathcal{B}_F(A_\star; \min\{r_{\text{cost}}(A_\star), r_\theta(A_\star)\})$, we can bound*

$$\|\mathcal{H}(A_\star) - \mathcal{H}(\hat{A})\|_{\text{op}} \leq \text{poly}(L_{\pi_\star}, \|A_\star\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x) \cdot \|\hat{A} - A_\star\|_{\text{op}}.$$

D.1 Proof of Smoothness of Nonlinear System

We let $f_w(\cdot)$ denote the density of the noise (which, by assumption, is simply an isotropic Gaussian density). We let $f_{A,\theta}(\cdot)$ denote the density over trajectories induced by playing controller θ on system A . We will overload notation somewhat and let $f_{A,\theta}(\mathbf{x}_{h+1} \mid \boldsymbol{\tau}_{1:h})$ denote the density over $\mathbf{x}_{h+1}$ induced by playing controller θ given trajectory $\boldsymbol{\tau}_{1:h}$. Note that $f_{A,\theta}(\mathbf{x}_{h+1} \mid \boldsymbol{\tau}_{1:h}) = f_w(\mathbf{x}_{h+1} - A\phi(\mathbf{x}_h^\tau, \pi_h^\theta(\boldsymbol{\tau}_{1:h})))$ and

$$f_{A,\theta}(\boldsymbol{\tau}) = \prod_{h=1}^H f_{A,\theta}(\mathbf{x}_{h+1} \mid \boldsymbol{\tau}_{1:h}).$$

Throughout this section we let $\mathbf{x}_h^\tau$ (resp. $\mathbf{u}_h^\tau$) denote the state (resp. input) at step h of trajectory $\boldsymbol{\tau}$. Under our regularity assumptions (Assumptions 1, 2, 4 and 5) and since the noise is Gaussian, we can swap derivatives and integrals, which we make use of throughout the following proofs.

Proof of Lemma D.1. Let $\text{cost}(\boldsymbol{\tau})$ denote the cost of trajectory $\boldsymbol{\tau}$. Then we have

$$\mathcal{J}(\theta; A) = \int \text{cost}(\boldsymbol{\tau}) f_{A,\theta}(\boldsymbol{\tau}) d\boldsymbol{\tau}.$$

Let $A_t := A + t_1 \Delta_1^A + t_2 \Delta_2^A + t_3 \Delta_3^A$ and $\theta_s := \theta + s_1 \Delta_1^\theta + s_2 \Delta_2^\theta + s_3 \Delta_3^\theta$, for some Δ_i^A and Δ_j^θ , which we assume satisfy $\|\Delta_i^A\|_{\text{op}}, \|\Delta_j^\theta\|_{\text{op}} \leq 1$. Rather than differentiating $\mathcal{J}(\theta; A)$ with respect to θ or A , we will differentiate $\mathcal{J}(\theta_s; A_t)$ with respect to some $x_1, x_2, x_3, x_4 \in \{t_1, t_2, t_3, s_1, s_2, s_3\}$. Note that, for example,

$$\frac{d}{dt_1} \mathcal{J}(\theta_s; A_t)|_{t=s=0} = \nabla_A \mathcal{J}(\theta, A)[\Delta_1^A],$$

i.e. the directional gradient of $\mathcal{J}(\theta, A)$ with respect to A in direction Δ_1^A , and that this similarly holds for gradients with respect to other t_i, s_j , or higher-order derivatives. Thus, if we can show that $\mathcal{J}(\theta_s; A_t)$ is differentiable with respect to any $x_1, x_2, x_3, x_4 \in \{t_1, t_2, t_3, s_1, s_2, s_3\}$, and this holds for any choice of $\Delta_i^A, \Delta_j^\theta$, then we have that $\mathcal{J}(\theta, A)$ is four-times differentiable with respect to θ and A . Furthermore, we can bound the operator norm of $\nabla_A \mathcal{J}(\theta, A)$, by bounding the value of $\frac{d}{dt_1} \mathcal{J}(\theta_s; A_t)|_{t=s=0}$ for all Δ_1^A satisfying $\|\Delta_1^A\|_{\text{op}} \leq 1$ (and we can similarly bound the operator norm of the higher order derivatives of $\mathcal{J}(\theta, A)$).

$\mathcal{J}(\boldsymbol{\theta}; A)$ is **Differentiable**. Let $x_1, x_2, x_3, x_4 \in \{t_1, t_2, t_3, s_1, s_2, s_3\}$. We have

$$\begin{aligned} \frac{d}{dx_1} \mathcal{J}(\boldsymbol{\theta}_s; A_t) &= \frac{d}{dx_1} \int \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &= \int \frac{f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau})}{f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau})} \frac{d}{dx_1} f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \text{cost}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &= \int \frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau}. \end{aligned}$$

Differentiating this gives

$$\begin{aligned} \frac{d}{dx_2} \frac{d}{dx_1} \mathcal{J}(\boldsymbol{\theta}_s; A_t) &= \frac{d}{dx_2} \int \frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &= \int \left(\frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \left(\frac{d}{dx_2} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &\quad + \int \frac{d}{dx_2} \frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau}, \end{aligned}$$

and

$$\begin{aligned} \frac{d}{dx_3} \frac{d}{dx_2} \frac{d}{dx_1} \mathcal{J}(\boldsymbol{\theta}_s; A_t) &= \int \left(\frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \left(\frac{d}{dx_2} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \left(\frac{d}{dx_3} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &\quad + \int \left(\frac{d}{dx_3} \frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \left(\frac{d}{dx_2} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &\quad + \int \left(\frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \left(\frac{d}{dx_3} \frac{d}{dx_2} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &\quad + \int \frac{d}{dx_3} \frac{d}{dx_2} \frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &\quad + \int \left(\frac{d}{dx_2} \frac{d}{dx_1} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \left(\frac{d}{dx_3} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) \right) \cdot \text{cost}(\boldsymbol{\tau}) f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) d\boldsymbol{\tau}. \end{aligned}$$

The fourth derivative of $\mathcal{J}(\boldsymbol{\theta}; A)$ can be similarly calculated by differentiating $\frac{d}{dx_3} \frac{d}{dx_2} \frac{d}{dx_1} \mathcal{J}(\boldsymbol{\theta}_s; A_t)$; we omit it for brevity. We have

$$\begin{aligned} \log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau}) &= \log \prod_{h=1}^H f_{A_t, \boldsymbol{\theta}_s}(\mathbf{x}_{h+1}^\tau \mid \boldsymbol{\tau}_{1:h}) \\ &= \sum_{h=1}^H \log f_{\mathbf{w}}(\mathbf{x}_{h+1}^\tau - A_t \boldsymbol{\phi}(\mathbf{x}_h^\tau, \pi^{\boldsymbol{\theta}_s}(\boldsymbol{\tau}_{1:h}))) \\ &= \sum_{h=1}^H -\frac{1}{2\sigma_{\mathbf{w}}^2} \|\mathbf{x}_{h+1}^\tau - A_t \boldsymbol{\phi}(\mathbf{x}_h^\tau, \pi^{\boldsymbol{\theta}_s}(\boldsymbol{\tau}_{1:h}))\|_2^2 + C \end{aligned}$$

for some C which does not depend on t or s . Given that $\boldsymbol{\phi}(\mathbf{x}, \mathbf{u})$ is four-times differentiable in $\mathbf{u}$ and $\pi_h^{\boldsymbol{\theta}_s}(\boldsymbol{\tau}_{1:h})$ is four-times differentiable in $\mathbf{x}$ (which hold by Assumption 4 and Assumption 5), it is clear that $\log f_{A_t, \boldsymbol{\theta}_s}(\boldsymbol{\tau})$ is four-times differentiable in t_i or s_i , regardless of the choice of Δ_i^A or Δ_i^θ . This proves the first result.

Norm Bounds on Gradient. Note that

$$\begin{aligned}\frac{d}{dt_i} \log f_{A_t, \theta_s}(\boldsymbol{\tau})|_{t=s=0} &= \sum_{h=1}^H \frac{1}{\sigma_w^2} (\mathbf{x}_{h+1}^\tau - A\phi(\mathbf{x}_h^\tau, \pi_h^\theta(\boldsymbol{\tau}_{1:h})))^\top \cdot \Delta_i^A \phi(\mathbf{x}_h^\tau, \pi_h^\theta(\boldsymbol{\tau}_{1:h})), \\ \frac{d}{ds_i} \log f_{A_t, \theta_s}(\boldsymbol{\tau})|_{t=s=0} &= \sum_{h=1}^H \frac{1}{\sigma_w^2} (\mathbf{x}_{h+1}^\tau - A\phi(\mathbf{x}_h^\tau, \pi_h^\theta(\boldsymbol{\tau}_{1:h})))^\top \cdot A \nabla_{\mathbf{u}} \phi(\mathbf{x}_h^\tau, \pi_h^\theta(\boldsymbol{\tau}_{1:h})) \cdot \nabla_{\boldsymbol{\theta}} \pi_h^\theta(\boldsymbol{\tau}_{1:h}) \cdot \Delta_i^\theta.\end{aligned}$$

Furthermore, differentiating these expressions further with respect to t_j or s_j will simply yield higher-order derivatives of $\phi(\mathbf{x}, \mathbf{u})$ and $\pi_h^\theta(\boldsymbol{\tau}_{1:h})$. Using the norm bounds on the gradient of $\phi(\mathbf{x}, \mathbf{u})$ and $\pi_h^\theta(\boldsymbol{\tau}_{1:h})$ given in Assumption 4 and Assumption 5, and the norm bound of $\phi(\mathbf{x}, \mathbf{u})$ given in Assumption 1, we can then bound

$$\|\nabla_A^{(i)} \nabla_{\boldsymbol{\theta}}^{(j)} \log f_{A, \theta}(\boldsymbol{\tau})\|_{\text{op}} \leq \text{poly}(\|A\|_{\text{op}}, B_\phi, L_\phi, L_\theta, \sigma_w^{-1}) \cdot \sum_{h=1}^H (1 + \|\mathbf{x}_{h+1}^\tau - A\phi(\mathbf{x}_h^\tau, \pi_h^\theta(\boldsymbol{\tau}_{1:h}))\|_2)$$

for $i, j \in \{0, 1, 2, 3, 4\}$ satisfying $1 \leq i + j \leq 4$ (where we have used the fact noted above that, to bound the operator norm of $\nabla_A^{(i)} \nabla_{\boldsymbol{\theta}}^{(j)} \log f_{A, \theta}(\boldsymbol{\tau})$, it suffices to bound the directional gradient in every direction). It follows that we can bound

$$\begin{aligned}\|\nabla_A^{(i)} \nabla_{\boldsymbol{\theta}}^{(j)} \mathcal{J}(\boldsymbol{\theta}; A)\|_{\text{op}} &\leq \text{poly}(\|A\|_{\text{op}}, B_\phi, L_\phi, L_\theta, \sigma_w^{-1}) \cdot \int \left(\sum_{h=1}^H (1 + \|\mathbf{x}_{h+1}^\tau - A\phi(\mathbf{x}_h^\tau, \pi_h^\theta(\boldsymbol{\tau}_{1:h}))\|_2) \right)^4 \cdot \text{cost}(\boldsymbol{\tau}) f_{A, \theta}(\boldsymbol{\tau}) d\boldsymbol{\tau} \\ &\stackrel{(a)}{\leq} \text{poly}(\|A\|_{\text{op}}, B_\phi, L_\phi, L_\theta, \sigma_w^{-1}) \cdot \sqrt{\int \text{cost}(\boldsymbol{\tau})^2 f_{A, \theta}(\boldsymbol{\tau}) d\boldsymbol{\tau}} \\ &\quad \cdot \sqrt{\int \left(\sum_{h=1}^H (1 + \|\mathbf{x}_{h+1}^\tau - A\phi(\mathbf{x}_h^\tau, \pi_h^\theta(\boldsymbol{\tau}_{1:h}))\|_2) \right)^8 f_{A, \theta}(\boldsymbol{\tau}) d\boldsymbol{\tau}} \\ &\stackrel{(b)}{\leq} \text{poly}(\|A\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x)\end{aligned}$$

where (a) follows from Cauchy-Schwarz, and (b) follows from Lemma A.1 and Assumption 2, since we have assumed $A \in \mathcal{B}_F(A_\star; r_{\text{cost}}(A_\star))$. $\square$

Proof of Proposition 5. Existence and Differentiability of $\boldsymbol{\theta}_\star$. By Lemma D.1 we have that $\mathcal{J}(\boldsymbol{\theta}; A)$ is four-times differentiable in its arguments. By the Implicit Function Theorem, since $\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)}) > 0$ by assumption, we have that there exists some $r'_\theta(A_\star) > 0$ and unique function $\boldsymbol{\theta}_\star(A)$ defined on $\mathcal{B}_F(A_\star; r'_\theta(A_\star))$ such that $\nabla_{\boldsymbol{\theta}} \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)} = 0$, and $\boldsymbol{\theta}_\star(A)$ is three-times differentiable (note that, while the Implicit Function Theorem is typically stated to give that the resulting function is only one-time differentiable, it can be extended to k -times differentiable, assuming the implicit equation is k -times differentiable (Dieudonné, 2011)).

By Lemma D.1 and the continuity of eigenvalues, it follows that for A close enough to $A_\star$, we have $\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)}) \geq \frac{1}{2} \lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)}) > 0$. We can therefore apply the Implicit Function Theorem as above to any A satisfying this, to get that there exists some unique $\tilde{\boldsymbol{\theta}}_\star(A')$ defined for all A' near A such that $\nabla_{\boldsymbol{\theta}} \mathcal{J}(\boldsymbol{\theta}; A')|_{\boldsymbol{\theta}=\tilde{\boldsymbol{\theta}}_\star(A')} = 0$ and $\tilde{\boldsymbol{\theta}}_\star(A')$ is differentiable. By the uniqueness of $\boldsymbol{\theta}_\star(A)$ on $\mathcal{B}_F(A_\star; r'_\theta(A_\star))$, it follows that any $\tilde{\boldsymbol{\theta}}_\star(A')$ defined in this way must be

identical to $\boldsymbol{\theta}_\star(A)$ on $\mathcal{B}_F(A_\star; r'_\theta(A_\star))$ (assuming the regions on which they are defined overlaps). We can therefore define $\boldsymbol{\theta}_\star(A)$ to simply be the extension of $\boldsymbol{\theta}_\star(A)$ to all such $\tilde{\boldsymbol{\theta}}_\star(A)$, defined for all A near $A_\star$ such that $\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)}) \geq \frac{1}{2} \lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)})$, and will have that $\boldsymbol{\theta}_\star(A)$ is three-times differentiable and satisfies $\nabla_{\boldsymbol{\theta}} \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)} = 0$ for all such A .

We then choose $r_\theta(A_\star)$ to be defined such that, for all $A \in \mathcal{B}_F(A_\star; r_\theta(A_\star))$, we have $\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)}) \geq \frac{1}{2} \lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)})$. By Lemma D.1, we know that $\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)$ is continuous and furthermore we know that eigenvalues are continuous. Using the gradient bounds given in Lemma D.1 to bound the Lipschitz constant of $\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)$, it follows that we can take

$$r_\theta(A_\star) = \text{poly} \left(\frac{1}{\lambda_{\min}(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)})}, \|A_\star\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x \right)^{-1}.$$

Bounding Norm of Gradients. Fix $A \in \mathcal{B}_F(A_\star; r_\theta(A_\star))$. We know that $\boldsymbol{\theta}_\star(A)$ satisfies

$$\nabla_{\boldsymbol{\theta}} \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)} = 0.$$

We wish to differentiate $\boldsymbol{\theta}_\star(A)$ with respect to A , and bound the magnitude of up to the third derivative. Similar to the proof of Lemma D.1, we let $A_t := A + t_1 \Delta_1^A + t_2 \Delta_2^A + t_3 \Delta_3^A$ for some Δ_i^A satisfying $\|\Delta_i^A\|_{\text{op}} \leq 1$. As noted in the proof of Lemma D.1, we have

$$\frac{d}{dt_i} \boldsymbol{\theta}_\star(A_t)|_{t=0} = \nabla_A \boldsymbol{\theta}_\star(A)[\Delta_i^A]$$

(and similarly for higher-order derivatives). Thus, to show the result, it suffices to show that $\boldsymbol{\theta}_\star(A_t)$ is differentiable in t_1, t_2, t_3 for all Δ_i^A , and to bound the magnitude of this derivative for all Δ_i^A with $\|\Delta_i^A\|_{\text{op}} \leq 1$. We have

$$\begin{aligned} \frac{d}{dt_1} \nabla_{\boldsymbol{\theta}} \mathcal{J}(\boldsymbol{\theta}; A_t)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_t)}|_{t=0} &= 0 \\ \implies \underbrace{\nabla_{A'} \nabla_{\boldsymbol{\theta}} \mathcal{J}(\boldsymbol{\theta}; A')|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A), A'=A}[\Delta_1^A]}_{=: G_1(A, \Delta_1^A)} + \nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)} \cdot \nabla_A \boldsymbol{\theta}_\star(A)[\Delta_1^A] &= 0 \end{aligned} \quad (\text{D.1})$$

which implies

$$\nabla_A \boldsymbol{\theta}_\star(A)[\Delta_1^A] = -(\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)})^{-1} \cdot G_1(A, \Delta_1^A)$$

which is well-defined since we have assumed that $\nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)}$ is full-rank, and $\mathcal{J}$ is differentiable in both its arguments by Lemma D.1. To compute the second derivative of $\boldsymbol{\theta}_\star$, we differentiate through (D.1) which gives

$$\begin{aligned} \frac{d}{dt_2} (G_1(A_t, \Delta_1^A) + \nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_t)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_t)} \cdot \nabla_A \boldsymbol{\theta}_\star(A_t)[\Delta_1^A])|_{t=0} &= 0 \\ \implies \underbrace{\frac{d}{dt_2} (G_1(A_t, \Delta_1^A) + \nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A_t)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_t)} \cdot \nabla_A \boldsymbol{\theta}_\star(A_t)[\Delta_1^A])|_{t=0}}_{=: G_2(A, \Delta_1^A, \Delta_2^A)} \\ + \nabla_{\boldsymbol{\theta}}^2 \mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)} \cdot \nabla_A^2 \boldsymbol{\theta}_\star(A)[\Delta_1^A, \Delta_2^A] &= 0. \end{aligned}$$

Note that $G_2(A, \Delta_1^A, \Delta_2^A)$ involves at most a third-order derivative of $\mathcal{J}(\theta; A)$ and first-order derivative of $\theta_*(A)$, both of which we know exist by Lemma D.1 and what we showed above. This then further implies

$$\nabla_A^2 \theta_*(A)[\Delta_1^A, \Delta_2^A] = -(\nabla_\theta^2 \mathcal{J}(\theta; A)|_{\theta=\theta_*(A)})^{-1} \cdot G_2(A, \Delta_1^A, \Delta_2^A),$$

which is well-defined since we have assumed that $\nabla_\theta^2 \mathcal{J}(\theta; A)|_{\theta=\theta_*(A)}$ is full-rank. Finally, we compute

$$\begin{aligned} & \frac{d}{dt_3} (G_2(A_t, \Delta_1^A, \Delta_2^A) + \nabla_\theta^2 \mathcal{J}(\theta; A_t)|_{\theta=\theta_*(A_t)} \cdot \nabla_A^2 \theta_*(A_t)[\Delta_1^A, \Delta_2^A]) \big|_{t=0} = 0 \\ \implies & \underbrace{\frac{d}{dt_3} (G_2(A_t, \Delta_1^A, \Delta_2^A) + \nabla_\theta^2 \mathcal{J}(\theta; A_t)|_{\theta=\theta_*(A_t)} \cdot \nabla_A^2 \theta_*(A_t)[\Delta_1^A, \Delta_2^A]) \big|_{t=0}}_{=: G_3(A, \Delta_1^A, \Delta_2^A, \Delta_3^A)} \\ & + \nabla_\theta^2 \mathcal{J}(\theta; A)|_{\theta=\theta_*(A)} \cdot \nabla_A^3 \theta_*(A)[\Delta_1^A, \Delta_2^A, \Delta_3^A] = 0. \end{aligned}$$

Note that $G_3(A, \Delta_1^A, \Delta_2^A, \Delta_3^A)$ involves at most a fourth-order derivative of $\mathcal{J}(\theta; A)$ and second-order derivative of $\theta_*(A)$, both of which we know exist by Lemma D.1 and what we showed above. We therefore have

$$\nabla_A^3 \theta_*(A)[\Delta_1^A, \Delta_2^A, \Delta_3^A] = -(\nabla_\theta^2 \mathcal{J}(\theta; A)|_{\theta=\theta_*(A)})^{-1} \cdot G_3(A, \Delta_1^A, \Delta_2^A, \Delta_3^A)$$

which is well-defined since we have assumed that $\nabla_\theta^2 \mathcal{J}(\theta; A)|_{\theta=\theta_*(A)}$ is full-rank. As each of these expressions is defined for all choice of Δ_i^A , the differentiability of $\theta_*(A)$ follows.

Note that the above expressions for $\nabla_A \theta_*(A)[\Delta_1^A]$, $\nabla_A^2 \theta_*(A)[\Delta_1^A, \Delta_2^A]$, and $\nabla_A^3 \theta_*(A)[\Delta_1^A, \Delta_2^A, \Delta_3^A]$ all depend on at most a fourth derivative of $\mathcal{J}(\theta; A)$, as well as $(\nabla_\theta^2 \mathcal{J}(\theta; A)|_{\theta=\theta_*(A)})^{-1}$. The norm bounds are then a direct consequence of Lemma D.1.

□

Proof of Proposition 6. By Lemma D.1 we have that $\mathcal{J}(\theta; A)$ is four-times differentiable in its arguments. Since we have assumed $\nabla_\theta^2 \mathcal{J}(\theta; A_*)|_{\theta=\theta_*(A_*)} \succ 0$, by the Implicit Function Theorem (Dieudonné, 2011), it follows that there exists some $r'_\theta > 0$ and mapping $\tilde{\theta}(A)$ such that, for all $A \in \mathcal{B}_F(A_*; r'_\theta)$, $\nabla_\theta \mathcal{J}(\theta; A)|_{\theta=\tilde{\theta}(A)} = 0$, and $\tilde{\theta}(A)$ is three-times differentiable.

Our goal is now to show that $\tilde{\theta}(A) = \theta_*(A)$ for A close enough to A_* . By the continuity of eigenvalues, $\mathcal{J}(\theta; A)$, and $\tilde{\theta}(A)$, we have that there exists some r and r''_θ such that, for all $\theta \in \mathcal{B}_F(\theta_*(A_*); r)$ and $A \in \mathcal{B}_F(A_*; r''_\theta)$, we have $\nabla_\theta^2 \mathcal{J}(\theta; A) \succ 0$ and, furthermore, $\tilde{\theta}(A) \in \mathcal{B}_2(\theta_*(A_*); r/2)$ for all $A \in \mathcal{B}_F(A_*; r''_\theta)$. This implies that $\tilde{\theta}(A)$ is strict local minimum of $\mathcal{J}(\theta; A)$ and, in particular, that

$$\mathcal{J}(\theta; A) > \mathcal{J}(\tilde{\theta}(A); A), \quad \forall \theta \in \mathcal{B}_2(\theta_*(A_*); r), \theta \neq \tilde{\theta}(A).$$

Let $\Delta^* := \min_{\theta \notin \mathcal{B}_2(\theta_*(A_*); r)} \mathcal{J}(\theta; A_*) - \mathcal{J}(\theta_*(A_*); A_*)$ and note that, since we have assumed the global minimum of $\mathcal{J}(\theta; A_*)$ is unique, we have $\Delta^* > 0$.

Fix some $A \in \mathcal{B}_F(A_*; r''_\theta)$ and assume that $\tilde{\theta}(A)$ is not the global minimum of $\mathcal{J}(\theta; A)$. This implies that $\theta_*(A)$, the global minimum of $\mathcal{J}(\theta; A)$, is outside of $\mathcal{B}_2(\theta_*(A_*); r)$. Furthermore, by the continuity of $\mathcal{J}(\theta; A)$, we have, for some $L, L' > 0$,

$$\begin{aligned} \mathcal{J}(\theta_*(A); A_*) & \leq \mathcal{J}(\theta_*(A); A) + L\|A - A_*\|_F \\ & \leq \mathcal{J}(\tilde{\theta}(A); A) + L\|A - A_*\|_F \end{aligned}$$

$$\begin{aligned}
&\leq \mathcal{J}(\tilde{\boldsymbol{\theta}}(A); A_\star) + 2L\|A - A_\star\|_F \\
&\leq \mathcal{J}(\boldsymbol{\theta}_\star(A_\star); A_\star) + 2L\|A - A_\star\|_F + L\|\tilde{\boldsymbol{\theta}}(A) - \boldsymbol{\theta}_\star(A_\star)\|_2 \\
&\leq \mathcal{J}(\boldsymbol{\theta}_\star(A_\star); A_\star) + L'\|A - A_\star\|_F.
\end{aligned}$$

This implies that

$$L'r''_{\boldsymbol{\theta}} \geq L'\|A - A_\star\|_F \geq \mathcal{J}(\boldsymbol{\theta}_\star(A); A_\star) - \mathcal{J}(\boldsymbol{\theta}_\star(A_\star); A_\star) \geq \Delta^\star.$$

However, for $r''_{\boldsymbol{\theta}}$ small enough, this is a contradiction. Thus, it follows that $\tilde{\boldsymbol{\theta}}(A)$ is the global minimum of $\mathcal{J}(\boldsymbol{\theta}; A)$, so $\tilde{\boldsymbol{\theta}}(A) = \boldsymbol{\theta}_\star(A)$.

The result then follows since we already have that $\tilde{\boldsymbol{\theta}}(A)$ is three-times differentiable and satisfies $\nabla_{\boldsymbol{\theta}}\mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\tilde{\boldsymbol{\theta}}(A)} = 0$, and by taking $r_{\boldsymbol{\theta}}(A_\star)$ to be the minimum of $r'_{\boldsymbol{\theta}}$ and $r''_{\boldsymbol{\theta}}$. The boundedness of $L_{\pi_\star}$ follows as in the proof of Proposition 5. $\square$

Proof of Proposition 7. Note that, by convexity and the KKT conditions, the solutions to $\arg \min_{\boldsymbol{\theta} \in \mathbb{R}^{d_\theta}} \mathcal{J}(\boldsymbol{\theta}; A)$ are described by

$$\nabla_{\boldsymbol{\theta}}\mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)} = 0.$$

Thus, an equivalent definition for $\boldsymbol{\theta}_\star(A)$ is that it satisfies $\nabla_{\boldsymbol{\theta}}\mathcal{J}(\boldsymbol{\theta}; A)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A)} = 0$. Assumption 13 can then be shown to hold by an argument analogous to Proposition 5. $\square$

Proof of Lemma D.2. Let $A(t) = t\hat{A} + (1-t)A_\star$ and $g(t) := \mathcal{J}(\boldsymbol{\theta}_\star(A(t)); A_\star)$. By Lemma D.1 and under Assumption 5, we have that both $\mathcal{J}(\boldsymbol{\theta}; A)$ and $\boldsymbol{\theta}_\star(A)$ are three-times differentiable for all $A = t\hat{A} + (1-t)A_\star, t \in [0, 1]$, so it follows that $g(t)$ is three-times differentiable in t . We can therefore apply Taylor's Theorem to expand $g(1)$ about the point $t = 0$ to get:

$$\begin{aligned}
g(1) &= g(0) + \nabla_{\boldsymbol{\theta}}\mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)} \cdot \nabla_A\boldsymbol{\theta}_\star(A)|_{A=A_\star}[\hat{A} - A_\star] \\
&\quad + \nabla_A\boldsymbol{\theta}_\star(A)|_{A=A_\star}^\top \nabla_{\boldsymbol{\theta}}^2\mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)} \nabla_A\boldsymbol{\theta}_\star(A)|_{A=A_\star}[\hat{A} - A_\star, \hat{A} - A_\star] \\
&\quad + \nabla_{\boldsymbol{\theta}}\mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)} \cdot \nabla_A^2\boldsymbol{\theta}_\star(A)|_{A=A_\star}[\hat{A} - A_\star, \hat{A} - A_\star] \\
&\quad + \nabla_A^3\mathcal{J}(\boldsymbol{\theta}_\star(A); A_\star)|_{A=A'}[\hat{A} - A_\star, \hat{A} - A_\star, \hat{A} - A_\star]
\end{aligned}$$

where $A' = A(t')$ for some $t' \in [0, 1]$. Under Assumption 13, we have that $\nabla_{\boldsymbol{\theta}}\mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)} = 0$, which implies that, plugging in the definition of $g(1)$ and $g(0)$,

$$\begin{aligned}
\mathcal{J}(\boldsymbol{\theta}_\star(\hat{A}); A_\star) &= \mathcal{J}(\boldsymbol{\theta}_\star(A_\star); A_\star) + \nabla_A\boldsymbol{\theta}_\star(A)|_{A=A_\star}^\top \nabla_{\boldsymbol{\theta}}^2\mathcal{J}(\boldsymbol{\theta}; A_\star)|_{\boldsymbol{\theta}=\boldsymbol{\theta}_\star(A_\star)} \nabla_A\boldsymbol{\theta}_\star(A)|_{A=A_\star}[\hat{A} - A_\star, \hat{A} - A_\star] \\
&\quad + \nabla_A^3\mathcal{J}(\boldsymbol{\theta}_\star(A); A_\star)|_{A=A'}[\hat{A} - A_\star, \hat{A} - A_\star, \hat{A} - A_\star].
\end{aligned}$$

We can bound

$$|\nabla_A^3\mathcal{J}(\boldsymbol{\theta}_\star(A); A_\star)|_{A=A'}[\hat{A} - A_\star, \hat{A} - A_\star, \hat{A} - A_\star]| \leq \|\nabla_A^3\mathcal{J}(\boldsymbol{\theta}_\star(A); A_\star)|_{A=A'}\|_{\text{op}} \cdot \|\hat{A} - A_\star\|_{\text{op}}^3.$$

The expression for $\nabla_A^3\mathcal{J}(\boldsymbol{\theta}_\star(A); A_\star)$ contains up to the third derivative of both $\mathcal{J}(\boldsymbol{\theta}; A_\star)$ and $\boldsymbol{\theta}_\star(A)$. By Lemma D.1 and under Assumption 13, since $A' \in \mathcal{B}_F(A_\star; \min\{r_{\text{cost}}(A_\star), r_{\boldsymbol{\theta}}(A_\star)\})$ by construction, we can then bound

$$\|\nabla_A^3\mathcal{J}(\boldsymbol{\theta}_\star(A); A_\star)|_{A=A'}\|_{\text{op}} \leq \text{poly}(L_{\pi_\star}, \|A_\star\|_{\text{op}}, B_\phi, L_\phi, L_{\boldsymbol{\theta}}, L_{\text{cost}}, \sigma_w^{-1}, H, d_x).$$

The result follows by the definition of $\mathcal{H}(A_\star)$. $\square$

Proof of Lemma D.3. Recall that $\mathcal{H}(\hat{A}) = \nabla_A^2 \mathcal{J}(\theta_*(A); \hat{A})|_{A=\hat{A}}$. To prove this, we will use that this is differentiable by Lemma D.1, and will apply Taylor's Theorem.

First, note that by Taylor's Theorem we have

$$\nabla_A^2 \mathcal{J}(\theta_*(A); \hat{A})|_{A=\hat{A}} = \nabla_A^2 \mathcal{J}(\theta_*(A); \hat{A})|_{A=A_*} + \nabla_{A'}^3 \mathcal{J}(\theta_*(A); \hat{A})|_{A=A_*} [\hat{A} - A_*]$$

for $A' = t\hat{A} + (1-t)A_*$ for some $t \in [0, 1]$. The third derivative of $\mathcal{J}(\theta_*(A); \hat{A})$ will involve up to the third derivative of both $\mathcal{J}(\theta; \hat{A})$ and $\theta_*(A)$, so using Lemma D.1 and Assumption 13, since $A' \in \mathcal{B}_F(A_*; \min\{r_{\text{cost}}(A_*), r_\theta(A_*)\})$ by assumption, we can bound

$$\|\nabla_{A'}^3 \mathcal{J}(\theta_*(A); \hat{A})|_{A=A'} [\hat{A} - A_*]\|_{\text{op}} \leq \text{poly}(L_{\pi_*}, \|A_*\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x) \cdot \|\hat{A} - A_*\|_{\text{op}}.$$

Next, we wish to relate $\nabla_A^2 \mathcal{J}(\theta_*(A); \hat{A})|_{A=A_*}$ to $\nabla_A^2 \mathcal{J}(\theta_*(A); A_*)|_{A=A_*} = \mathcal{H}(A_*)$. Again applying Taylor's Theorem, we have

$$\nabla_A^2 \mathcal{J}(\theta_*(A); \hat{A})|_{A=A_*} = \nabla_A^2 \mathcal{J}(\theta_*(A); A_*)|_{A=A_*} + \nabla_{A'} \nabla_A^2 \mathcal{J}(\theta_*(A); A')|_{A=A_*, A'=A''} [\hat{A} - A_*]$$

for $A'' = t\hat{A} + (1-t)A_*$ for some $t \in [0, 1]$. By Lemma D.1 and Assumption 13, we can bound

$$\begin{aligned} & \|\nabla_{A'} \nabla_A^2 \mathcal{J}(\theta_*(A); A')|_{A=A_*, A'=A''} [\hat{A} - A_*]\|_{\text{op}} \\ & \leq \text{poly}(L_{\pi_*}, \|A_*\|_{\text{op}}, L_\phi, L_\theta, L_{\text{cost}}, \sigma_w^{-1}, H, d_x) \cdot \|\hat{A} - A_*\|_{\text{op}}. \end{aligned}$$

The result follows. $\square$

Lemma D.4. *Under Assumptions 1, 2, 4, 5 and 13, for all $A \in \mathcal{B}_F(A_*; \min\{r_{\text{cost}}(A_*), r_\theta(A_*)\})$, we can bound*

$$\|\mathcal{H}(A)\|_{\text{op}} \leq \text{poly}(\|A_*\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, L_{\pi_*}, \sigma_w^{-1}, H, d_x)$$

Proof. Recall that $\mathcal{H}(\hat{A}) = \nabla_A^2 \mathcal{J}(\theta_*(A); \hat{A})|_{A=\hat{A}}$. The bound then follows from Lemma D.1 and Assumption 13. $\square$

E High-Probability Regret Bounds in Nonlinear Systems

In this section, we modify the proof the main result of Kakade et al. (2020) slightly to show a high probability regret bound for LC³. For the sake of brevity, we omit details that are identical to the proof given in Kakade et al. (2020). We will need the following assumption.

Assumption 14 (Bounded Cost). *We assume that, for all trajectories τ , we have $\text{cost}(\tau) \leq c_{\text{max}}$.*

We adopt the notation used in this work, modifying somewhat the notation from Kakade et al. (2020). In particular, we let $\mathcal{J}(\pi; A)$ denote the expected cost of playing policy π under system A , and we set

$$\Sigma_t = \sum_{s=1}^t \sum_{h=1}^H \phi(x_h^s, u_h^s) \phi(x_h^s, u_h^s)^\top + \lambda I$$

denote the covariates obtained by the first t episodes of LC³ (plus a regularizer). We let π^t denote the policy played at episode t of LC³. For a policy set Π , we define regret as

$$\mathcal{R}_T(\Pi) := \sum_{t=1}^T \mathcal{J}(\pi^t; A_*) - T \cdot \min_{\pi \in \Pi} \mathcal{J}(\pi; A_*).$$

We will also denote $\pi_\star := \arg \min_{\pi \in \Pi} \mathcal{J}(\pi; A_\star)$.

In addition to these notational changes, we modify LC^3 slightly to use the parameter

$$\beta^t := \sqrt{\lambda} B_A + \sqrt{8d_{\mathbf{x}} \log 5 + 8 \log(T \det(\mathbf{\Sigma}_t) \det(\mathbf{\Sigma}_0)^{-1}/\delta)}$$

in the construction of the confidence set, BALL^t .

Besides the aforementioned changes, in the following proofs we adopt the same notation as [Kakade et al. \(2020\)](#). We have the following result.

Theorem 6. *Under Assumptions 1 and 14 and with any policy class Π , with probability at least $1 - \delta$, LC^3 has regret bounded as*

$$\mathcal{R}_T(\Pi) \leq C \cdot c_{\max} H \sqrt{d_\phi \cdot (d_\phi + d_{\mathbf{x}} + B_A + \log \frac{1}{\delta}) \cdot T \cdot \log(1 + B_\phi H T / \sigma_{\mathbf{w}})}$$

for a universal constant C .

Proof of Theorem 6. By Lemma E.2, we have that the event $\mathcal{E}_1$ holds with probability at least $1 - \delta$. We therefore assume $\mathcal{E}_1$ holds for the remainder of the proof.

By the definition of the confidence set in LC^3 , on $\mathcal{E}_1$ we have that $A_\star$ is the in confidence set for all $t \leq T$. It follows that on $\mathcal{E}_1$,

$$\begin{aligned} \mathcal{R}_T &= \sum_{t=1}^T [\mathcal{J}(\pi^t; A_\star) - \mathcal{J}(\pi_\star; A_\star)] \\ &\stackrel{(a)}{\leq} \sum_{t=1}^T [\mathcal{J}(\pi^t; A_\star) - \mathcal{J}(\pi_\star; \hat{A}^t)] \\ &\stackrel{(b)}{\leq} \sum_{t=1}^T c_{\max} \cdot \mathbb{E}_{A_\star, \pi^t} \left[\sum_{h=1}^H \min \left\{ \frac{1}{\sigma_{\mathbf{w}}} \|(A_\star - \hat{A}^t) \cdot \phi(\mathbf{x}_h, \mathbf{u}_h)\|_2, 1 \right\} \right] \end{aligned} \quad (\text{E.1})$$

where (a) follows from the optimistic property of LC^3 when $A_\star \in \text{BALL}^t$, and (b) follows from Lemma E.1. On $\mathcal{E}_1$, we have

$$\begin{aligned} \|(A_\star - \hat{A}^t) \phi(\mathbf{x}_h, \mathbf{u}_h)\|_2 &\leq \|(A_\star - \hat{A}^t) \mathbf{\Sigma}_t^{1/2}\|_2 \|\mathbf{\Sigma}_t^{-1/2} \phi(\mathbf{x}_h, \mathbf{u}_h)\|_2 \\ &\leq \left(\|(A_\star - \bar{A}^t) \mathbf{\Sigma}_t^{1/2}\|_2 + \|(\bar{A}^t - \hat{A}^t) \mathbf{\Sigma}_t^{1/2}\|_2 \right) \cdot \|\mathbf{\Sigma}_t^{-1/2} \phi(\mathbf{x}_h, \mathbf{u}_h)\|_2 \\ &\leq 2\beta^t \|\phi(\mathbf{x}_h, \mathbf{u}_h)\|_{\mathbf{\Sigma}_t^{-1}} \end{aligned}$$

where the last inequality follows from the definition of BALL^t since $\hat{A}^t \in \text{BALL}^t$ by construction, and by the definition of $\mathcal{E}_1$. This gives

$$(\text{E.1}) \leq \sum_{t=1}^T c_{\max} \cdot \mathbb{E}_{A_\star, \pi^t} \left[\sum_{h=1}^H \min \left\{ \frac{2\beta^t}{\sigma_{\mathbf{w}}} \|\phi(\mathbf{x}_h, \mathbf{u}_h)\|_{\mathbf{\Sigma}_t^{-1}}, 1 \right\} \right].$$

By Lemma E.3, with probability $1 - \delta$ we can bound this as

$$\leq \underbrace{\frac{2c_{\max}\beta^T}{\sigma_{\mathbf{w}}} \cdot \sum_{t=1}^T \sum_{h=1}^H \min \left\{ \|\phi(\mathbf{x}_h^t, \mathbf{u}_h^t)\|_{\mathbf{\Sigma}_t^{-1}}, 1 \right\}}_{(a)} + 4c_{\max} H \sqrt{T \log 1/\delta}.$$

By Cauchy-Schwarz, we can bound (a) as

$$(a) \leq \frac{2c_{\max}\beta^T}{\sigma_{\mathbf{w}}} \cdot \sqrt{T} \sqrt{\sum_{t=1}^T \sum_{h=1}^H \min \left\{ \|\phi(\mathbf{x}_h^t, \mathbf{u}_h^t)\|_{\Sigma_t^{-1}}^2, 1 \right\}}.$$

We have

$$\begin{aligned} \sum_{t=1}^T \sum_{h=1}^H \min \left\{ \|\phi(\mathbf{x}_h^t, \mathbf{u}_h^t)\|_{\Sigma_t^{-1}}^2, 1 \right\} H &\leq \sum_{t=1}^T \min \left\{ \sum_{h=1}^H \|\phi(\mathbf{x}_h^t, \mathbf{u}_h^t)\|_{\Sigma_t^{-1}}^2, 1 \right\} \\ &\leq 2H \log(\det(\Sigma_T) \det(\Sigma_0)^{-1}) \end{aligned}$$

where the last inequality uses Lemma B.6 of [Kakade et al. \(2020\)](#). Putting all of this together, we have shown that with probability at least $1 - 2\delta$, we have

$$\mathcal{R}_T \leq \frac{2c_{\max}\beta^T}{\sigma_{\mathbf{w}}} \cdot \sqrt{T} \cdot \sqrt{2H \log(\det(\Sigma_T) \det(\Sigma_0)^{-1})} + 4c_{\max}H \sqrt{T \log 1/\delta}.$$

It remains to bound β^T and $\log(\det(\Sigma_T) \det(\Sigma_0)^{-1})$. We have $\Sigma_0 = \lambda I$, so $\det(\Sigma_0) = \lambda^{d_\phi}$. Furthermore, if $\|\phi(\mathbf{x}, \mathbf{u})\|_2 \leq B_\phi$, then we can bound $\det(\Sigma_T) \leq (\lambda + B_\phi^2 TH)^{d_\phi}$. Putting this together we have

$$\log(\det(\Sigma_T) \det(\Sigma_0)^{-1}) \leq d_\phi \cdot \log(1 + B_\phi^2 TH/\lambda).$$

Recalling that

$$\beta^T = \sqrt{\lambda} B_A + \sigma_{\mathbf{w}} \sqrt{8d_{\mathbf{x}} \log 5 + 8 \log(T \det(\Sigma_T) \det(\Sigma_0)^{-1}/\delta)}$$

we can similarly bound

$$\begin{aligned} \beta^T / \sigma_{\mathbf{w}} &\leq \sqrt{\lambda} B_A / \sigma_{\mathbf{w}} + \sqrt{8d_{\mathbf{x}} \log 5 + 8d_\phi \cdot \log(1 + B_\phi^2 TH/\lambda) + 8 \log(T/\delta)} \\ &\leq \sqrt{\lambda} B_A / \sigma_{\mathbf{w}} + c \sqrt{d_{\mathbf{x}} + d_\phi \log(1 + B_\phi^2 TH/\lambda) + \log 1/\delta}. \end{aligned}$$

Choosing $\lambda = \sigma_{\mathbf{w}}^2$ completes the proof. $\square$

E.1 Supporting Lemmas

Lemma E.1. *Under Assumption 14, we can bound*

$$\mathcal{J}(\pi; A_\star) - \mathcal{J}(\pi; A) \leq c_{\max} \cdot \mathbb{E}_{A_\star, \pi} \left[\sum_{h=1}^H \min \left\{ \frac{1}{\sigma_{\mathbf{w}}} \|(A_\star - A)\phi(\mathbf{x}_h, \mathbf{u}_h)\|_2, 1 \right\} \right].$$

Proof. Following the proof of Lemma B.3 of [Kakade et al. \(2020\)](#), and adopting the same notation, we have

$$\mathcal{J}(\pi; A_\star) - \mathcal{J}(\pi; A) \leq \sum_{h=1}^H \mathbb{E}_{A_\star, \pi} \left[\sqrt{A_h} \min \left\{ \frac{1}{\sigma_{\mathbf{w}}} \|(A_\star - A)\phi(\mathbf{x}_h, \mathbf{u}_h)\|_2, 1 \right\} \right].$$

Under Assumption 14 we have $A_h \leq c_{\max}^2$. Plugging this in gives the result. $\square$

Lemma E.2. Let $\beta^t := \sqrt{\lambda}B_A + \sigma_w \sqrt{8d_x \log 5 + 8 \log(T \det(\Sigma_t) \det(\Sigma_0)^{-1}/\delta)}$ and let $\mathcal{E}_1$ denote the event

$$\mathcal{E}_1 := \left\{ \forall t \leq T : \left\| (\bar{A}^t - A_\star) \Sigma_t^{1/2} \right\|_{\text{op}} \leq \beta^t \right\}.$$

Then running LC^3 we have $\mathbb{P}_{A_\star}[\mathcal{E}_1] \geq 1 - \delta$.

Proof. The proof of Lemma B.5 of [Kakade et al. \(2020\)](#) shows that with probability at least $1 - \delta$,

$$\left\| (\bar{A}^t - A_\star) \Sigma_t^{1/2} \right\|_{\text{op}} \leq \sqrt{\lambda} \|A_\star\|_{\text{op}} + \sigma_w \sqrt{8d_x \log 5 + 8 \log(\det(\Sigma_t) \det(\Sigma_0)^{-1}/\delta)}.$$

The result then follows from this, since $\|A_\star\|_{\text{op}} \leq B_A$, and a union bound. $\square$

Lemma E.3. With probability $1 - \delta$, we have

$$\begin{aligned} \sum_{t=1}^T \mathbb{E}_{A_\star, \pi^t} \left[\sum_{h=1}^H \min \left\{ \frac{2\beta^t}{\sigma_w} \|\phi(\mathbf{x}_h, \mathbf{u}_h)\|_{\Sigma_t^{-1}}, 1 \right\} \right] &\leq \frac{2\beta^T}{\sigma_w} \sum_{t=1}^T \sum_{h=1}^H \min \left\{ \|\phi(\mathbf{x}_h^t, \mathbf{u}_h^t)\|_{\Sigma_t^{-1}}, 1 \right\} \\ &\quad + 4H \sqrt{T \log 1/\delta}. \end{aligned}$$

Proof. This is an immediate consequence of Azuma-Hoeffding, since $\sum_{h=1}^H \min \left\{ \frac{2\beta^t}{\sigma_w} \|\phi(\mathbf{x}_h, \mathbf{u}_h)\|_{\Sigma_t^{-1}}, 1 \right\} \leq H$ almost surely, and from upper bounding

$$\min \left\{ \frac{2\beta^t}{\sigma_w} \|\phi(\mathbf{x}_h^t, \mathbf{u}_h^t)\|_{\Sigma_t^{-1}}, 1 \right\} \leq \frac{2\beta^T}{\sigma_w} \min \left\{ \|\phi(\mathbf{x}_h^t, \mathbf{u}_h^t)\|_{\Sigma_t^{-1}}, 1 \right\}.$$

$\square$

F Lower Bounds on Learning in Nonlinear Systems

In this section, we assume that $\theta_\star$ and $\pi_\star$ correspond to the global minimizer:

$$\theta_\star(A) := \arg \min_{\theta \in \mathbb{R}^{d_\theta}} \mathcal{J}(\theta; A), \quad \pi_\star(A) := \arg \min_{\pi \in \Pi^\star} \mathcal{J}(\pi; A). \quad (\text{F.1})$$

Here we formally state the additional assumptions needed in [Section 4.2](#), and provide a formal version of [Theorem 2](#).

Assumption 15. There exists some $r_\mu(A_\star) > 0$ such that, for all $A \in \mathcal{B}_F(A_\star, r_\mu(A_\star))$, $\pi_\star(A)$ is unique and, furthermore, there exists some $\mu > 0$ such that

$$\mathcal{J}(\theta; A) \geq \mathcal{J}(\theta_\star(A); A) + \frac{\mu}{2} \|\theta - \theta_\star(A)\|_2^2.$$

Assumption 15 will be satisfied in cases where $\mathcal{J}(\pi^\theta; A)$ is strongly convex in θ , but may hold even when this is not the case. Intuitively, it requires that our controller class is not overparameterized—moving θ away from its optimal value will cause the loss to increase. We will additionally make the following regularity assumptions on policies in Π_{exp} and their induced covariates set, Ω .

Assumption 16. There exists some $\underline{\lambda} > 0$ such that, for each $\Lambda \in \Omega$, we have $\lambda_{\min}(\Lambda) \geq \underline{\lambda}$.

Assumption 16 requires that *every* exploration policy we consider excites all directions in ϕ space (in contrast, Assumption 3 only assumes there *exists* some distribution over policies in Π_{exp} which excite all directions). We remark that this assumption is relatively mild if Assumption 3 holds. As we show in Appendix C.2, under Assumption 3, a mixture over policies, ω , satisfying $\lambda_{\min}(\mathbb{E}_{\pi \sim \omega}[\mathbf{A}_\pi]) > 0$ can be learned using only a number of samples scaling polynomially in problem parameters. Given ω , a policy class Π_{exp} satisfying Assumption 16 can be obtained by simply mixing ω with every other exploration policy. We are now ready to state our main lower bound.

Theorem 7 (Formal Version of Theorem 2). *Under Assumptions 1, 2, 4, 5, 13, 15 and 16 and if $\pi_\star$ is defined as in (F.1), as long as $T \geq C_{\text{lb}}$, for any $\omega_{\text{exp}} \in \Delta_{\Pi_{\text{exp}}}$, we have*

$$\min_{\hat{\pi}} \max_{A \in \mathcal{B}_T} \mathbb{E}_{\mathcal{D}_T \sim A, \omega_{\text{exp}}} [\mathcal{J}(\hat{\pi}(\mathcal{D}_T); A) - \mathcal{J}(\pi_\star(A); A)] \geq \frac{\sigma_w^2}{3T} \cdot \min_{\mathbf{A} \in \Omega} \text{tr}(\mathcal{H}(\mathbf{A}_\star) \check{\mathbf{A}}^{-1}) - \frac{C_{\text{lb}}}{T^{5/4}}$$

for $\mathcal{B}_T := \{A : \|A - A_\star\|_F^2 \leq 5d_x d_\phi / (\lambda d_x T H)^{5/6}\}$, $\mathbb{E}_{\mathcal{D}_T \sim A, \omega_{\text{exp}}}[\cdot] = \mathbb{E}_{\pi \sim \omega_{\text{exp}}}[\mathbb{E}_{\mathcal{D}_T \sim A, \pi}[\cdot]]$ denotes the expectation over trajectories generated by running policies π drawn according to ω_{exp} on system A for T episodes, $\hat{\pi}$ any mapping from observations to policies in $\Pi^\star$, and

$$C_{\text{lb}} := \text{poly}\left(d_\phi, d_x, H, \|A_\star\|_{\text{op}}, B_\phi, L_\phi, L_\theta, L_{\text{cost}}, L_{\pi_\star}, \sigma_w, \sigma_w^{-1}, \frac{1}{\lambda}, \frac{1}{\mu}, \frac{1}{r_{\text{cost}}(A_\star)}, \frac{1}{r_\theta(A_\star)}, \frac{1}{r_\mu(A_\star)}\right).$$

Proof of Theorem 7. This proof follows immediately from Lemma F.1, by lower bounding the right-hand side of (F.2) by the min over all policies in Π . $\square$

Lemma F.1. *Under Assumptions 1, 2, 4, 5, 13, 15 and 16 and if $\theta_\star$ is defined as in (F.1), as long as*

$$T \geq \text{poly}(\|A_\star\|_{\text{op}}, L_\phi, L_\theta, L_{\pi_\star}, L_{\text{cost}}, \sigma_w, \sigma_w^{-1}, B_\phi, H, d_x, d_\phi, \lambda^{-1}, \mu^{-1}, r_{\text{cost}}(A_\star)^{-1}, r_\theta(A_\star)^{-1}, r_\mu(A_\star)^{-1}),$$

for any $\omega_{\text{exp}} \in \Delta_\Pi$, we have

$$\min_{\hat{\theta}} \max_{A \in \mathcal{B}_T} \mathbb{E}_{\mathcal{D}_T \sim A, \omega_{\text{exp}}} [\mathcal{J}(\hat{\theta}(\mathcal{D}_T); A) - \mathcal{J}(\theta_\star(A); A)] \geq \frac{\sigma_w^2}{3T} \cdot \text{tr}(\mathcal{H}(A_\star) \mathbb{E}_{\pi \sim \omega_{\text{exp}}}[\check{\mathbf{A}}_\pi]^{-1}) - \frac{C_{\text{lb}}}{T^{5/4}} \quad (\text{F.2})$$

for $\mathcal{B}_T := \{A : \|A - A_\star\|_F^2 \leq 5d_x d_\phi / (\lambda T H)^{5/6}\}$, where $\mathbb{E}_{\mathcal{D}_T \sim A, \omega_{\text{exp}}}[\cdot] = \mathbb{E}_{\pi \sim \omega_{\text{exp}}}[\mathbb{E}_{\mathcal{D}_T \sim A, \pi}[\cdot]]$ denotes the expectation over trajectories generated by running policies π drawn according to ω_{exp} on system A for T episodes, and

$$C_{\text{lb}} := \text{poly}(\|A_\star\|_{\text{op}}, L_\phi, L_\theta, L_{\pi_\star}, L_{\text{cost}}, \sigma_w, \sigma_w^{-1}, B_\phi, H, d_x, d_\phi, \lambda^{-1}, \mu^{-1}, r_{\text{cost}}(A_\star)^{-1}, r_\theta(A_\star)^{-1}, r_\mu(A_\star)^{-1}).$$

Proof. This result is a direct consequence of Theorem 6.1 of Wagenmaker et al. (2021)—to obtain the result we must only verify that the assumptions of this result are met. We verify each assumption below.

Verifying Assumption 3 of Wagenmaker et al. (2021). Part 1 of Assumption 3 of Wagenmaker et al. (2021) is met by Assumption 15 within diameter $r_\mu(A_\star)$. Furthermore, under Assumptions 1, 2, 4, 5 and 13 and by Lemma D.1, the additional parts of Assumption 3 of Wagenmaker et al. (2021) are also met with diameter $\min\{r_{\text{cost}}(A_\star), r_\theta(A_\star)\}$ and smoothness constant $\text{poly}(\|A_\star\|_{\text{op}}, L_\phi, L_\theta, L_{\pi_\star}, L_{\text{cost}}, \sigma_w^{-1}, H, d_x)$.

Verifying Assumption 4 and Assumption 5 of Wagenmaker et al. (2021). Assumption 4 of Wagenmaker et al. (2021) is immediately met by Assumption 16. Furthermore, Assumption 5 is met by Lemma F.2 with $c_{\text{cov}} = 1$, $L_{\text{cov}}(\theta_*, \gamma^2) = \frac{H^2 B_\phi^3}{\sigma_w^2} \cdot \text{poly}(d_x)$, and $C_{\text{cov}} = 0$.

Given that these assumptions are met, the result follows noting that, if we run for T episodes, then the effective horizon is $d_x T H$ (using the mapping from the setting of (1.1) to the martingale regression setting described in Appendix A.1.1). Note that the final bound scales with $\frac{1}{T}$ instead of $\frac{1}{d_x T H}$ as we are able to bring the $d_x H$ factor into the $\check{\mathbf{\Lambda}}_{\pi_{\text{exp}}}$ term, since $\mathbf{\Lambda}_{\pi_{\text{exp}}}$ is not normalized by $d_x H$. $\square$

Lemma F.2. *Under Assumption 1, for any policy distribution $\omega \in \Delta_{\Pi_{\text{exp}}}$ and A, A' , we have*

$$\mathbb{E}_{\pi \sim \omega}[\mathbf{\Lambda}_{A, \pi}] \preceq \mathbb{E}_{\pi \sim \omega}[\mathbf{\Lambda}_{A', \pi}] + \frac{H^2 B_\phi^3}{\sigma_w^2} \cdot \text{poly}(d_x) \cdot \|A - A'\|_F \cdot I.$$

Proof. We will prove that the desired bound follows for a particular $\pi \in \Pi_{\text{exp}}$, which immediately implies that it holds for $\omega \in \Delta_{\Pi_{\text{exp}}}$. By definition we have

$$\mathbf{\Lambda}_{A, \pi} = \int \left(\sum_{h=1}^H \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau) \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau)^\top \right) \cdot f_{A, \pi}(\tau) d\tau$$

and

$$f_{A, \pi}(\tau) = \prod_{h=1}^H f_A(\mathbf{x}_{h+1}^\tau \mid \mathbf{x}_h^\tau, \mathbf{u}_h^\tau) \pi_h(\mathbf{u}_h^\tau \mid \mathbf{x}_h^\tau).$$

Fix some $\mathbf{v} \in \mathcal{S}^{d_\phi - 1}$, and note that, given the expression above, we have

$$\mathbf{v}^\top \mathbf{\Lambda}_{A, \pi} \mathbf{v} = \int \sum_{h=1}^H (\mathbf{v}^\top \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau))^2 \cdot f_{A, \pi}(\tau) d\tau.$$

It follows that

$$\begin{aligned} \nabla_A \mathbf{v}^\top \mathbf{\Lambda}_{A, \pi} \mathbf{v} &= \int \sum_{h=1}^H (\mathbf{v}^\top \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau))^2 \cdot \nabla_A f_{A, \pi}(\tau) d\tau \\ &= \int \sum_{h=1}^H (\mathbf{v}^\top \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau))^2 \cdot f_{A, \pi}(\tau) \nabla_A \log f_{A, \pi}(\tau) d\tau. \end{aligned}$$

As in the proof of Lemma D.1, we have, for any Δ ,

$$\nabla_A \log f_{A, \pi}(\tau)[\Delta] = \sum_{h=1}^H \frac{1}{\sigma_w^2} (\mathbf{x}_{h+1}^\tau - A\phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau)) \cdot \Delta \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau),$$

so we can bound

$$\|\nabla_A \log f_{A, \pi}(\tau)\|_{\text{op}} \leq \frac{B_\phi}{\sigma_w^2} \sum_{h=1}^H \|\mathbf{x}_{h+1}^\tau - A\phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau)\|_2.$$

Furthermore, we can also bound $(\mathbf{v}^\top \phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau))^2 \leq B_\phi^2$. We therefore have

$$\begin{aligned} \|\nabla_A \mathbf{v}^\top \mathbf{\Lambda}_{A,\pi} \mathbf{v}\|_{\text{op}} &\leq H B_\phi^3 \int \sum_{h=1}^H \|\mathbf{x}_{h+1}^\tau - A\phi(\mathbf{x}_h^\tau, \mathbf{u}_h^\tau)\|_2 \cdot f_{A,\pi}(\tau) d\tau \\ &\leq \frac{H^2 B_\phi^3}{\sigma_w^2} \cdot \text{poly}(d_x) \end{aligned}$$

where the last inequality follows from Lemma A.1. It follows from the Mean Value Theorem that

$$|\mathbf{v}^\top \mathbf{\Lambda}_{A,\pi} \mathbf{v} - \mathbf{v}^\top \mathbf{\Lambda}_{A',\pi} \mathbf{v}| \leq \frac{H^2 B_\phi^3}{\sigma_w^2} \cdot \text{poly}(d_x) \cdot \|A - A'\|_F.$$

As this holds for all $\mathbf{v} \in \mathcal{S}^{d-1}$, it follows that

$$\|\mathbf{\Lambda}_{A,\pi} - \mathbf{\Lambda}_{A',\pi}\|_{\text{op}} \leq \frac{H^2 B_\phi^3}{\sigma_w^2} \cdot \text{poly}(d_x) \cdot \|A - A'\|_F.$$

□

G Additional Experimental Details

In this section, we provide additional details on our experimental results presented in Section 6. All experiments were run on a machine with 56 Intel(R) Xeon(R) CPU E5-2690 v4 @ 2.60GHz CPUs, and 64GB RAM. All code was implemented in PyTorch.

G.1 Details on Problem Settings and Controller Parameterizations

We first expand on the precise definitions of the systems considered. As noted in Section 6, for the drone and car examples we set $H = 50$, and for the system of Section 1.1 we set $H = 10$. In addition, for all examples the noise is distributed as $\mathbf{w}_h \sim \mathcal{N}(0, 0.1 \cdot I)$. In all cases we set $\gamma^2 = 10H$ (where γ^2 is a bound on $\mathbb{E}_{\pi_{\text{exp}}}[\sum_{h=1}^H \mathbf{u}_h^\top \mathbf{u}_h]$), and we therefore let Π_{exp} denote the set of all policies satisfying $\mathbb{E}_{\pi_{\text{exp}}}[\sum_{h=1}^H \mathbf{u}_h^\top \mathbf{u}_h] \leq \gamma^2$.

G.1.1 System of Section 1.1 (Figure 1)

The dynamics for this system are given by

$$\mathbf{x}_{h+1} = 0.8\mathbf{x}_h + \mathbf{u}_h - \sum_{i=1}^{10} 3\phi_i(\mathbf{x}_h) + \mathbf{w}_h$$

for $\phi_i(\mathbf{x}) = \max\{1 - 100(\mathbf{x} - c_i)^2, 0\}$, and $\text{cost}(\mathbf{x}, \mathbf{u}) = (\mathbf{x} - c_1)^2 + 100^{-1} \cdot \mathbf{u}^2$. We set

$$c_1 = 10, c_2 = -14, c_3 = -11, c_4 = -8, c_5 = -5, c_6 = -2, c_7 = 1, c_8 = 4, c_9 = 7.$$

This then corresponds to a system in the form (1.1) with

$$A_\star = [0.8, 1, -3, \dots, -3], \quad \phi(\mathbf{x}, \mathbf{u}) = [\mathbf{x}, \mathbf{u}, \phi_1(\mathbf{x}), \dots, \phi_{10}(\mathbf{x})].$$

For this system, we parameterize our controller class $\Pi^\star$ as, for any $\pi^\theta \in \Pi^\star$ with parameter θ ,

$$\pi^\theta(x) = \theta_1 x + \sum_{i=1}^{10} \theta_{i+1} \phi_i(x) + \theta_{12}.$$

Note that the form of this controller lets us simply “match” the parameters of the system, and cancel undesirable parameters. Given this, for this system we let $\pi_\star(A)$ be the controller which sets $\theta_{1:11}$ to cancel the dynamics of the system A , and set $\theta_{12} = c_1 = 10$.

See Appendix G.1.3 for details on the computation of $\mathcal{H}(A)$ on this system.

G.1.2 Drone System (Figure 2)

The dynamics of this system are given by

$$\mathbf{x}_{h+1} = \begin{bmatrix} 1 & 0 & 0 & 0.1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0.1 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0.1 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix} \mathbf{x}_h + \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0.1 & 0 & 0 \\ 0 & 0.1 & 0 \\ 0 & 0 & 0.1 \end{bmatrix} \mathbf{u}_h + \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ -0.98 \end{bmatrix} + \mathbf{w}_h. \quad (\text{G.1})$$

Here we interpret $[\mathbf{x}]_{1:3}$ as the x, y , and z positions, respectively, and $[\mathbf{x}]_{4:6}$ as the x, y, z velocities. This system is therefore equivalent to three double integrator systems, with an affine term (which we interpret as “gravity”) affecting only the z coordinate. We set the cost to

$$\text{cost}(\mathbf{x}, \mathbf{u}) = \frac{0.1}{5} \cdot \sum_{i=1}^{d_x} [\mathbf{x}]_i^2 + \frac{1}{5} \cdot [\mathbf{u}]_1^2 + [\mathbf{u}]_2^2 + [\mathbf{u}]_3^2$$

This then corresponds to a system in the form (1.1) with

$$A_\star = \begin{bmatrix} 1 & 0 & 0 & 0.1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0.1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0.1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0.1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0.1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0.1 & -0.98 \end{bmatrix}, \quad \phi(\mathbf{x}, \mathbf{u}) = [\mathbf{x}, \mathbf{u}, 1].$$

For this system, we parameterize our controller class $\Pi^\star$ as, for any $\pi^\theta \in \Pi^\star$ with parameter θ ,

$$\pi_h^\theta(\mathbf{x}) = \theta_h^{\text{fb}} \mathbf{x} + \theta_h^{\text{offset}}$$

where $\theta_h^{\text{fb}} \in \mathbb{R}^{3 \times 6}$ is the state-feedback portion of the controller, and $\theta_h^{\text{offset}} \in \mathbb{R}^3$ is an offset term. It can be shown that the optimal controller for a system of the form (G.1) can be parameterized in this way Yu et al. (2020a). Furthermore, the optimal parameters can be computed in closed-form. As such, for this system we set $\pi_\star(A)$ to be with the optimal parameters, computed using this closed-form solution.

In addition to computing the optimal controller in closed-form, we can also compute the cost of a controller, $\mathcal{J}(\pi; A)$, in closed-form. To compute $\mathcal{H}(A)$ in this example, we then simply apply the `torch.autograd.functional.hessian` function to $\mathcal{J}(\pi_\star(A); A)$.

G.1.3 Car System (Figure 3)

The dynamics of this system are given by

$$\mathbf{x}_{h+1} = \begin{bmatrix} 1 & 0 & 0.1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0.1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0.1 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix} \mathbf{x}_h + \begin{bmatrix} 0 & 0 \\ 0 & 0 \\ 0.1 \cdot \cos([\mathbf{x}_h]_5) & 0 \\ 0.1 \cdot \sin([\mathbf{x}_h]_5) & 0 \\ 0 & 0 \\ 0 & 0.1 \end{bmatrix} \mathbf{u}_h + \mathbf{w}_h \quad (\text{G.2})$$

where $[\mathbf{x}_h]_5$ denotes the 5th element of $\mathbf{x}_h$. Here we interpret $[\mathbf{x}_h]_1$ as the x position, $[\mathbf{x}_h]_2$ as the y position, $[\mathbf{x}_h]_3$ as the x velocity, $[\mathbf{x}_h]_4$ as the y velocity, $[\mathbf{x}_h]_5$ as the angle of orientation (that is, the direction the car is facing), and $[\mathbf{x}_h]_6$ as the angular velocity. The first control dimension, then, corresponds to the “gas”, the power given to the car to move forward or backward, and the second control dimension corresponds to altering the direction of the steering wheel. Similar to the drone system, we set the cost to

$$\text{cost}(\mathbf{x}, \mathbf{u}) = \begin{bmatrix} \mathbf{x} \\ \mathbf{u} \end{bmatrix}^\top Q \begin{bmatrix} \mathbf{x} \\ \mathbf{u} \end{bmatrix} \quad \text{with} \quad Q = \frac{0.1 \cdot I + \mathbf{v}_1 \mathbf{v}_1^\top + \mathbf{v}_2 \mathbf{v}_2^\top}{\|0.1 \cdot I + \mathbf{v}_1 \mathbf{v}_1^\top + \mathbf{v}_2 \mathbf{v}_2^\top\|_{\text{op}}}$$

for some $\mathbf{v}_1, \mathbf{v}_2$. To write this in the form of (1.1), in order to make the problem more challenging we choose an overparameterized $\phi(\mathbf{x}, \mathbf{u})$:

$$\phi(\mathbf{x}, \mathbf{u}) = [\mathbf{x}, \mathbf{u}, \cos([\mathbf{x}]_5), \sin([\mathbf{x}]_5), [\mathbf{u}]_1 \cdot \cos([\mathbf{x}]_5), [\mathbf{u}]_1 \cdot \sin([\mathbf{x}]_5), [\mathbf{u}]_2 \cdot \cos([\mathbf{x}]_5), [\mathbf{u}]_2 \cdot \sin([\mathbf{x}]_5)]$$

and set

$$A_\star = \begin{bmatrix} 1 & 0 & 0.1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0.1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0.1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0.1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0.1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0.1 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix}.$$

For the car system, the controller class $\Pi^\star$ is a hierarchical controller parameterized by some $\boldsymbol{\theta} \in \mathbb{R}^4$. This controller first uses PD control to compute a “goal input”, the direction we would like to modify the state in, as:

$$\mathbf{u}_{\text{goal}}(\mathbf{x}) = -\boldsymbol{\theta}_1[\mathbf{x}]_{1:2} - \boldsymbol{\theta}_2[\mathbf{x}]_{3:4}.$$

Given the underactuated structure of the system in (G.2), we cannot directly push the state in the direction of $\mathbf{u}_{\text{goal}}(\mathbf{x})$. Instead, we set $\mathbf{u}$ to the following:

$$\begin{bmatrix} \mathbf{u}_{\text{goal}}(\mathbf{x})^\top \begin{bmatrix} \cos([\mathbf{x}]_5) \\ \sin([\mathbf{x}]_5) \end{bmatrix} \\ -\boldsymbol{\theta}_3([\mathbf{x}]_5 - \beta_{\text{goal}}(\mathbf{x})) - \boldsymbol{\theta}_4[\mathbf{x}]_6 \end{bmatrix} \quad \text{for} \quad \beta_{\text{goal}}(\mathbf{x}) = \tan^{-1}([\mathbf{u}_{\text{goal}}(\mathbf{x})]_2 / [\mathbf{u}_{\text{goal}}(\mathbf{x})]_1).$$

Given the complex form of this controller and the dynamics, there does not exist a closed-form way to set $\boldsymbol{\theta}$ optimally. Instead, for this system, we rely on a simple random search procedure to compute $\pi_\star(A)$. To find an optimal controller for system A , we randomly sample parameters $\boldsymbol{\theta}$, compute

the cost they incur on system A , and then set $\pi_\star(A)$ to the randomly generated controller with lowest cost. Note that this procedure is not differentiable, but we require $\pi_\star(A)$ is differentiable. To remedy this, in situations where a differentiable $\pi_\star(A)$ is needed (in particular, in the computation of $\mathcal{H}(A)$), rather than returning a single controller, we return the softmax distribution over all controllers sampled, weighting each controller by its estimated cost. As the softmax distribution can be differentiated, this parameterization of $\pi_\star(A)$ is differentiable.

For this system, there does not exist a closed-form expression for $\mathcal{J}(\pi; A)$ and, as such, to compute $\mathcal{J}(\pi; A)$, we simply perform many roll-outs of policy π on system A and average the cost. Given this and the search-based implementation of $\pi_\star(A)$ outlined above, we found that computing the hessian $\mathcal{H}(A)$ using the `torch.autograd.functional.hessian` as in Appendix G.1.2 was very memory-intensive. Instead, we computed the Jacobian $G(A) := \nabla_{A'} \mathcal{J}(\pi_\star(A'); A)|_{A'=A}$, and then, in place of $\mathcal{H}(A)$, we use $G(A)G(A)^\top$. To compute $G(A)$, we use the `torch.autograd.functional.jacobian` function. While using $G(A)G(A)^\top$ in place of $\mathcal{H}(A)$ is not justified by our theoretical analysis, if we are in settings where $\pi_\star(A)$ is not precisely the minimum of $\mathcal{J}(\pi; A)$ (which will likely be the case here since we are relying on a sampling-based implementation of $\pi_\star(A)$, which will incur some small error), then we argue that this is a reasonable metric to use. In particular, in this setting, the approximation of $\mathcal{J}(\pi_\star(\hat{A}); A_\star)$ given in Proposition 1 should have an additional first-order term of the form $G(A)^\top \text{vec}(A_\star - \hat{A})$. As we can upper bound

$$G(A)^\top \text{vec}(A_\star - \hat{A}) \leq \sqrt{\|\text{vec}(A_\star - \hat{A})\|_{G(A_\star)G(A_\star)^\top}^2},$$

optimizing for the metric $G(A)G(A)^\top$ instead of $\mathcal{H}(A)$ can be seen as minimizing the first-order Taylor-approximation of the excess loss. Intuitively, this metric quantifies the sensitivity of the loss to particular parameters in $A_\star$, and in practice we found that optimizing this metric produced significant improvements over existing methods. The implementation of the example from Section 1.1 relied on this same approximation.

G.2 Implementation Details

For all methods considered, our implementation follows the basic structure of Algorithm 1: at every epoch, we explore so as to minimize some exploration objective, form an estimate of $A_\star$ on the collected data, and then compute $\pi_\star(\hat{A}_t)$ on our estimate. Our main experimental results (Figures 1 to 3) show the loss of $\pi_\star(\hat{A}_t)$ as the time horizon t increases. For each method, to collect an initial set of data, we begin each trial by exploring randomly for some fixed number of episodes (10 for the drone example, 100 for the car example). The first point in each plot then corresponds to the performance after this initial random exploration. Each aspect of our implementation is modular, and any given component can be easily replaced. Below we highlight our implementation of the exploration routine, and choice of exploration objective, for the various approaches we consider.

G.2.1 Implementation of DynamicOED

Implementing the exploration procedure, DYNAMICOED, requires access to a regret minimization oracle. While in principle the LC³ algorithm of Kakade et al. (2020) could be applied to this problem to give such an oracle, the LC³ algorithm requires access to a computation oracle which is not clear how to implement in practice. To remedy this, we implement a Thompson Sampling-inspired modification to the LC³ algorithm of Kakade et al. (2020).

The primary computational challenge of implementing the LC³ algorithm is the computation of

the *optimistic* policy:

$$\arg \min_{\pi \in \Pi_{\text{exp}}} \min_{A \in \text{BALL}^t} \mathcal{J}^{\text{exp}}(\pi; A)$$

where $\mathcal{J}^{\text{exp}}(\pi; A)$ denotes the exploration cost that is minimized in LC³ (i.e. the expected cost on the cost function $\text{cost}_h^n(\mathbf{x}, \mathbf{u}) \leftarrow \frac{1}{M} \cdot \phi(\mathbf{x}, \mathbf{u})^\top (\Xi_n) \phi(\mathbf{x}, \mathbf{u})$ set in DYNAMICOED), and BALL^t the confidence set for $A_\star$ at iteration t .

To avoid solving this optimization, we adopt a Thompson Sampling-inspired variation of this procedure. In particular, at iteration t , we sample $\tilde{A}_t \sim \mathcal{N}(\hat{A}_t, \mathbf{\Lambda}_t^{-1})$. Standard Thompson Sampling would then compute $\arg \min_{\pi \in \Pi_{\text{exp}}} \mathcal{J}^{\text{exp}}(\pi; \tilde{A}_t)$, but even this can be challenging, so we instead rely on a sampling MPC-inspired approach. Given that we are at state $\mathbf{x}_h$ and have played inputs $\mathbf{u}_1, \dots, \mathbf{u}_{h-1}$, we aim to approximately solve the following optimization:

$$\begin{aligned} \min_{\mathbf{u}_h, \mathbf{u}_{h+1}, \dots, \mathbf{u}_H \in \mathbb{R}^{d_u}} \quad & \sum_{h'=h}^H \text{cost}_h^n(\mathbf{x}_{h'}, \mathbf{u}_{h'}) \\ \text{s.t.} \quad & \mathbf{x}_{h+1} = \tilde{A}_t \phi(\mathbf{x}_h, \mathbf{u}_h), \sum_{h=1}^H \mathbf{u}_h^\top \mathbf{u}_h \leq \gamma^2. \end{aligned} \tag{G.3}$$

To solve this approximately, we sample many possible $\mathbf{u}$ randomly, compute the value of the objective of (G.3) on the trajectories induced by these $\mathbf{u}$, and finally choose the input that minimizes this objective. Rather than playing the entire sequence of chosen inputs, however, we simply play the first input in the sequence, observe the new state on the actual system, and re-solve (G.3) on this new state. Note that the implementation of LC³ used for the experiments given in Kakade et al. (2020) relies on a similar Thompson Sampling-based approximation to the LC³ algorithm.

G.2.2 Implementation of Uniform Exploration

The goal of the procedure we have referred to as **Uniform Exploration** is to collect data which will result in the estimation error, $\|A_\star - \hat{A}\|_{\text{op}}$, being minimized, the goal of the method given in Mania et al. (2022). It can be shown that this is equivalent to maximizing $\lambda_{\min}(\mathbf{\Lambda}_T)$, so this method reduces to choosing inputs that maximize $\lambda_{\min}(\mathbf{\Lambda}_T)$. To implement this procedure, we rely on the same sampling-based MPC approach as we outlined above, with the primary difference being that instead of minimizing the objective of (G.3), we choose the inputs that maximize

$$\lambda_{\min} \left(\mathbf{\Lambda}_t + \sum_{h=1}^H \phi(\mathbf{x}_h, \mathbf{u}_h) \phi(\mathbf{x}_h, \mathbf{u}_h)^\top \right),$$

where $\mathbf{\Lambda}_t$ denote the covariates we have obtained so far at iteration t . While very similar in spirit to the algorithm of Mania et al. (2022), the implementation details are somewhat different than the algorithm proposed in that work. We found that in practice our implementation performed better than directly implementing (a sampling-based variant of) the algorithm from Mania et al. (2022), and all reported results for **Uniform Exploration** are therefore on this version.

G.2.3 Exploring via Cost Minimization

A natural point of comparison to our methods would be to forsake the system identification phase entirely, and simply run standard policy optimization algorithms such as TRPO or PPO (Schulman

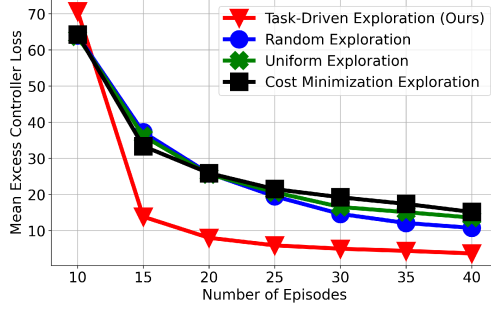


Figure 4: Performance on drone with LC^3 Exploration

et al., 2015, 2017), to obtain a controller $\hat{\pi}_T$. The primary difficulty with these approaches in the settings we consider is that these algorithms are *on-policy*, meaning that they primarily roll out trajectories using their current estimate of the optimal policy, $\hat{\pi}_t$, and using the collected data to do policy improvement on $\hat{\pi}_t$. In contrast, our setting is *off-policy*, in the sense that the learner must explore by playing policies in Π_{exp} , but return some policy in Π^* . Since in the settings we consider $\Pi_{\text{exp}} \neq \Pi^*$, on-policy approaches are simply exploring very differently, and therefore cannot be compared with directly.

This is particularly an issue in our setting where *stability* may come into play. Indeed, it may be the case that some controller in Π^* will destabilize the system, and cause the norm of the state to increase exponentially in h . Inducing such trajectories significantly improves one’s ability to perform system identification as the signal-to-noise ratio also then increases exponentially. However, to induce this trajectory with a state-feedback controller, the power of the input played by this controller will also increase exponentially in h . Since we choose Π_{exp} to include only policies with bounded power, this is not a fair comparison (and, furthermore, is likely not an algorithm one would want to run in practice).

While direct comparison with such approaches is therefore not possible, it is possible to compare against algorithms that, instead of collecting data that minimizes or maximizes objectives such as $\text{tr}(\mathcal{H}\tilde{\mathbf{A}}^{-1})$ or $\lambda_{\min}(\mathbf{A})$, instead simply aims to play policies minimizing $\mathcal{J}(\pi; A)$. In principle, such algorithms are similar to approaches such as TRPO in how they perform their exploration—both collect data by aiming to minimize the actual cost we are attempting to find a controller to minimize.

To implement this approach, we rely on a sampling-based MPC algorithm similar to that described in Appendix G.2.1, but where the goal is now to solve

$$\min_{\pi \in \Pi_{\text{exp}}} \mathcal{J}(\pi; A).$$

Note that the key difference between this approach and approaches such as TRPO is that we still only play $\pi \in \Pi_{\text{exp}}$. Using this objective to induce exploration, we then simply estimate A_* on this collected data, and return $\pi_*(\hat{A})$. The results of this approach on the drone system are given in Figure 4 (with this cost minimization approach denoted as **Cost Minimization Exploration**). As this illustrates, this approach is significantly worse than Algorithm 1, and is also outperformed by **Uniform Exploration** or **Random Exploration** when the number of episodes is large enough.

G.3 Additional Results

Finally, in this section we present versions of Figures 1 to 3 with error bars in Figures 5 to 7. In all figures, errors bars denote one standard error.

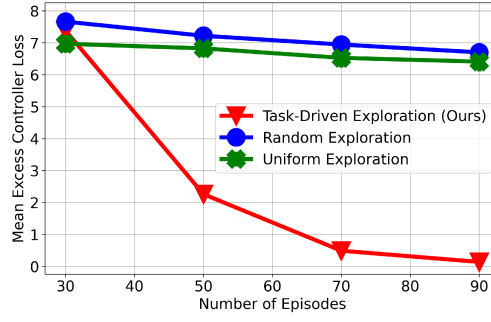


Figure 5: Mean excess controller loss on instance of Section 1.1 with error bars

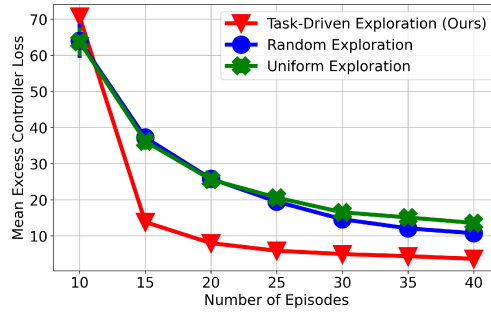


Figure 6: Mean excess controller loss on drone with error bars

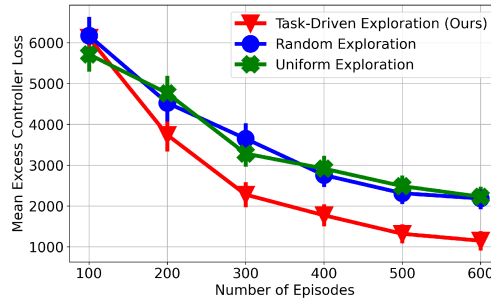


Figure 7: Mean excess controller loss on car with error bars