

Distributionally Robust Probabilistic Prediction for Stochastic Dynamical Systems

Tao Xu and Jianping He

Abstract—Probabilistic prediction of stochastic dynamical systems (SDSs) aims to accurately predict the conditional probability distributions of future states. However, accurate probabilistic predictions tightly hinge on accurate distributional information from a nominal model, which is hardly available in practice. To address this issue, we propose a novel functional-maximin-based distributionally robust probabilistic prediction (DRPP) framework. In this framework, one can design probabilistic predictors that have worst-case performance guarantees over a pre-defined ambiguity set of SDSs. Nevertheless, DRPP requires optimizing over the space of probability measures with density functions with respect to the Lebesgue measure, which is generally intractable. We develop a methodology that equivalently transforms the original maximin from function spaces to Euclidean spaces. Although it remains intractable to seek a global optimal solution, two suboptimal solutions are derived. By relaxing the constraints on the ambiguity set, we obtain a suboptimal predictor called Noise-DRPP. Relaxing the constraints on the predictor yields another suboptimal predictor, Eig-DRPP. Moreover, optimality gaps between the proposed predictors and the global optimal predictor are derived. Finally, we conduct elaborate numerical simulations to compare the performance of different predictors under different SDSs.

Index Terms—Stochastic Dynamical System, Distributionally Robust Optimization, Uncertainty Quantification, Probabilistic Prediction.

I. INTRODUCTION

A. Background

A probabilistic prediction is a forecast that assigns probabilities to different possible outcomes, enabling better uncertainty quantification, risk assessment, and decision-making. Due to the ever-increasing demand for both accurate and reliable predictions against uncertainty [1], it has wide applications in fields such as epidemiology [2], climatology [3], and robotics [4], etc. The performance of a probabilistic prediction is evaluated by both sharpness and calibration of the predictive probability density function (pdf). The sharpness reflects how concentrated (e.g., low-variance or low-entropy) the predictive pdf is, subject to the intrinsic stochasticity of the underlying system. The calibration concerns the statistical compatibility between the predictive pdf and the realized outcomes [5]. To simultaneously assess both sharpness and calibration, one usually assigns a numerical score to the predictive pdf and the realized value of the prediction target. If one has a further requirement that the expected scoring rule is maximized only when the predictive pdf equals the target's real pdf almost everywhere, a strictly proper scoring rule (see in [6]) is needed.

The probabilistic prediction of a stochastic dynamical system (SDS) is to predict the conditional pdfs of future states based on some prior information about the system. One's prior information is usually in the form of a nominal model, which is a simplified or idealized representation of the underlying system. While probabilistic prediction for a linear system with Gaussian noises is straightforward, it becomes computationally expensive once the nominal SDS is nonlinear or the noises are non-Gaussian. Therefore, a lot of research contributes to approximating the future pdf to balance the requirement of prediction precision and computation complexity [7].

B. Motivation

Despite the significant achievements in the current prediction algorithms for SDSs, they are not guaranteed to perform well if the distributional information provided by a nominal model is inaccurate. It is ubiquitous for a nominal SDS to possess both inaccurate state evolution functions and incomplete distributional information of noises. For example, the nominal state evolution function may be a local linearization from a nonlinear model [8]. Additionally, one's nominal information about the distributions of system noises is usually partial. Most of the time, only the estimated mean and covariance are available in the nominal model, which is far from uniquely determining a pdf [9]. A nominal predictor can lead to misleading predictive pdfs even though the uncertainties on SDS are not prominent (as demonstrated by experiments in Sec. VIII). Addressing this issue is crucial for the widespread application of probabilistic prediction in model-uncertain scenarios. In this paper, the aim is to develop a framework to design probabilistic predictors for SDSs with worst-case performance guarantees.

To design probabilistic predictors against distributional uncertainties of the model, we inherit the key notions from the distributionally robust optimization (DRO) [10]. DRO is a minimax-based mathematical framework for decision-making under uncertainty. It aims to find the decision that performs well across a range of possible pdfs in an ambiguity set. Nevertheless, what we need is to find the optimal probabilistic predictor that performs well over a set of SDSs. In DRPP, we generalize the idea of ambiguity sets to jointly describe one's uncertainties of noises' pdfs and state evolution functions.

The proposed DRPP framework has direct applicability in control and decision-making systems. Reliable predictive pdfs and confidence regions produced by DRPP can be integrated into downstream modules such as risk-aware model predictive control, robot trajectory planning, and safety assessment in power-system frequency regulation. For instance, understanding the evolution of robust confidence regions enables safe

This work was supported by the National Natural Science Foundation of China under Grant 62373247, 625B2117. (Corresponding author: Jianping He.) The authors are with the Department of Automation, Shanghai Jiao Tong University, and Key Laboratory of System Control and Information Processing, Ministry of Education of China, Shanghai 200240, China. E-mail: {Zerken, jphe}@sjtu.edu.cn.

motion planning for manipulators and informs constraint tightening in safety-critical control.

C. Contributions

Motivated by the aforementioned considerations, we propose a novel functional-maximin-based distributionally robust probabilistic prediction (DRPP) framework. In DRPP, the objective is the expected score under the worst-case SDS within an ambiguity set, and the predictor optimizes the worst-case objective over the space of admissible predictive pdfs. The main contributions of this paper are summarized as follows:

- To the best of our knowledge, this is the first work on probabilistic prediction of SDSs that optimizes the worst-case performance over an ambiguity set. We have generalized the classical conic moment-based ambiguity set such that both the uncertainties of state evolution functions and system noises are quantified.
- We develop a methodology to greatly reduce the complexity of solving DRPP. First, we use the principle of dynamic programming to derive the Bellman equation. Second, we exploit a necessary optimality condition to equivalently transform the maximin problem from function spaces to Euclidean spaces.
- We design two suboptimal DRPP predictors by relaxing the constraints. Noise-DRPP is the optimal predictor when the nominal state evolution function is accurate, and Eig-DRPP is the optimal predictor when the predictor is constrained by an eigenvector restriction. Moreover, an upper bound of the optimality gap of the proposed predictors is obtained.

D. Organization

The rest of this article is organized as follows. Section II provides a brief review of the related work and Section III introduces preliminaries on probabilistic prediction and DRO. In Section IV, we formulate the proposed DRPP problem. In Section V, we introduce the main methodology to address the main challenges in the DRPP. Then, we relax the constraints to solve two suboptimal DRPPs with upper and lower bounds in Section VI and Section VII. Next, we conduct elaborate experiments and provide an application in Section VIII, followed by a conclusion in Section IX. Please see Fig. 1 for the overall structure of this paper.

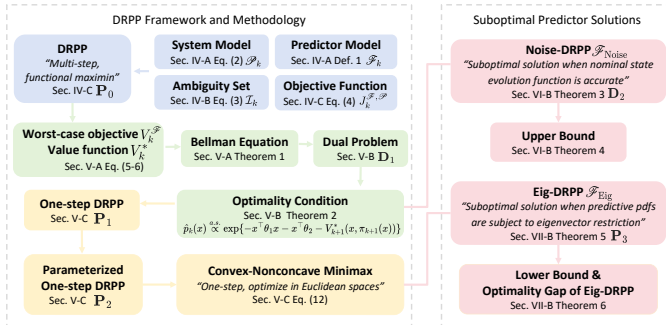


Fig. 1: An illustration of the roadmap.

II. RELATED WORK

The study of DRPP for SDSs integrates ideas from fields including probabilistic prediction, control theory, minimax optimization, and DRO. In this section, we briefly summarize the most relevant work as follows.

a) *Probabilistic prediction of SDS*: According to how the predictive pdfs are approximated, existing literature in the control community can be classified into the following groups. First, approximation by the first several central moments. Particularly, there are many methods using the mean and covariance to describe the predictive distribution. Classical local approximation methods include the Taylor polynomial in the extended Kalman filter [11], [12], the Stirling's interpolation [13] and the unscented transformation [14] as representatives for derivative-free estimation methods [15].

Second, approximation by a weighted sum of basis functions. The Gaussian mixture model (GMM) [16] approximates a pdf by a sum of weighted Gaussian distributions, and it can achieve any desired approximation accuracy if there are sufficiently large number of Gaussians [17]. Using GMM, the probabilistic prediction is equivalent to propagating the first two moments and weights of each basis [18]. The polynomial chaos expansion (PCE) [19] approximates a state evolution function by a sum of weighted polynomial basis functions, and these basis functions are orthogonal regarding the input pdf. In this way, the mean and covariance of the output can be conveniently approximated by the weights of PCE [20].

Third, approximation by a finite number of sampling points. The numerical solution using Monte Carlo (MC) samplings [21] can approximate any statistics of the output distribution if the sampling number is sufficiently high. For a further review of the modern MC methods in approximated probabilistic prediction, we recommend the review [22].

Nevertheless, assuming the accuracy of a nominal model is restrictive, especially in the control engineering practice. As reviewed in [7], an important application of probabilistic prediction for SDSs is to formulate chance constraints [23] in stochastic model predictive control (SMPC) [24], [25]. In these situations, a misleading predictive pdf can lead to unsafe control policies. To take the distributional uncertainties into account, Coulson et al [26] have proposed a novel distributionally robust data-enabled predictive control algorithm, Coppens and Patrinos [27] develop a data-driven distributionally robust MPC scheme using generalized moment-based ambiguity sets. However, these works aim to robustify the control performance, while the problem of distributionally robust probabilistic prediction remains open.

b) *Probabilistic prediction in machine learning*: Taking a historical view of the development of PP in the statistics and machine learning communities, we can find that the research in the early phase also focuses on approximating the posterior pdf, especially from a Bayesian perspective. The advent of the scoring rule has provided a convenient scalar way to measure the performance of probabilistic predictions [5]. Then, the research focus is shifting toward training probabilistic predictors to optimize the empirical scores. Classical statistical methods include Bayesian statistical models [28], approximate

Bayesian forecasting [29], quantile regression [30], etc. Recent machine learning algorithms include distributional regression [31], boosting methods such as NGBoost [32], and autoregressive recurrent networks like DeepAR [33], etc. We recommend [34] for a more detailed review of machine learning algorithms for probabilistic prediction.

c) *Distributionally robust optimization*: To guarantee the prediction performance for SDSs when the nominal model is not accurate, the primary task is to describe the ambiguity set for SDSs, which is the key concept in the DRO community. In the seminal work [35], the Chebyshev ambiguity set is defined, where only the mean and covariance are known. Using this ambiguity set, several extensions are further studied in [36]. Next, Delage and Ye [10] generalize the Chebyshev ambiguity set to the conic moment-based ambiguity set, which allows the mean and covariance matrix to be also unknown. Shapiro and Pichler [37] further extend the conic moment-based ambiguity set to the conditional version, which is suitable for multistage decision-making problems. Overall, viewing from the outer optimization, DRO still optimizes in an Euclidean space, while DRPP optimizes in a functional space of all pdfs. Since classical tools in DRO cannot be directly applied to solve DRPP, new methods are needed.

III. PRELIMINARIES

In this paper, we denote random variables in **bold fonts** to distinguish them from constant variables. Given a random variable $\mathbf{x}$ taking values in $\mathcal{X}$, we denote its probability measure as $P_{\mathbf{x}}(\cdot)$ and its pdf as $p_{\mathbf{x}}(\cdot)$. Consider another random variable $\mathbf{y}$, we denote the conditional probability measure of $\mathbf{y}$ given $\mathbf{x}$ as $P_{\mathbf{y}|\mathbf{x}}(\cdot | \cdot)$ and the conditional pdf as $p_{\mathbf{y}|\mathbf{x}}(\cdot | \cdot)$. Let $\mathcal{P}(\mathcal{X})$ be the space of probability measures on $\mathcal{X}$, equipped with the standard weak*-topology. Let $\mathcal{P}_2(\mathcal{X}) \subset \mathcal{P}(\mathcal{X})$ be the space of probability measures with finite second moments. Let $\mathcal{M}_+(\mathcal{X})$ be the set of positive measures on $\mathcal{X}$. Given a measurable map $h : \mathcal{X} \rightarrow \mathbb{R}$, the expectation of $h(\mathbf{x})$ is denoted as $\mathbb{E}_{P_{\mathbf{x}}} h(\mathbf{x})$ or $\mathbb{E}_{\mathbf{x} \sim P_{\mathbf{x}}} h(\mathbf{x})$. We also denote a sequence $\{(\cdot)_k\}_{k=1}^T$ by $(\cdot)_{1:T}$. We denote the set of symmetric matrices with dimension n as S^n , whose subset of positive semi-definite matrices is S_+^n . For a function $f : \mathcal{X} \rightarrow \mathbb{R}$, the operator norm $\|f\|_2 := \inf\{c \geq 0 : \|f(v)\|_2 \leq c\|v\|_2 \ \forall v \in \mathbb{R}^n\}$. For $\rho \in \mathcal{P}(\mathcal{X})$, we denote $\langle \rho, f \rangle := \int_{\mathcal{X}} f(x) d\rho(x)$. Given $A, B \in \mathbb{R}^{n \times n}$, we use $A \succeq B$ to indicate $A - B \in S_+^n$, and we denote $\langle A, B \rangle := \text{tr}(A^T B)$.

A. Probabilistic Prediction

Given a random variable $\mathbf{x}$ taking value on $\mathcal{X}$, consider a predictive space, $\mathcal{F}$, of probability measures on $\mathcal{X}$ that admits a Lebesgue density, $\hat{p}_{\mathbf{x}}$, for each element. After the value of $\mathbf{x}$ is materialized as x , a scoring rule,

$$\mathcal{S}(\hat{p}_{\mathbf{x}}, x) : \mathcal{F} \times \mathcal{X} \rightarrow \mathbb{R},$$

assigns a numerical score $\mathcal{S}(\hat{p}_{\mathbf{x}}, x)$ to measure the quality of the predictive distribution $\hat{p}_{\mathbf{x}}$ on the realized value x . Notice that the scoring rule $\mathcal{S}(\hat{p}_{\mathbf{x}}, x)$ is a random variable because it depends on the realization of $\mathbf{x}$. A scoring rule $\mathcal{S}$ is proper with respect to the predictive space $\mathcal{F}$ if $\mathbb{E}\mathcal{S}(\hat{p}_{\mathbf{x}}, x) \leq \mathbb{E}\mathcal{S}(p_{\mathbf{x}}, x)$

holds for all $\hat{p}_{\mathbf{x}}, p_{\mathbf{x}} \in \mathcal{F}$. It is strictly proper if the equality holds only when $\hat{p}_{\mathbf{x}}$ equals $p_{\mathbf{x}}$ almost everywhere. For example, the logarithm score,

$$\mathcal{L}(\hat{p}_{\mathbf{x}}, x) := \log \hat{p}_{\mathbf{x}}(x),$$

is one of the most celebrated scoring rules for being essentially the only local proper scoring rule up to equivalence [38].

B. Distributionally Robust Optimization

Traditional stochastic optimization approaches assume a known probability distribution P and solve problems

$$\inf_{x \in \mathcal{X}} \mathbb{E}_{\xi \sim P} [h(x, \xi)],$$

where $x \in \mathcal{X}$ represents the decision variable, ξ is sampled from the distribution P , $h(x, \xi)$ is the objective function. However, in many real-world problems, the true probability distribution of the uncertain parameters is unknown or partially known. DRO addresses this issue by considering a family of distributions, known as an ambiguity set, rather than a single distribution. The goal is to find a solution that performs well for the worst-case distribution within this set. A general DRO problem is formulated as:

$$\inf_{x \in \mathcal{X}} \sup_{Q \in \mathcal{D}} \mathbb{E}_{\xi \sim Q} [h(x, \xi)],$$

where $\mathcal{D}$ is the ambiguity set, a set of probability distributions that are close to the nominal distribution P in some sense.

The ambiguity set $\mathcal{D}$ can be defined in various ways. One of the most popular ones is the moment-based sets, where distributions are constrained by their moments (e.g., mean, variance). For example, the celebrated conic moment-based ambiguity set is used in the seminal work [10] such that

$$\mathcal{D}_1(\mathcal{X}, \mu_0, \Sigma_0, \gamma_1, \gamma_2) := \left\{ p_{\xi} \left| \begin{array}{l} \mathbb{P}(\xi \in \mathcal{X}) = 1 \\ (\mathbb{E}[\xi] - \mu_0)^T \Sigma_0^{-1} (\mathbb{E}[\xi] - \mu_0) \leq \gamma_1 \\ \mathbb{E}[(\xi - \mu_0)(\xi - \mu_0)^T] \preceq \gamma_2 \Sigma_0 \end{array} \right. \right\}, \quad (1)$$

where $\mathcal{X}$ is the nominal support, μ_0 and Σ_0 are the nominal mean and covariance, $\gamma_1, \gamma_2 \in \mathbb{R}_+$ quantifies the uncertainties on the mean and covariance. A thorough review of popular ambiguity sets is provided in [39, Sec. 3].

IV. PROBLEM FORMULATION

To begin with, we define the system and probabilistic predictor model, respectively. Next, we generalize the classical conic moment-based ambiguity set (1) for SDSs such that one's uncertainties of both state evolution functions and system noises are jointly described. Finally, we formulate the DRPP problem as a multistep functional maximin optimization.

A. System and Predictor Model

System model $(\mathcal{X}, \mathcal{U}, \mathcal{P}, \pi, T)$. Consider a discrete-time SDS with horizon $T \in \mathbb{Z}_+$, denoted as $\mathcal{P}$, whose dynamics (or kernel) at time step $k \in \{0, \dots, T-1\}$ is

$$\mathcal{P}_k : \mathbf{x}_{k+1} = f_k(\mathbf{x}_k, \mathbf{u}_k) + \mathbf{w}_k, \quad (2)$$

where $\mathcal{P}_k$ denotes the system kernel at time step k , $\mathbf{x}_k$ is the system state taking values in the state space $\mathcal{X} := \mathbb{R}^{d_x}$, the initial state is given as $\mathbf{x}_0$. $\mathbf{u}_k$ is the control taking values in the input space $\mathcal{U} := \mathbb{R}^{d_u}$, which comes from a deterministic state-feedback policy $\pi_k : \mathcal{X} \rightarrow \mathcal{U}$, i.e., $\mathbf{u}_k = \pi_k(\mathbf{x}_k)$. The subsequent problem formulation thus evaluates predictive performance conditioned on a fixed policy π . We denote the state-control pair $(\mathbf{x}_k, \mathbf{u}_k)$ as $\mathbf{z}_k \in \mathcal{Z} := \mathcal{X} \times \mathcal{U}$, then $f_k : \mathcal{Z} \rightarrow \mathcal{X}$ is the state evolution function. $\mathbf{w}_k$ is an exogenous noise vector taking values in $\mathbb{R}^{d_x}$.

Assumption 1 (SDS). At each time step $k \in \{0, \dots, T-1\}$,

- f_k has finite operator norms, i.e., $\|f_k\|_2$ is bounded, which is satisfied by many nonlinear systems in practice.
- the noise $\mathbf{w}_k$ has a bounded second moment, i.e., $P_{\mathbf{w}_k} \in \mathcal{P}_2(\mathbb{R}^{d_x})$, and $\mathbf{w}_{0:T-1}$ are independent but not necessarily identically distributed.

Predictor model ($\mathcal{F}$). Starting from the system's initial state $\mathbf{x}_0$, a probabilistic predictor keeps observing the states and the control inputs, aiming to predict the conditional pdfs of the next state.

Definition 1 (Predictive space). Let $\mathcal{F} \subset \mathcal{P}(\mathcal{X})$ be the set of continuous positive pdfs on $\mathcal{X}$ with tail envelope of order 2, i.e.,

$$\mathcal{F} := \left\{ \hat{p} \mid \begin{array}{l} \hat{p} \text{ is a strictly positive pdf on } \mathcal{X}, \\ \exists C > 0 \text{ s.t. } |\log \hat{p}(x)| \leq C(1 + \|x\|_2^2) \forall x \in \mathcal{X} \end{array} \right\}.$$

A probabilistic predictor policy for $\mathcal{P}$ is a sequence $\mathcal{F} = (\mathcal{F}_0, \dots, \mathcal{F}_{T-1})$ where each $\mathcal{F}_k : \mathcal{Z} \rightarrow \mathcal{F}$ is a measurable map from the previous state and control input to a predictive conditional pdf for the next state. Let $\mathfrak{F}$ be the collection of such predictor policies.

Assumption 2 (Predictor's prior information). At each time step $k \in \{0, \dots, T-1\}$, the probabilistic predictor has only access to the following information:

- the state trajectory $\mathbf{x}_{0:k}$, control input sequence $\mathbf{u}_{0:k}$,
- a nominal state evolution function $\bar{f}_k : \mathcal{Z} \rightarrow \mathcal{X}$,
- a nominal noise mean $\bar{\mu}_k \in \mathcal{U}$, and a nominal noise covariance $\bar{\Sigma}_k \in S_+^{d_x}$.

The predictor's prior information usually does not align with the ground truth. Therefore, to further quantify the uncertainty between the nominal model and the real model, we need to define an ambiguity set of SDSs.

B. Ambiguity Set

The crucial spirit of an ambiguity set is that although a nominal model rarely identifies with the ground truth, one can still have additional information or belief that the real model is located around the nominal model. For example, an upper bound for the operator norm $\|f_k - \bar{f}_k\|_2$ may be available when the nominal function is derived from a system identification procedure. Additionally, estimations and confidence regions (CRs) of the mean and covariance of system noise may be statistically inferred from historical data. Even if no supportive evidence exists to construct a reasonable ambiguity set, one can still subjectively select an ambiguity set.

Jointly considering the uncertainties of both state evolution functions and noises, we define the conditional conic moment-based ambiguity sets for $\mathcal{P}$.

Definition 2 (Conditional conic moment-based ambiguity set). For each time step $k \in \{0, \dots, T-1\}$ and state-control pair $\mathbf{z}$, the conditional conic moment-based ambiguity set is a subset of $\mathcal{P}_2(\mathcal{X})$ defined as

$$\mathcal{I}_k(\mathbf{z} \mid \bar{f}_k, \bar{\mu}_k, \bar{\Sigma}_k, \gamma_0, \gamma_1, \gamma_2, \gamma_3) =: \left\{ P_{\mathbf{x}_{k+1}|\mathbf{z}_k}(\cdot \mid \mathbf{z}) \mid \begin{array}{l} \mathbf{x}_{k+1} = f_k(\mathbf{z}) + \mathbf{w}_k, \\ \|f_k(\mathbf{z}) - \bar{f}_k(\mathbf{z})\|_2^2 \leq \gamma_0(\mathbf{z}) \\ \|\mathbb{E}(\mathbf{w}_k) - \bar{\mu}_k\|_{\bar{\Sigma}_k^{-1}}^2 \leq \gamma_1 \\ \gamma_3 \bar{\Sigma}_k \preceq \mathbb{E}(\mathbf{w}_k - \bar{\mu}_k)(\mathbf{w}_k - \bar{\mu}_k)^\top \preceq \gamma_2 \bar{\Sigma}_k \end{array} \right\}, \quad (3)$$

where $\bar{f}_k, \bar{\mu}_k, \bar{\Sigma}_k$ are defined in Assumption 2, and predictor's uncertainties are quantified by $\gamma_0 : \mathcal{Z} \rightarrow \mathbb{R}_+$ and $\gamma_j \in \mathbb{R}_+$ for $j = 1, 2, 3$. When there is no confusion about the parameters that define an ambiguity set, we use $\mathcal{I}_k(\mathbf{z})$ for simplification.

The conic moment-based ambiguity set covers a wide family of distributions with bounded second moments. This choice not only facilitates tractable dual reformulations but also provides a statistically interpretable description of moment uncertainty, leading to guaranteed robustness without assuming a specific parametric form.

Assumption 3 (Regularity). At each time step $k \in \{0, \dots, T-1\}$, conditioned on the state-control pair $\mathbf{z}_k = \mathbf{z}_k$, the system chooses a conditional probability measure $P_k(\cdot \mid \mathbf{z}_k) \in \mathcal{I}_k(\mathbf{z}_k)$ independent of the previous choices.

Let $\mathfrak{P}$ be a collection of SDSs subject to the conditional conic moment-based ambiguity set and regularity assumption.

C. Problem in Interest

Objective function ($J_k^{\mathcal{F}, \mathcal{P}}$). Starting from the time step $k \in \{0, \dots, T-1\}$, given an initial state-control pair $\mathbf{z} \in \mathcal{Z}$, a predictor policy $\mathcal{F} \in \mathfrak{F}$, and a system kernel $\mathcal{P} \in \mathfrak{P}$, the prediction performance of $\mathcal{F}$ is characterized by the expectation of cumulative log-score over the state trajectory. We define it by the objective function as follows:

$$J_k^{\mathcal{F}, \mathcal{P}}(\mathbf{z}) := \mathbb{E}_{\mathcal{P}_{k:T-1}} \left[\sum_{t=k}^{T-1} \mathcal{L}(\mathcal{F}_t(\mathbf{z}_t), \mathbf{x}_{t+1}) \mid \mathbf{z}_k = \mathbf{z} \right], \quad (4)$$

where the notation $\mathbb{E}_{\mathcal{P}}$ indicates that the expectation is taken over the random trajectory $(\mathbf{x}_k, \dots, \mathbf{x}_T)$ generated by the kernel sequence $\mathcal{P}_{k:T-1}$.

Distributionally robust probabilistic prediction ($\mathbf{P}_0$). Given a set of probabilistic predictors $\mathfrak{F}$ and a set of system kernels $\mathfrak{P}$ for the SDS $\mathcal{P}$ with an initial state-control pair $\mathbf{z}_0$, we are interested in finding probabilistic predictors in $\mathfrak{F}$ that are distributionally robust with respect to $\mathfrak{P}$. In other words, a distributionally robust probabilistic predictor should have a worst-case performance guarantee for any system kernel in $\mathfrak{P}$. To this end, a multistep functional maximin optimization is formulated as follows:

$$(\mathbf{P}_0) : \sup_{\mathcal{F} \in \mathfrak{F}} \inf_{\mathcal{P} \in \mathfrak{P}} J_0^{\mathcal{F}, \mathcal{P}}(\mathbf{z}_0).$$

V. MAIN METHODOLOGY

To begin with, we derive the Bellman equation of DRPP based on the principle of dynamic programming, which shows that solving the DRPP problem $\mathbf{P}_0$ is equivalent to solving a sequence of Bellman equations in a backward manner. Then, we canonicalize the ambiguity set of SDS and transform the inner optimization into an infinite-dimensional conic linear program. Next, we take the dual form of inner minimization with a strong duality guarantee. A necessary optimality condition is developed to handle the computational intractability of functional optimization. Exploiting this optimality condition, we equivalently transform the one-step DRPP in function spaces into a convex-nonconcave minimax problem in Euclidean spaces. Finally, a discussion on the computational complexity is provided.

A. Bellman Equation

The concept of robust value function originates from the study of robust Markov decision processes (RMDPs), where the optimal worst-case performance of a control problem is quantified. Here, we apply the same idea to $\mathbf{P}_0$. Starting from the time step $k \in \{0, \dots, T-1\}$, given an initial state-control pair $z \in \mathcal{Z}$ and a predictor policy $\mathcal{F} \in \mathfrak{F}$, the worst-case objective function for $\mathcal{F}$ is defined as

$$V_k^{\mathcal{F}}(z) := \inf_{\mathcal{P} \in \mathfrak{P}} J_k^{\mathcal{F}, \mathcal{P}}(z). \quad (5)$$

Then, the robust value function is defined as the optimal worst-case objective function, i.e.,

$$V_k^*(z) := \sup_{\mathcal{F} \in \mathfrak{F}} V_k^{\mathcal{F}}(z). \quad (6)$$

Leveraging the principle of dynamic programming, one can derive the following recursive equations for V_k^* .

Theorem 1 (Bellman equation). *The set of robust value functions $\{V_k^*, k = 0, \dots, T\}$ satisfies the following Bellman equation: for any state-control pair $z \in \mathcal{Z}$, $V_T^*(z) = 0$, and for $k \in \{0, \dots, T-1\}$, there is*

$$V_k^*(z) = \sup_{\hat{p}_k \in \mathcal{F}} \inf_{\rho_k \in \mathcal{I}_k(z)} \int_{\mathcal{X}} \mathcal{L}(\hat{p}_k, x) + V_{k+1}^*(x, \pi_{k+1}(x)) d\rho_k(x). \quad (7)$$

Proof. Please see Appendix A. $\square$

Remark 1. *The Bellman equation implies that different control policies can lead to different robust value functions even when the system kernel family $\mathfrak{P}$ and predictor family $\mathfrak{F}$ are fixed. Thus, the DRPP $\mathbf{P}_0$ can be extended to jointly optimize the control policy and the predictor policy in future work.*

Theorem 1 shows that solving the DRPP problem $\mathbf{P}_0$ is equivalent to solving the Bellman equation (7). This result generalizes the classical Bellman equation for RMDPs [40] to the DRPP setting.

Let $\nu_k = f_k(z)$ and $\bar{\nu}_k = \bar{f}_k(z)$, the ambiguity constraint of f_k is equivalent to describing the distance between ν_k and $\bar{\nu}_k$, i.e., $\|\nu_k - \bar{\nu}_k\|_2^2 = \|f_k(z) - \bar{f}_k(z)\|_2^2 \leq \gamma_0(z)$. To deal with the

remaining two ambiguity constraints on the first two moments of w_k , let $\mu_k = \mathbb{E}w_k$, and the r.h.s. of (7) is transformed as

$$\begin{aligned} & \sup_{\hat{p}_k \in \mathcal{F}} \inf_{\nu_k, \mu_k, \rho_k} \int_{\mathcal{X}} \mathcal{L}(\hat{p}_k, x) + V_{k+1}^*(x, \pi_{k+1}(x)) d\rho_k(x) \\ & \text{s.t.} \begin{cases} \rho_k \in \mathcal{M}_+(\mathcal{X}), \mathbb{E}_{\rho_k}[1] = 1, \mathbb{E}_{x \sim \rho_k}[x] = \mu_k + \nu_k \\ \gamma_3 \bar{\Sigma}_k \preceq \mathbb{E}_{x \sim \rho_k} \left[(x - \nu_k - \bar{\mu}_k)(x - \nu_k - \bar{\mu}_k)^\top \right] \preceq \gamma_2 \bar{\Sigma}_k \\ \begin{bmatrix} \bar{\Sigma}_k & (\mu_k - \bar{\mu}_k) \\ (\mu_k - \bar{\mu}_k)^\top & \gamma_1 \end{bmatrix} \succeq 0 \\ \begin{bmatrix} I & (\nu_k - \bar{\nu}_k) \\ (\nu_k - \bar{\nu}_k)^\top & \gamma_0(z) \end{bmatrix} \succeq 0, \end{cases} \end{aligned} \quad (8)$$

where the inner problem is canonicalized into an infinite-dimensional conic linear program.

B. Optimality Condition

Directly solving the worst-case probability measure in the positive measure space $\mathcal{M}_+(\mathcal{X})$ is computationally intractable. As a first step towards dealing with this challenge, we leverage the dual analysis for (8). If a strong duality holds, one can optimize the Lagrange multipliers in the dual form. For the inner optimization of (8), a classical conclusion is that $\bar{\Sigma}_k \succ 0$ is a sufficient condition for strong duality to hold [10, p. 5]. Let $r \in \mathbb{R}, q \in \mathbb{R}^{d_x}, Q_i \in S_+^{d_x}, \kappa_i = \{P_i \in S_+^{d_x}, p_i \in \mathbb{R}^{d_x}, s_i \in \mathbb{R}\}$ for $i = 1, 2$ be the Lagrange multipliers, the dual form of (8) is

$$\begin{aligned} (\mathbf{D}_1) : & \sup_{\hat{p}_k \in \mathcal{F}, r, q, Q_1, Q_2, \kappa_1, \kappa_2} G(r, q, Q_1, Q_2, \kappa_1, \kappa_2) \\ & \text{s.t.} \begin{cases} x^\top (Q_1 - Q_2)x + x^\top q + r + \log \hat{p}_k(x) \\ + V_{k+1}^*(x, \pi_{k+1}(x)) \geq 0 \quad \forall x \in \mathcal{X} \\ \begin{bmatrix} P_1 & p_1 \\ p_1^\top & s_1 \end{bmatrix} \succeq 0, \begin{bmatrix} P_2 & p_2 \\ p_2^\top & s_2 \end{bmatrix} \succeq 0. \end{cases} \end{aligned}$$

where $G(r, q, Q_1, Q_2, \kappa_1, \kappa_2) = -r + \bar{\mu}_k^\top (Q_1 - Q_2) \bar{\mu}_k - (\gamma_2 Q_1 - \gamma_3 Q_2 + P_1) \cdot \bar{\Sigma}_k - P_2 \cdot I + 2p_1^\top \bar{\mu}_k - s_1 \gamma_1 + 2p_2^\top \bar{\nu}_k - s_2 \gamma_0(z) + g(q, p_1, p_2, Q_1, Q_2)$, and $g(q, p_1, p_2, Q_1, Q_2) = \inf_{\nu_k} -(q + 2p_2)^\top \nu_k - \nu_k^\top (Q_1 - Q_2) \nu_k$ s.t. $q + 2p_1 + 2(Q_1 - Q_2)(\nu_k + \bar{\mu}_k) = 0$. Please see Appendix B for a detailed derivation.

Unfortunately, we have to deal with an infinite number of constraints indexed by $x \in \mathcal{X}$ in problem $\mathbf{D}_1$. Although the dual problem cannot lead to a direct solution, it helps to develop an optimality condition for the probabilistic predictors in the following theorem.

Theorem 2 (Optimality condition). *A necessary optimality condition for the DRPP $\mathbf{P}_0$ is that $\hat{p}_k$ belongs to the following exponential family almost everywhere on $\mathbb{R}^{d_x}$,*

$$\hat{p}_k(x) \stackrel{a.s.}{\propto} \exp\{-x^\top \theta_1 x - x^\top \theta_2 - V_{k+1}^*(x, \pi_{k+1}(x))\}, \quad (9)$$

where $\propto$ means the right-hand side gives the density up to a normalizing constant, $\theta_1 \in S^{d_x}, \theta_2 \in \mathbb{R}^{d_x}$, and $k \in \{0, \dots, T-1\}$. Particularly for $k = T-1$, the optimal $\hat{p}_{T-1}$ is subject to a Gaussian distribution almost everywhere, i.e., $\exists \hat{\mu}_k \in \mathbb{R}^{d_x}$ and $\hat{\Sigma}_k \in S_+^{d_x}$ such that

$$\hat{p}_k \stackrel{a.s.}{\sim} \mathcal{N}(\hat{\mu}_k, \hat{\Sigma}_k). \quad (10)$$

Proof. Please see Appendix C. $\square$

Theorem 2 has reduced the optimization complexity from the original function space $\mathcal{F}$ to a specific exponential distribution family. For $k = T - 1$, the space is further reduced to $S_+^{d_x} \times \mathbb{R}^{d_x}$, where $\mathbf{D}_1$ can be transformed into a finite-dimensional optimization problem as follows

$$\begin{aligned} & \max_{\hat{\mu}_k, \hat{\Sigma}_k, Q_1, Q_2, \kappa_1, \kappa_2} -\frac{1}{2}d_x \log(2\pi) + \frac{1}{2} \log \det(\hat{\Sigma}_k^{-1}) \\ & - (\gamma_2 Q_1 - \gamma_3 Q_2) \cdot \bar{\Sigma}_k - P_1 \cdot \bar{\Sigma}_k - P_2 \cdot I - s_1 \gamma_1 \\ & - s_2 \gamma_0(z^2) + 2p_1^\top \hat{\Sigma}_k(2p_2 - p_1) - 2p_2^\top (\hat{\mu}_k - \bar{\mu}_k - \bar{\nu}_k) \quad (11) \\ \text{s.t. } & \begin{cases} Q_1 \succeq 0, Q_2 \succeq 0, Q_1 - Q_2 = \frac{1}{2} \hat{\Sigma}_k^{-1}, \hat{\Sigma}_k \succ 0 \\ \begin{bmatrix} P_1^\top & p_1 \\ p_1^\top & s_1 \end{bmatrix} \succeq 0, \begin{bmatrix} P_2^\top & p_2 \\ p_2^\top & s_2 \end{bmatrix} \succeq 0. \end{cases} \end{aligned}$$

Please see Appendix D for a detailed derivation. However, this problem is now a nonconvex semidefinite optimization with respect to $\hat{\Sigma}_k$, solving a globally optimal solution is NP-hard in general. Even worse, the lack of a closed-form solution of V_{T-1}^* makes it impossible to recursively parameterize and solve $\hat{p}_k$ backwardly based on the Bellman equation (7).

C. One-step DRPP

At time step $k \in \{0, \dots, T-1\}$, the Bellman equation (7) implies that the optimal prediction policy maximizes the one-step score $\mathcal{L}(\hat{p}_k, x)$ plus the expected best future cumulative scores $V_{k+1}^*(x, \pi_{k+1}(x))$ from the next state x onward. If one does not take into account the future scores, Theorem 2 shows that the optimal predictive pdf can be parametrized by a Gaussian distribution. Given $z \in \mathcal{Z}$, consider the following **one-step DRPP** problem:

$$(\mathbf{P}_1) : \sup_{\hat{p}_k \in \mathcal{F}} \inf_{\rho_k \in \mathcal{I}_k(z)} \int_{\mathcal{X}} \mathcal{L}(\hat{p}_k, x) d\rho_k(x),$$

we substitute $\hat{p}_k$ by (10) and the objective follows as

$$\mathbb{E}_{x \sim \rho_k} -\frac{1}{2} [d_x \log(2\pi) + \log \det(\hat{\Sigma}_k) + (x - \hat{\mu}_k) \hat{\Sigma}_k^{-1} (x - \hat{\mu}_k)^\top].$$

Next, notice that $\mathbb{E}_{x \sim \rho_k}[x] = \mu_k + \nu_k$ and $\mathbb{E}_{x \sim \rho_k}(x - \nu_k - \mu_k)(x - \nu_k - \mu_k)^\top = \Sigma_k$, one has $\mathbb{E}_{x \sim \rho_k}(x - \hat{\mu}_k) \hat{\Sigma}_k^{-1} (x - \hat{\mu}_k)^\top$ equals to $[\Sigma_k + (\mu_k + \nu_k - \hat{\mu}_k)(\mu_k + \nu_k - \hat{\mu}_k)^\top] \cdot \hat{\Sigma}_k^{-1}$. The functional minimax $\mathbf{P}_1$ is then equivalent to a minimax optimization in Euclidean spaces as follows

$$\begin{aligned} (\mathbf{P}_2) : & \min_{\hat{\Sigma}_k, \hat{\mu}_k} \max_{\Sigma_k, \mu_k, \nu_k} -\log \det(\hat{\Sigma}_k^{-1}) \\ & + [\Sigma_k + (\mu_k + \nu_k - \hat{\mu}_k)(\mu_k + \nu_k - \hat{\mu}_k)^\top] \cdot \hat{\Sigma}_k^{-1} \\ \text{s.t. } & \begin{cases} \|\nu_k - \bar{\nu}_k\|_2^2 \leq \gamma_0(z), \|\mu_k - \bar{\mu}_k\|_{\hat{\Sigma}_k^{-1}}^2 \leq \gamma_1 \\ \gamma_3 \bar{\Sigma}_k \preceq \Sigma_k + (\mu_k - \bar{\mu}_k)(\mu_k - \bar{\mu}_k)^\top \preceq \gamma_2 \bar{\Sigma}_k. \end{cases} \end{aligned}$$

Lemma 1. Given any μ_k , the one-step DRPP $\mathbf{P}_2$ has explicit solutions for Σ_k and $\hat{\mu}_k$ as

- i) $\Sigma_k^* = \gamma_2 \bar{\Sigma}_k - (\mu_k - \bar{\mu}_k)(\mu_k - \bar{\mu}_k)^\top$.
- ii) $\hat{\mu}_k^* = \bar{\mu}_k + \bar{\nu}_k$.

Proof. Please see Appendix E. $\square$

Substituting the expressions of Σ_k^* and $\hat{\mu}_k^*$ into $\mathbf{P}_2$, it is equivalent to a convex-nonconcave minimax optimization:

$$\begin{aligned} & \min_{\hat{\Sigma}_k \succ 0} \max_{a, b} -\log \det(\hat{\Sigma}_k^{-1}) + \gamma_2 \bar{\Sigma}_k \cdot \hat{\Sigma}_k^{-1} \\ & + \|a + b\|_{\hat{\Sigma}_k^{-1}}^2 - \|a\|_{\hat{\Sigma}_k^{-1}}^2 \quad (12) \\ \text{s.t. } & \|a\|_{\hat{\Sigma}_k^{-1}}^2 \leq \gamma_1, \|b\|_2^2 \leq \gamma_0(z), \end{aligned}$$

where $a = \mu_k - \bar{\mu}_k, b = \nu_k - \bar{\nu}_k$. Notice that the inner maximization is an indefinite quadratic constrained quadratic programming (QCQP) problem, and there are still short of efficient algorithms to solve this kind of minimax optimization.

Given any $\hat{\Sigma}_k \succ 0$ and $b \neq 0$, one can first solve a by considering the following quadratic constrained linear programming (QCLP) problem:

$$\max_a \|a + b\|_{\hat{\Sigma}_k^{-1}}^2 - \|a\|_{\hat{\Sigma}_k^{-1}}^2 \quad \text{s.t. } \|a\|_{\hat{\Sigma}_k^{-1}}^2 \leq \gamma_1. \quad (13)$$

Lemma 2. The solution of the QCLP problem (13) is

$$a^* = \frac{\sqrt{\gamma_1} \bar{\Sigma}_k \hat{\Sigma}_k^{-1} b}{\|\bar{\Sigma}_k \hat{\Sigma}_k^{-1} b\|_{\hat{\Sigma}_k^{-1}}}.$$

Proof. Using the KKT condition, there are

$$\begin{cases} \partial_a \left[\|a + b\|_{\hat{\Sigma}_k^{-1}}^2 - \|a\|_{\hat{\Sigma}_k^{-1}}^2 - s_1 (\|a\|_{\hat{\Sigma}_k^{-1}}^2 - \gamma_1) \right] = 0 \\ s_1 (\|a\|_{\hat{\Sigma}_k^{-1}}^2 - \gamma_1) = 0, s_1 \geq 0 \end{cases}$$

The first condition reveals that $s_1 \neq 0$ and $a = s_1^{-1} \bar{\Sigma}_k \hat{\Sigma}_k^{-1} b$, thus the second condition leads to $\|a\|_{\hat{\Sigma}_k^{-1}}^2 = \gamma_1$. Then,

$$s_1 = \frac{\|\bar{\Sigma}_k \hat{\Sigma}_k^{-1} b\|_{\hat{\Sigma}_k^{-1}}}{\|a\|_{\hat{\Sigma}_k^{-1}}} = \frac{\|\bar{\Sigma}_k \hat{\Sigma}_k^{-1} b\|_{\hat{\Sigma}_k^{-1}}}{\sqrt{\gamma_1}},$$

and the proof is completed. $\square$

Substituting the optimal a^* into the problem, the inner optimization of (12) can be further reformulated as a convex maximization problem:

$$\max_b \alpha \|b\|_A + \|b\|_B^2 \quad \text{s.t. } \|b\|_2^2 \leq \gamma_0(z), \quad (14)$$

where $A = \hat{\Sigma}_k^{-1} \bar{\Sigma}_k \hat{\Sigma}_k^{-1}$, $B = \hat{\Sigma}_k^{-1}$, and $\alpha = 2\sqrt{\gamma_1}$.

D. Complexity Analysis

As we have shown, solving the one-step DRPP $\mathbf{P}_1$ is equivalent to finding a global minimax point (Stackelberg equilibrium) for a convex-nonconcave minimax problem where the outer minimization is a semidefinite programming and the inner problem is to maximize a non-smooth convex function subject to a quadratic constraint. It is known that convex maximization is NP-hard in very simple cases, such as quadratic maximization over a hypercube, and even verifying local optimality is NP-hard [41]. Thus, finding a global minimax point of a one-step DRPP problem is also NP-hard.

Even worse, finding a local minimax point of a general constrained minimax optimization is of fundamental hardness. Not only may the local minimax point not exist [42], but any gradient-based algorithm needs exponentially many queries in the dimension and ϵ^{-1} to compute an ϵ -approximate first-order

local stationary point [43]. Finally, the local stationary point is not always guaranteed to be a local minimax point [44], [45].

In summary, globally finding a minimax point of a one-step DRPP is computationally intractable, and gradient-based algorithms aiming for local minimax points are limited by both convergence speed and performance guarantee. Rather than pursuing the optimal solutions of $\mathbf{P}_0$, a more reasonable goal is to derive suboptimal solutions for the one-step DRPP $\mathbf{P}_1$. Based on these suboptimal solutions, we can then derive upper and lower bounds for the robust value functions V_k^* .

VI. VALUE FUNCTION UPPER BOUND

To get an upper bound of $V_k^*(z)$, one can either enlarge the feasible region of the outer maximization or shrink the feasible region of the inner minimization. In this paper, we choose to shrink the ambiguity set $\mathcal{I}_k(z)$. If the first constraint is strengthened as $\nu_k = \bar{\nu}_k$, one is constrained to a smaller ambiguity subset of $\mathcal{I}_k(z)$. If a global minimax point can be solved for the relaxed one-step DRPP, one will get a suboptimal predictor and a performance upper bound.

A. Problem Reformulation

If we shrink the ambiguity set of one-step DRPP $\mathbf{P}_1$ by assuming $\bar{\nu}_k = \nu_k$, there is $\mathbf{w}_k = \mathbf{x}_{k+1} - \bar{\nu}_k$, which means that $\mathbf{w}_k$ can be precisely known after $\mathbf{x}_{k+1}$ been observed. Now that there is no uncertainty of the state evolution ν_k , predicting the state's conditional probability measure $P_{\mathbf{x}_{k+1}|\mathbf{z}_k}$ is equivalent to predicting the noise' probability measure $P_{\mathbf{w}_k}$, and the one-step DRPP $\mathbf{P}_1$ is then reformulated as

$$\begin{aligned} & \sup_{\hat{p}_k \in \mathcal{F}} \inf_{\mu_k, \bar{p}_k} \int_{\mathcal{X}} \log[\hat{p}_k(w + \bar{\nu}_k)] d\bar{p}_k(w) \\ & \text{s.t.} \begin{cases} \bar{p}_k \in \mathcal{M}_+(\mathcal{X}), \mathbb{E}_{w \sim \bar{p}_k}[1] = 1, \mathbb{E}_{w \sim \bar{p}_k}[w] = \mu_k \\ \gamma_3 \bar{\Sigma}_k \preceq \mathbb{E}_{w \sim \bar{p}_k}[(w - \bar{\mu}_k)(w - \bar{\mu}_k)^\top] \preceq \gamma_2 \bar{\Sigma}_k \\ \begin{bmatrix} \bar{\Sigma}_k & (\mu_k - \bar{\mu}_k) \\ (\mu_k - \bar{\mu}_k)^\top & \gamma_1 \end{bmatrix} \succeq 0. \end{cases} \end{aligned} \quad (15)$$

With the outer maximization imposed on $\hat{p}_k$, we can get a reformulated one-step DRPP whose optimal worst-case value function is an upper bound of the objective of $\mathbf{P}_1$.

B. Noise-DRPP

Taking the dual form of the inner problem of (15) and joining it with the outer maximization, we transform the primal maximin problem into the following maximization problem:

$$\begin{aligned} (\mathbf{D}_2) : & \sup_{\hat{p}_k \in \mathcal{F}, r, q, Q_1, Q_2, P, p, s} G(r, q, Q_1, Q_2, P, p, s) \\ & \text{s.t.} \begin{cases} w^\top(Q_1 - Q_2)w + w^\top q + r + \log \hat{p}_k(w + \bar{\nu}_k) \geq 0 \\ \quad \quad \quad \forall w \in \mathbb{R}^{d_x} \\ q + 2(Q_1 - Q_2)\bar{\mu}_k + 2p = 0 \\ Q_1 \succeq 0, Q_2 \succeq 0 \\ \begin{bmatrix} P & p \\ p^\top & s \end{bmatrix} \succeq 0, \end{cases} \end{aligned}$$

where $G(r, q, Q_1, Q_2, P, p, s) = -r + \bar{\mu}_k^\top(Q_1 - Q_2)\bar{\mu}_k - (\gamma_2 Q_1 - \gamma_3 Q_2 + P) \cdot \bar{\Sigma}_k + 2p^\top \bar{\mu}_k - s\gamma_1$. Please see Appendix F for a detailed derivation of the dual problem.

Theorem 3 (Noise-DRPP). *Given $\mathbf{z}_k = z \in \mathcal{Z}$ at time step $k \in \{0, \dots, T-1\}$, if the probabilistic predictor has no ambiguity of the one-step state evolution, i.e., $\bar{\nu}_k = \nu_k$, the solution to the one-step DRPP $\mathbf{P}_1$ is:*

i) *The optimal predictive pdf is*

$$\hat{p}_k^* \sim \mathcal{N}(\bar{\nu}_k + \bar{\mu}_k, \gamma_2 \bar{\Sigma}_k).$$

ii) *The worst-case conditional measure ρ_k^* belongs to the set*

$$\left\{ P_{\mathbf{x}_{k+1}|\mathbf{z}_k}(\cdot | z) \mid \begin{matrix} \mathbf{x}_{k+1} = \bar{f}_k(z) + \mathbf{w}_k \\ \mathbb{E}_{w \sim P_{\mathbf{w}_k}}[(w - \bar{\mu}_k)(w - \bar{\mu}_k)^\top] = \gamma_2 \bar{\Sigma}_k \end{matrix} \right\}.$$

iii) *The objective function at $(\hat{p}_k^*, \rho_k^*)$ is*

$$-\frac{1}{2} [d_x \log(2\pi) + d_x + \log \det(\gamma_2 \bar{\Sigma}_k)].$$

Proof. Please see Appendix G. $\square$

First, an immediate application of Theorem 3 is to utilize the explicit solution of $\hat{p}_k^*$ to develop a probabilistic prediction algorithm. Because this predictor solely focuses on the distributional robustness against the ambiguity of system noises, we name it Noise-DRPP. The pseudo-code is in Algorithm 1.

Second, notice that the optimal predictive pdf is unique, but the worst-case SDS is not. It contains any SDS that has a proper one-step state evolution and a noise whose first two moments satisfy an equation.

Finally, which is also our original goal, the objective function at $(\hat{p}_k^*, \rho_k^*)$ is an upper bound of the optimal objective of $\mathbf{P}_1$. Since the one-step upper bound does not depend on z , one can recursively use the Bellman equation (7) to get an upper bound of the robust value function V_k^* .

Theorem 4 (Upper bound of V_k^*). *Given an initial state-control pair $z \in \mathcal{Z}$ at time step $k \in \{0, \dots, T-1\}$, an upper bound of the robust value function $V_k^*(z)$ is*

$$V_k^*(z) \leq \sum_{t=k}^{T-1} -\frac{1}{2} [d_x \log(2\pi) + d_x + \log \det(\gamma_2 \bar{\Sigma}_t)].$$

Remark 2. *We highlight that this upper bound can be computed offline because it only depends on the parameters of ambiguity sets $\mathcal{I}_{0:T-1}$. Furthermore, the conservatism of this upper bound can be evaluated by the gap between it and another lower bound, as developed in the next section.*

VII. VALUE FUNCTION LOWER BOUND

To get a lower bound of $V_k^*(z)$, one can either shrink the feasible region of the outer maximization or enlarge the feasible region of the inner minimization. Because enlarging the ambiguity set does not essentially change the structure to alleviate the difficulty of a one-step DRPP $\mathbf{P}_2$, we choose to restrict $\hat{p}_k$ to a smaller pdf family by forcing the eigenvectors of the predictive covariance $\hat{\Sigma}_k$ to be the same as the nominal covariance Σ_k . In this section, we elaborate on how the eigenvector restriction contributes to a well-performed algorithm and a lower bound.

Algorithm 1 Noise-DRPP $\mathcal{F}_{\text{Noise}}$

Input: time horizon T , ambiguity sets $\mathcal{I}_{0:T-1}$, control policy $\pi_{0:T-1}$, initial state x_0 .

- 1: $s_0 \leftarrow 0$. $\triangleright$ score at time step 0
- 2: **for** $k = 0, \dots, T-1$ **do**
- 3: Predictor updates ambiguity set $\mathcal{I}_k$ as (3).
- 4: SDS generates control input $u_k \sim \pi_k(\cdot | x_k)$.
- 5: $z_k \leftarrow (x_k, u_k)$, $\bar{\nu}_k \leftarrow \bar{f}_k(z_k)$.
- 6: Predictor predicts $\hat{p}_k^* \sim \mathcal{N}(\bar{\nu}_k + \bar{\mu}_k, \gamma_2 \bar{\Sigma}_k)$.
- 7: SDS generates the next state x_{k+1} .
- 8: $s_{k+1} \leftarrow s_k + \mathcal{L}(\hat{p}_k^*, x_{k+1})$.
- 9: **end for**

Output: states $x_{0:T}$, predictions $\hat{p}_{0:T-1}^*$, scores $s_{0:T}$.

A. Problem Reformulation

By adding auxiliary constraints to the predictor in (12), one can still get a suboptimal prediction performance guarantee in the worst case, which is also a performance lower bound to the original one-step DRPP. Nevertheless, the choice of relaxation methods is delicate, which can greatly affect both the optimality gap and the computation efficiency.

Suppose the spectral decomposition for $\bar{\Sigma}_k$ is $Q_k \Lambda_k Q_k^\top$, where $Q_k = [\mathbf{v}_{1,k}, \dots, \mathbf{v}_{d_x,k}]$ is an orthogonal matrix whose columns are eigenvectors of $\bar{\Sigma}_k$ and $\Lambda_k = \text{diag}\{\lambda_{1,k}, \dots, \lambda_{d_x,k}\}$. Similarly, we let the spectral decomposition for $\hat{\Sigma}_k$ be $\hat{Q}_k \hat{\Lambda}_k \hat{Q}_k^\top$, where $\hat{Q}_k$ is an orthogonal matrix whose columns are eigenvectors of $\hat{\Sigma}_k$ and $\hat{\Lambda}_k = \text{diag}\{\hat{\lambda}_{1,k}, \dots, \hat{\lambda}_{d_x,k}\}$. Then, the eigenvector restriction for (12) is to impose the following constraint on the predictor:

$$\hat{Q}_k = Q_k. \quad (16)$$

To explain why we chose the eigenvector restriction, some supporting lemmas need to be derived.

Lemma 3. *The optimal b^* to problem (14) is an eigenvector of $\frac{\alpha}{\|b^*\|_A} A + 2B$ with its ℓ^2 -norm being $\sqrt{\gamma_0(z)}$.*

Proof. There are necessary optimality conditions for (14):

$$\begin{cases} \left(\frac{\alpha}{\|b\|_A} A + 2B - 2sI \right) b = 0 \\ \|b\|_2^2 = \gamma_0(z), \end{cases}$$

where the first one comes from the KKT condition and the second one holds because the objective increases monotonously with $\|b\|_2$. It follows that the optimal b should be an eigenvector of $\frac{\alpha}{\|b\|_A} A + 2B$ with its ℓ^2 -norm being $\sqrt{\gamma_0(z)}$. $\square$

If the eigenvectors of $\frac{\alpha}{\|b^*\|_A} A + 2B$ have explicit expressions, one may have explicit expressions for a^* , b^* . Then, substituting these expressions into (12) can transform the original minimax optimization problem into a semidefinite-constrained maximization, where tractable optimization solvers may be available. However, it is always difficult to explicitly derive the eigenvectors when $\hat{Q}_k$ does not identify with Q_k .

Lemma 4. *When $\hat{Q}_k = Q_k$ is supplemented to the constraints of (14), the solution is*

$$b^* = \sqrt{\gamma_0(z)} \mathbf{v}_{j_k,k},$$

where $j_k = \arg \max_i (2\sqrt{\gamma_0(z)} \gamma_1 \lambda_{i,k} + \gamma_0(z)) \hat{\lambda}_{i,k}^{-1}$.

Proof. When $\hat{Q}_k = Q_k$ is satisfied, there is

$$\frac{\alpha}{\|b\|_A} A + 2B = Q_k \left[\frac{\alpha}{\|b\|_A} \hat{\Lambda}_k^{-2} \Lambda_k + \hat{\Lambda}_k^{-1} \right] Q_k^\top,$$

whose eigenvectors are the columns of Q_k . Lemma 3 indicates that the optimal b can be expressed as $\sqrt{\gamma_0(z)} \mathbf{v}_i$ for certain $i \in \{1, \dots, d_x\}$, then we have the objective of (14) is $(2\sqrt{\gamma_0(z)} \gamma_1 \lambda_{i,k} + \gamma_0(z)) \hat{\lambda}_{i,k}^{-1}$. Since $j_k = \arg \max_i (2\sqrt{\gamma_0(z)} \gamma_1 \lambda_{i,k} + \gamma_0(z)) \hat{\lambda}_{i,k}^{-1}$, problem (14) is solved as $b^* = \sqrt{\gamma_0(z)} \mathbf{v}_{j_k}$. $\square$

Utilizing the expression of b^* in Lemma 4, the minimax optimization (12) under eigenvector restriction (16) can be transformed into

$$\begin{aligned} (\mathbf{P}_3) : \min_{\hat{\Lambda}_k \succeq 0, j_k} \sum_{i=1}^{d_x} & \left[-\log(\hat{\lambda}_{i,k}^{-1}) + \gamma_2 \lambda_{i,k} \hat{\lambda}_{i,k}^{-1} \right] + \\ & \left(2\sqrt{\gamma_0(z)} \gamma_1 \lambda_{j_k,k} + \gamma_0(z) \right) \hat{\lambda}_{j_k,k}^{-1} \\ \text{s.t. } \hat{\lambda}_{j_k,k} & \leq \frac{2\sqrt{\gamma_0(z)} \gamma_1 \lambda_{j_k,k} + \gamma_0(z)}{2\sqrt{\gamma_0(z)} \gamma_1 \lambda_{i,k} + \gamma_0(z)} \hat{\lambda}_{i,k} \\ & \forall i \in \{1, \dots, d_x\}. \end{aligned}$$

B. Eig-DRPP

Notice that optimization $\mathbf{P}_3$ is convex for any fixed j_k . Therefore, $\mathbf{P}_3$ can be solved by taking the minimum of d_x different convex optimization problems, which is numerically efficient. Let $(\hat{\Lambda}_k^* = \text{diag}\{\hat{\lambda}_{1,k}^*, \dots, \hat{\lambda}_{d_x,k}^*\}, j_k^*)$ be the solution of $\mathbf{P}_3$, we summarize the solution of a one-step DRPP under eigenvector restriction as follows.

Theorem 5 (Eig-DRPP). *Given $z_k = z \in \mathcal{Z}$ at time step $k \in \{0, \dots, T-1\}$, if the probabilistic predictor is constrained by the eigenvector restriction (16), the solution to the one-step DRPP $\mathbf{P}_1$ is:*

i) *The optimal predictive pdf $\hat{p}_k^*$ is*

$$\hat{p}_k^* \sim \mathcal{N}(\bar{\nu}_k + \bar{\mu}_k, Q_k \hat{\Lambda}_k^* Q_k^\top).$$

ii) *The worst-case conditional measure ρ_k^* belongs to the set*

$$\left\{ P_{\mathbf{x}_{k+1}|z_k}(\cdot | z) \left| \begin{array}{l} \mathbf{x}_{k+1} = f_k(z) + \mathbf{w}_k \\ f_k(z) = \sqrt{\gamma_0(z)} \mathbf{v}_{j_k^*,k} + \bar{\nu}_k \\ \mathbb{E}[\mathbf{w}_k] = \bar{\mu}_k + \sqrt{\gamma_1 \lambda_{j_k^*,k}} \mathbf{v}_{j_k^*,k} \\ \text{Cov}[\mathbf{w}_k] = \gamma_2 \bar{\Sigma}_k - \gamma_1 \lambda_{j_k^*,k} \mathbf{v}_{j_k^*,k} \mathbf{v}_{j_k^*,k}^\top \end{array} \right. \right\}.$$

iii) *The objective function at $(\hat{p}_k^*, \rho_k^*)$ is*

$$-\frac{1}{2} \left\{ d_x \log(2\pi) + \sum_{i=1}^{d_x} \left[\log(\hat{\lambda}_{i,k}^*) + \gamma_2 \frac{\lambda_{i,k}}{\hat{\lambda}_{i,k}^*} \right] + \frac{2\sqrt{\gamma_0(z)} \gamma_1 \lambda_{j_k^*,k} + \gamma_0(z)}{\hat{\lambda}_{j_k^*,k}^*} \right\}.$$

Proof. The proof is completed by combining the results in Lemma 1, Lemma 2, and the solutions of $\mathbf{P}_3$. $\square$

First, an immediate application of Theorem 5 is to utilize the explicit solution of $\hat{p}_k^*$ to develop a probabilistic prediction algorithm, named Eig-DRPP. The pseudo-code is presented in Algorithm 2. Second, although the optimal predictive pdf is unique, the worst-case conditional measure is not. Compared to Theorem 3, there are relatively more restrictions on the set of worst-case conditional measures. Third, when $\gamma_0 = 0$, i.e., the predictor has no ambiguity of f_k , the solution of $\mathbf{P}_3$ is $\hat{\lambda}_{i,k}^* = \gamma_2 \lambda_{i,k}$ and j_k^* can be any feasible index. In this case, both the optimal predictive pdf and objective value in Theorem 5 coincide with those in Theorem (3). However, the sets of worst-case conditional measures are not the same, due to the extra constraint for the expectation in Theorem 5.

Finally, the objective function at $(\hat{p}_k^*, \rho_k^*)$ is a lower bound of the optimal objective of $\mathbf{P}_1$. Since the objective value depends on z through $\gamma_0(z)$, one cannot directly use the Bellman equation (7) recursively to get a lower bound of V_k^* .

Theorem 6 (Lower bound of V_k^* and optimality gap of $\mathcal{F}_{\text{Eig}}$). *Given an initial state-control pair $z \in \mathcal{Z}$ at time step $k \in \{0, \dots, T-1\}$, and an upper bound of γ_0 such that $\gamma_0(z) \leq \Gamma_0 \in \mathbb{R}_+$, $\forall z \in \mathcal{Z}$, there is*

i) *A lower bound of the robust value function $V_k^*(z)$ is*

$$V_k^*(z) \geq \sum_{t=k}^{T-1} -\frac{1}{2} \left\{ d_x \log(2\pi) + \sum_{i=1}^{d_x} \left[\log(\hat{\lambda}_{i,t}^*) + \gamma_2 \frac{\lambda_{i,t}}{\hat{\lambda}_{i,t}^*} \right] + \frac{2\sqrt{\Gamma_0 \gamma_1 \lambda_{j_t^*,t}} + \Gamma_0}{\hat{\lambda}_{j_t^*,t}^*} \right\},$$

where $\hat{\lambda}_{1,t}^*, \dots, \hat{\lambda}_{d_x,t}^*, j_t^*$ is the solution to $\mathbf{P}_3$ with k being replaced by t and all $\gamma_0(z)$ being replaced by Γ_0 .

ii) *The optimality gap between Eig-DRPP $\mathcal{F}_{\text{Eig}}$ and the global optimal predictor is upper bounded as follows,*

$$V_k^*(z) - V_k^{\mathcal{F}_{\text{Eig}}}(z) \leq \sum_{t=k}^{T-1} -\frac{1}{2} \left\{ \underbrace{d_x + \log \det(\gamma_2 \bar{\Sigma}_t)}_{\textcircled{1}} - \underbrace{\sum_{i=1}^{d_x} \left[\log(\hat{\lambda}_{i,t}^*) + \gamma_2 \frac{\lambda_{i,t}}{\hat{\lambda}_{i,t}^*} \right] - \frac{2\sqrt{\Gamma_0 \gamma_1 \lambda_{j_t^*,t}} + \Gamma_0}{\hat{\lambda}_{j_t^*,t}^*}}_{\textcircled{2}} \right\}. \quad (17)$$

Proof. i) Notice that as $\gamma_0(z)$ increases, the ambiguity set enlarges, thus the optimal value of one-step DRPP under eigenvector constraint should monotonously decrease. Since the upper bound of $\gamma_0(z)$, i.e., Γ_0 , is independent of z , one can still use (7) to derive a lower bound.

ii) Subtracting this lower bound from the upper bound in Theorem 4, one immediately gets an upper bound of the optimality gap for Eig-DRPP. As illustrated in (17), $\textcircled{1}$ comes from the upper bound, which is determined by γ_2 and the determinant of $\bar{\Sigma}_t$. While $\textcircled{2}$ comes from the lower bound, which is jointly determined by $\gamma_0, \gamma_1, \gamma_2$ and the eigenvalues of $\bar{\Sigma}_t$ by solving $\mathbf{P}_3$. $\square$

The proposed DRPP predictors can be naturally integrated into SMPC. In particular, the Noise-DRPP predictor provides a closed-form Gaussian predictive model that can be used

to construct distributionally robust chance constraints. The Eig-DRPP predictor, while computationally more involved, can serve as an offline surrogate to characterize the trade-off between robustness and conservatism in SMPC design.

Algorithm 2 Eig-DRPP $\mathcal{F}_{\text{Eig}}$

Input: time horizon T , ambiguity sets $\mathcal{I}_{0:T-1}$, control policy $\pi_{0:T-1}$, initial state x_0 .

- 1: $s_0 \leftarrow 0$. $\triangleright$ score at time step 0
- 2: **for** $k = 0, \dots, T-1$ **do**
- 3: Predictor updates ambiguity set $\mathcal{I}_k$ as (3).
- 4: Do spectral decomposition for $\bar{\Sigma}_k = Q_k \Lambda_k Q_k^\top$.
- 5: SDS generates control input $u_k \sim \pi_k(\cdot | x_k)$.
- 6: $z_k \leftarrow (x_k, u_k)$, $\bar{\nu}_k \leftarrow \bar{f}_k(z_k)$.
- 7: $j_k^* \leftarrow 0$, $\text{val}^* \leftarrow \infty$, $\hat{\Lambda}_k^* \leftarrow \Lambda_k$.
- 8: **for** $j = 1, 2, \dots, d_x$ **do**
- 9: Fix $j_k = j$, solve $\mathbf{P}_3$, where the optimizer is $\hat{\Lambda}_k^{(j)}$ and the optimal value is $\text{val}^{(j)}$. $\triangleright$ For each j , this step is a convex optimization
- 10: **if** $\text{val}^{(j)} < \text{val}^*$ **then**
- 11: $j_k^* \leftarrow j$, $\text{val}^* \leftarrow \text{val}^{(j)}$, $\hat{\Lambda}_k^* \leftarrow \hat{\Lambda}_k^{(j)}$.
- 12: **end if**
- 13: **end for**
- 14: Predictor predicts $\hat{p}_k^* \sim \mathcal{N}(\bar{\nu}_k + \bar{\mu}_k, Q_k \hat{\Lambda}_k^* Q_k^\top)$.
- 15: SDS generates the next state x_{k+1} .
- 16: $s_{k+1} \leftarrow s_k + \mathcal{L}(\hat{p}_k^*, x_{k+1})$.
- 17: **end for**

Output: states $x_{0:T}$, predictions $\hat{p}_{0:T-1}^*$, scores $s_{0:T}$.

VIII. EXPERIMENTS

In this section, a series of experiments is conducted to explore three questions. i) Given an SDS subject to an ambiguity set, what are the performance advantages of Noise-DRPP and Eig-DRPP compared to a nominal predictor? ii) How much will the prediction performances be influenced by control strategies? iii) How can DRPP predictors be applied to providing high probability confidence regions of future states?

A. Experiment Setup

Ambiguity set: Conditioned on a state-control pair z_k at time step $k \in \{0, 1, \dots, 31\}$, the nominal state evolution function is given as

$$\bar{f}_k(x_k, u_k) = \begin{pmatrix} 1 & 0.1 \\ 0 & 1 \end{pmatrix} x_k + \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} u_k,$$

and the uncertainty between $\bar{f}_k$ and the true f_k is quantified by $\gamma_0(z_k) = \min\{0.3\|z_k\|_2, 5\}^2$. The nominal mean and covariance of the noise w_k conditioned on z_k are $\bar{\mu}_k = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$, $\bar{\Sigma}_k = \begin{pmatrix} 1 & 0.5 \\ 0.5 & 1.5 \end{pmatrix}$, and the uncertainty between $\bar{\mu}_k, \bar{\Sigma}_k$ and the real μ_k, Σ_k are quantified by $\gamma_1 = 0.5$ and $\gamma_2 = 3$, respectively. Given the ambiguity set specified above, the real SDS belongs to $\mathcal{I}_k$. Although the nominal model is linear time-invariant (LTI), the real model does not have to be LTI.

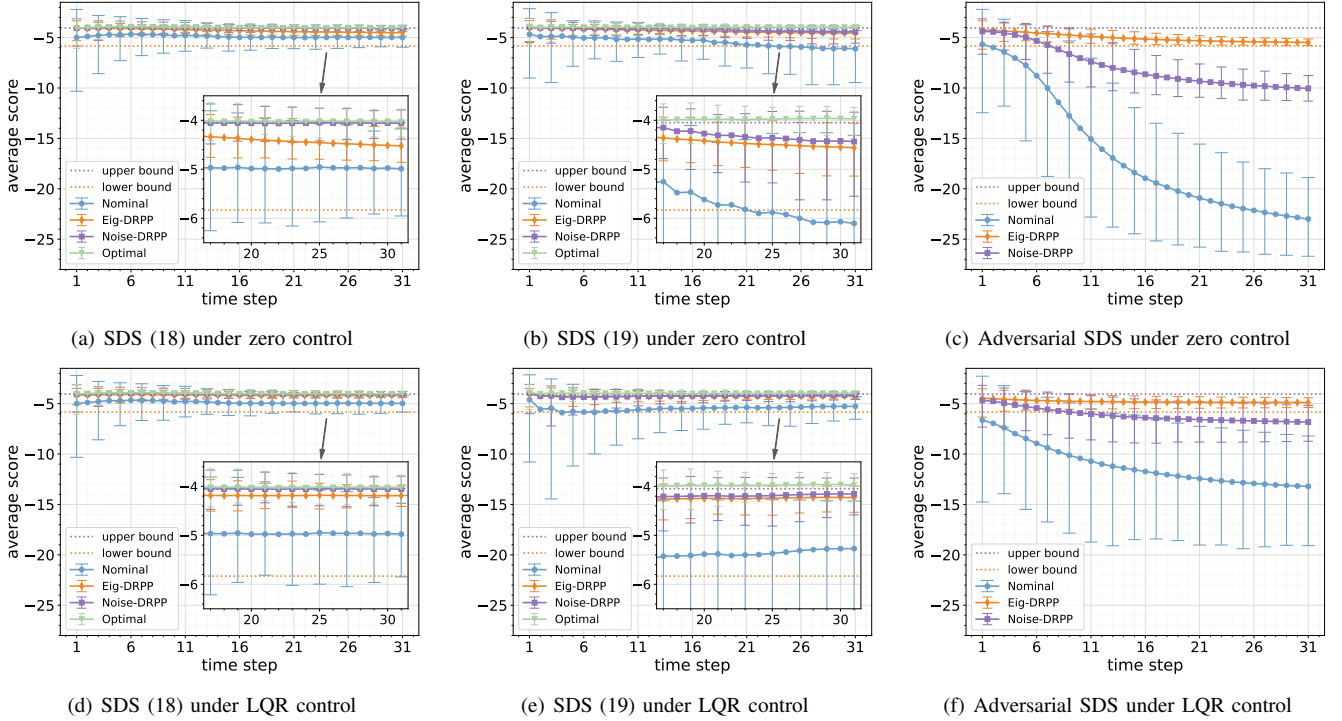


Fig. 2: Prediction performance of different probabilistic predictors on different SDSs under different control strategies.

Simulation mechanism: Three different simulation mechanisms are conducted to generate the underlying SDSs.

i) Let α_1, α_2 be two independent random variables subject to the uniform distribution on $[-1, 1]$, and let α_3 be another independent random variable subject to the uniform distribution on $[-1, 1]^2$. We randomly generate $\alpha_i \sim p_{\alpha_i}$ for $i = 1, 2, 3$, then we simulate the SDS as

$$\mathbf{x}_{k+1} = \begin{pmatrix} 1 & 0.1 + 0.3\alpha_1 \\ 0 & 1 \end{pmatrix} \mathbf{x}_k + \begin{pmatrix} 1 & 0.3\alpha_2 \\ 0 & 1 \end{pmatrix} \mathbf{u}_k + \mathbf{w}_k, \quad (18)$$

where $\mathbf{w}_k \sim \mathcal{N}(0.5\alpha_3, 3\bar{\Sigma}_k - 0.25\alpha_3\alpha_3^\top)$.

ii) For each $i = 1, 2, 3$, let $\{\alpha_{i,k}\}_{k=0}^{31}$ be random variables that are independent and identically distributed to p_{α_i} . We randomly generate $\alpha_{i,k} \sim p_{\alpha_i}$ for each i and k , then we simulate the SDS as

$$\mathbf{x}_{k+1} = \begin{pmatrix} 1 & 0.1 + 0.3\alpha_{1,k} \\ 0 & 1 \end{pmatrix} \mathbf{x}_k + \begin{pmatrix} 1 & 0.3\alpha_{2,k} \\ 0 & 1 \end{pmatrix} \mathbf{u}_k + \mathbf{w}_k, \quad (19)$$

where $\mathbf{w}_k \sim \mathcal{N}(0.5\alpha_{3,k}, 3\bar{\Sigma}_k - 0.25\alpha_{3,k}\alpha_{3,k}^\top)$.

iii) (Adversarial) In this case, the underlying SDS is allowed to adversarially choose an SDS in the ambiguity set at each step to degrade the prediction performance. Specifically, at each step, after a predictive pdf has been output by the predictor, the adversarial SDS uses the predicted covariance to choose a worst-case SDS as described in Theorem 5.

Predictors comparison: Two different control strategies are considered, the zero input and the linear quadratic regulation (LQR), where the zero input strategy means no control input is imposed, and the LQR control strategy is defined as a standard linear quadratic regulator based on the nominal linear model with Q and R all set as identity matrices. Under

each control strategy and SDS scenario, we randomly generate 1,000 trajectories starting from the same initial state $\mathbf{x}_0 = [2, 1]$. At each time step k , let the nominal predictor predicts $\bar{p}_k \sim \mathcal{N}(\bar{\nu}_k + \bar{\mu}_k, \bar{\Sigma}_k)$, and the optimal predictor predicts $p_k \sim \mathcal{N}(\nu_k + \mu_k, \Sigma_k)$. Then, we compare the prediction performance among the nominal predictor, Noise-DRPP, Eig-DRPP, and the optimal predictor. Notice that for the adversarial mechanism, there is no explicitly predefined optimal predictor because the real SDS depends on the predictor. Therefore, we should compare the prediction performance among the nominal predictor, Noise-DRPP, and Eig-DRPP.

B. Results and Discussions

In Fig. 2, the prediction performances of different probabilistic predictors on different SDSs under different control strategies are visualized and compared with each other. For each SDS, we evaluate the temporal average scores of each predictor at each time step on 1,000 randomly sampled trajectories. At each step, the mean of 1,000 average scores is plotted as a dot, which is an approximation to the expected average score. The 5% and 95% percentiles of these average scores are plotted as an error bar, reflecting how concentrated the average scores are.

Performance advantages: Under each setting of SDS and control strategy, both Noise-DRPP and Eig-DRPP possess larger expected average scores at each time step compared to the nominal predictor. Since the ambiguity set for our simulation does not have too large an uncertainty, the scores of both DRPP predictors are close to the optimal predictor. Additionally, both DRPP predictors have more concentrated average scores at each time step. As the uncertainty of

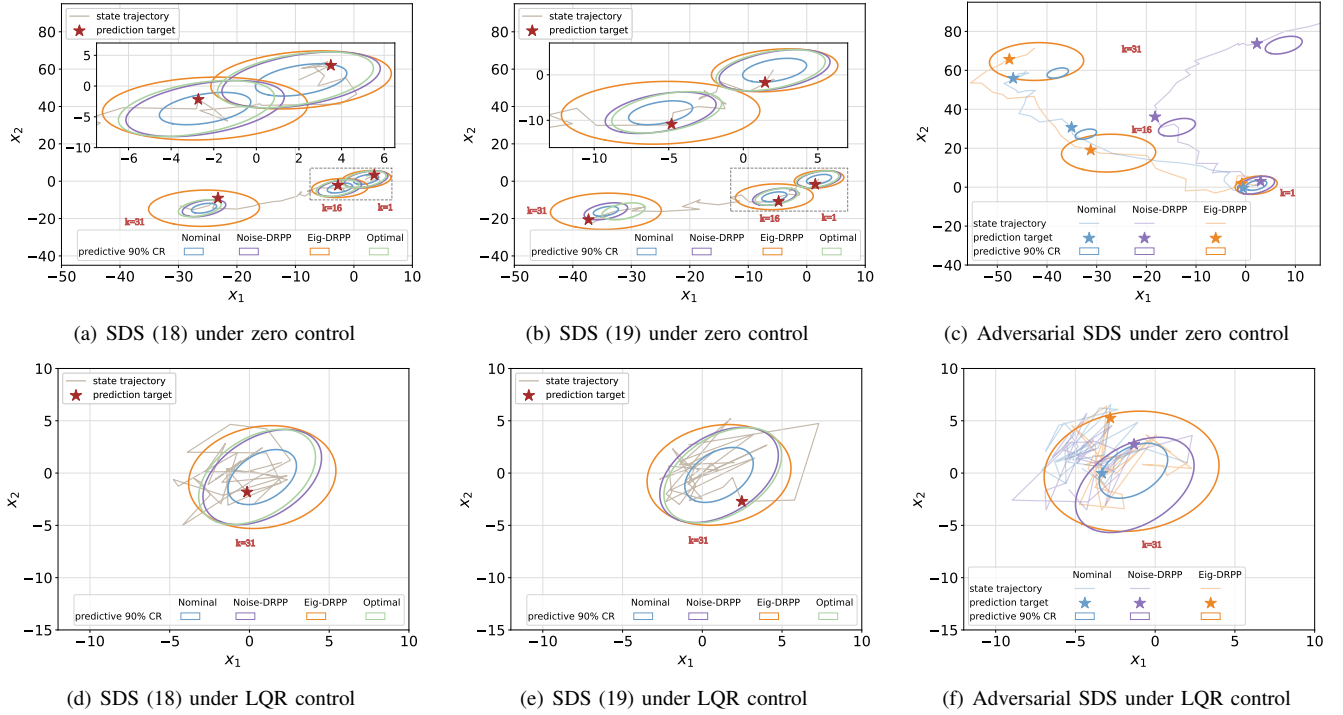


Fig. 3: Predictive 90% confidence regions of probabilistic predictors for different SDSs under different control strategies.

the real underlying SDS increases from LTI to adversarial, the performance gap between the DRPP predictors and the nominal predictor increases.

Influence of control strategies: When a predictor has no ambiguity about the system's state evolution function, predicting the states is equivalent to predicting the noises. In this case, the choice of control strategies does not affect any probabilistic predictor. However, when a predictor does have ambiguity of the system's state evolution function, different control strategies can greatly influence the performance of a probabilistic predictor. Since the ambiguity set at each step directly depends on the state-input pair, different control strategies will lead to quite different state-input trajectories, thus leading to quite different prediction performances. We can observe this phenomenon by comparing (a-c) with (d-f) in Fig. 2 respectively. An intuitive explanation is that the LQR control strategy regulates the system states towards the neighborhoods of 0, where the ambiguity of one-step state evolution is smaller than other states, thus the DRPP predictors perform better.

C. Application: Predict Robust Confidence Regions

An immediate application of probabilistic prediction is to predict high-probability confidence regions for future states. Since the predictive pdfs of both Noise-DRPP and Eig-DRPP are Gaussian, one can easily derive an elliptical β -confidence region for any $\beta \in (0, 1)$. In Fig. 3, we have visualized the predictive 90% confidence regions for each predictor at certain time steps of a randomly generated trajectory. For the zero control scenario, we visualize the predictions for time steps $k = 1, 16, 31$, while for the LQR control scenario, we visualize the prediction for time step $k = 31$.

Visualizing the behaviors of different probabilistic predictors on a trajectory, one can intuitively explain why DRPP predictors outperform the nominal predictor. The nominal predictor is overly confident, where the predictive pdfs are always too sharp to enclose the future states in their 90% CRs. In contrast, DRPP predictors have taken the ambiguity of noises into account, whose CRs are more robust to contain the prediction targets most of the time. If we further compare Noise-DRPP with Eig-DRPP, it seems that Eig-DRPP is more intelligent in adaptively changing the ellipses' shapes. This is because Eig-DRPP has also considered the ambiguity of one-step state evolution, and the eigenvalues of predictive covariances are attained based on minimax optimization.

D. Extension: Nonlinear System

We consider a 6-degree-of-freedom serial manipulator with revolute joints using the standard Euler-Lagrange rigid-body equations [46]. The coordinates are $q \in \mathbb{R}^6$ and velocities $\dot{q} \in \mathbb{R}^6$. The continuous-time dynamics are

$$M(q)\ddot{q} + C(q, \dot{q})\dot{q} + g(q) = \tau,$$

where $M(q)$ is the inertia matrix, $C(q, \dot{q})\dot{q}$ represents the Coriolis forces, $g(q)$ is the gravity vector, and $\tau \in \mathbb{R}^6$ denotes applied joint torques. The mass $m_i = 0.8\text{kg}$, link length $l_i = 0.5\text{m}$, and link center of mass (COM) distance $l_{c,i} = 0.25\text{m}$ for each link $i = 1, \dots, 6$. We write the full state $x = [q^\top, \dot{q}^\top]^\top \in \mathbb{R}^{12}$.

The nominal model adopts a lightweight numerical approximation that preserves the main physical structures while remaining efficient for repeated DRPP computations. For the inertia matrix, a diagonal-dominant rigid-body approximation

is constructed by summing per-link COM contributions, i.e., $M_{ij}(q) = \sum_{k=\max(i,j)}^6 m_k l_{c,k}^2$. For the Coriolis term, instead of computing full Christoffel symbols, we use a velocity-coupling surrogate, $C_{ij}(q, \dot{q}) = k_c (|\dot{q}_i| + |\dot{q}_j|)$, $k_c = 10^{-3}$, which captures the dominant dependence on joint velocities. The gravitational torque is computed from link COM heights, $g_i(q) = m_i g \sum_{k=1}^i l_{c,k} \cos(\sum_{j=1}^k q_j)$, where the gravitational acceleration $g = 9.81 \text{m/s}^2$. Finally, the nominal discrete-time dynamics are obtained by RK4 integration with sampling time $\Delta t = 0.1 \text{s}$:

$$x_{k+1} = F(x_k, u_k) + w_k,$$

where $u_k = \tau(k\Delta t)$ is the zero-holder of the control input τ in the time interval $[k\Delta t, (k+1)\Delta t)$ and w_k models the additive process noise in the state space. For the conditional concic moment-based ambiguity set, we let $\gamma_0(z) = \min\{0.3\|z\|_2, 5\}\Delta t^2$, $\gamma_1 = 0.5\Delta t^2$, $\gamma_2 = 3$, $\bar{\mu}_k = \mathbf{0}_{12 \times 1}$, the nominal covariance for each pair $(q_i, \dot{q}_i)$ is $\begin{pmatrix} \Delta t^2 & 0.5\Delta t^2 \\ 0.5\Delta t^2 & 1.5\Delta t^2 \end{pmatrix}$, and the full nominal covariance $\bar{\Sigma}_k$ is the blockwise combination of these 6 sub-covariances. We simulate the nonlinear SDS under zero control by adversarially choosing one from the ambiguity set at each step against the predictor. The prediction performances of different probabilistic predictors are visualized and compared in Fig. 4. Similar to the results in Fig. 2, both Noise-DRPP and Eig-DRPP outperform the nominal predictor significantly in terms of expected average score and score concentration. This experiment demonstrates the effectiveness and practicality of the proposed DRPP framework on nonlinear SDSs.

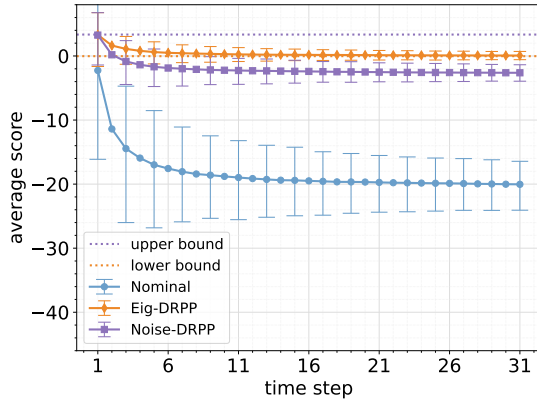


Fig. 4: Prediction performance of different probabilistic predictors on the nonlinear manipulator system.

IX. CONCLUSION

This paper presented a distributionally robust probabilistic prediction (DRPP) framework to optimize the worst-case prediction performance over a predefined ambiguity set of SDSs. To overcome the inherent intractability of optimizing over function spaces, Bellman equation is developed for the robust value function of DRPP, and optimality conditions are exploited to transform the original problem into Euclidean spaces. While achieving the global optimal solution remains computationally prohibitive, we developed two suboptimal

predictors: Noise-DRPP, obtained by relaxing the inner constraints, and Eig-DRPP, derived from relaxing the outer constraints. An explicit upper bound on the optimality gap was established for the proposed predictors. Numerical simulations demonstrated the effectiveness and practicality of the proposed predictors across various SDSs.

REFERENCES

- [1] F. S  vellec and S. S. Drijfhout, "A novel probabilistic forecast system predicting anomalously warm 2018-2022 reinforcing the long-term global warming trend," *Nature Communications*, vol. 9, no. 1, p. 3024, Aug. 2018.
- [2] E. Y. Cramer, E. L. Ray, V. K. Lopez, J. Bracher, A. Brennen, A. J. Castro Rivadeneira, A. Gerding, T. Gneiting, K. H. House, Y. Huang et al., "Evaluation of individual and ensemble probabilistic forecasts of covid-19 mortality in the united states," *Proceedings of the National Academy of Sciences*, vol. 119, no. 15, 2022.
- [3] R. Buizza, "The value of probabilistic prediction," *Atmospheric Science Letters*, vol. 9, no. 2, pp. 36–42, 2008.
- [4] J. Fisac, A. Bajcsy, S. Herbert, D. Fridovich-Keil, S. Wang, C. Tomlin, and A. Dragan, "Probabilistically safe robot planning with confidence-based human predictions," *Robotics: Science and Systems XIV*, 2018.
- [5] T. Gneiting and M. Katzfuss, "Probabilistic forecasting," *Annual Review of Statistics and Its Application*, vol. 1, no. 1, pp. 125–151, 2014.
- [6] T. Gneiting and A. E. Raftery, "Strictly proper scoring rules, prediction, and estimation," *Journal of the American Statistical Association*, vol. 102, no. 477, pp. 359–378, Mar. 2007.
- [7] D. Landgraf, A. V  lz, F. Berkel, K. Schmidt, T. Specker, and K. Graichen, "Probabilistic prediction methods for nonlinear systems with application to stochastic model predictive control," *Annual Reviews in Control*, vol. 56, Jan. 2023.
- [8] D. Hinrichsen and A. J. Pritchard, "Uncertain systems," in *Mathematical Systems Theory I: Modelling, State Space Analysis, Stability and Robustness*, D. Hinrichsen and A. J. Pritchard, Eds. Berlin, Heidelberg: Springer, 2005, pp. 517–713.
- [9] K. Schm  dgen, *The Moment Problem*, ser. Graduate Texts in Mathematics. Cham: Springer International Publishing, 2017, vol. 277.
- [10] E. Delage and Y. Ye, "Distributionally robust optimization under moment uncertainty with application to data-driven problems," *Operations Research*, Jan. 2010.
- [11] P. S. Maybeck, *Stochastic Models, Estimation, and Control*. Academic press, 1982.
- [12] M. Roth and F. Gustafsson, "An efficient implementation of the second order extended kalman filter," in *14th International Conference on Information Fusion*, Jul. 2011, pp. 1–6.
- [13] M. N  rgaard, N. K. Poulsen, and O. Ravn, "New developments in state estimation for nonlinear systems," *Automatica*, vol. 36, no. 11, pp. 1627–1638, Nov. 2000.
- [14] H. M. T. Menegaz, J. Y. Ishihara, G. A. Borges, and A. N. Vargas, "A systematization of the unscented kalman filter theory," *IEEE Transactions on Automatic Control*, vol. 60, no. 10, pp. 2583–2598, Oct. 2015.
- [15] M.   imandl and J. Dun  k, "Derivative-free estimation methods: New results and performance analysis," *Automatica*, vol. 45, no. 7, pp. 1749–1757, Jul. 2009.
- [16] H. W. Sorenson and D. L. Alspach, "Recursive bayesian estimation using gaussian sums," *Automatica*, vol. 7, no. 4, pp. 465–479, Jul. 1971.
- [17] D. Alspach and H. Sorenson, "Nonlinear bayesian estimation using gaussian sum approximations," *IEEE Transactions on Automatic Control*, vol. 17, no. 4, pp. 439–448, Aug. 1972.
- [18] R. Chen and J. S. Liu, "Mixture kalman filters," *Journal of the Royal Statistical Society Series B: Statistical Methodology*, vol. 62, no. 3, pp. 493–508, Sep. 2000.
- [19] N. Wiener, "The homogeneous chaos," *American Journal of Mathematics*, vol. 60, no. 4, pp. 897–936, 1938.
- [20] J. A. Paulson and A. Mesbah, "An efficient method for stochastic optimal control with joint chance constraints for nonlinear systems," *International Journal of Robust and Nonlinear Control*, vol. 29, no. 15, pp. 5017–5037, 2019.
- [21] S. S  rkk   and L. Svensson, *Bayesian Filtering and Smoothing*, 2nd ed., ser. Institute of Mathematical Statistics Textbooks. Cambridge: Cambridge University Press, 2023.
- [22] J. Zhang, "Modern monte carlo methods for efficient uncertainty quantification and propagation: A survey," *WIREs Computational Statistics*, vol. 13, no. 5, p. e1539, 2021.

- [23] M. Farina, L. Giulioni, and R. Scattolini, “Stochastic linear model predictive control with chance constraints – a review,” *Journal of Process Control*, vol. 44, pp. 53–67, Aug. 2016.
- [24] A. Mesbah, “Stochastic model predictive control: An overview and perspectives for future research,” *IEEE Control Systems Magazine*, vol. 36, no. 6, pp. 30–44, Dec. 2016.
- [25] R. D. McAllister and J. B. Rawlings, “Nonlinear stochastic model predictive control: Existence, measurability, and stochastic asymptotic stability,” *IEEE Transactions on Automatic Control*, vol. 68, no. 3, pp. 1524–1536, Mar. 2023.
- [26] J. Coulson, J. Lygeros, and F. Dörfler, “Distributionally robust chance constrained data-enabled predictive control,” *IEEE Transactions on Automatic Control*, vol. 67, no. 7, pp. 3289–3304, Jul. 2022.
- [27] P. Coppens and P. Patrinos, “Data-driven distributionally robust mpc for unconstrained stochastic systems,” *IEEE Control Systems Letters*, vol. 6, pp. 1274–1279.
- [28] C. P. Robert, *The Bayesian Choice*, ser. Springer Texts in Statistics. New York, NY: Springer, 2007.
- [29] D. T. Frazier, W. Maneesoonthorn, G. M. Martin, and B. P. M. McCabe, “Approximate bayesian forecasting,” *International Journal of Forecasting*, vol. 35, no. 2, pp. 521–539, Apr. 2019.
- [30] R. Koenker, “Quantile regression: 40 years on,” *Annual Review of Economics*, vol. 9, no. Volume 9, 2017, pp. 155–176, Aug. 2017.
- [31] L. Schlosser, T. Hothorn, R. Stauffer, and A. Zeileis, “Distributional regression forests for probabilistic precipitation forecasting in complex terrain,” *The Annals of Applied Statistics*, vol. 13, no. 3, pp. 1564–1589, Sep. 2019.
- [32] T. Duan, A. Anand, D. Y. Ding, K. K. Thai, S. Basu, A. Ng, and A. Schuler, “NGBoost: Natural gradient boosting for probabilistic prediction,” in *Proceedings of the 37th International Conference on Machine Learning*. PMLR, Nov. 2020, pp. 2690–2700.
- [33] D. Salinas, V. Flunkert, J. Gasthaus, and T. Januschowski, “DeepAR: Probabilistic forecasting with autoregressive recurrent networks,” *International Journal of Forecasting*, vol. 36, no. 3, pp. 1181–1191, Jul. 2020.
- [34] H. Tyralis and G. Papacharalampous, “A review of predictive uncertainty estimation with machine learning,” *Artificial Intelligence Review*, vol. 57, no. 4, p. 94, Mar. 2024.
- [35] H. E. Scarf, K. J. Arrow, and S. Karlin, *A Min-Max Solution of an Inventory Problem*. Rand Corporation Santa Monica, 1957.
- [36] G. Gallego and I. Moon, “The distribution free newsboy problem: Review and extensions,” *Journal of the Operational Research Society*, vol. 44, no. 8, pp. 825–834, Aug. 1993.
- [37] A. Shapiro and A. Pichler, “Conditional distributionally robust functionals,” *Operations Research*, vol. 72, no. 6, pp. 2745–2757.
- [38] M. Parry, A. P. Dawid, and S. Lauritzen, “Proper local scoring rules,” *The Annals of Statistics*, vol. 40, no. 1, Feb. 2012.
- [39] F. Lin, X. Fang, and Z. Gao, “Distributionally robust optimization: A review on theory and applications,” *Numerical Algebra, Control and Optimization*, vol. 12, no. 1, pp. 159–212, 2022.
- [40] G. N. Iyengar, “Robust dynamic programming,” *Mathematics of Operations Research*, vol. 30, no. 2, pp. 257–280, May 2005.
- [41] P. M. Pardalos and G. Schnitzler, “Checking local optimality in constrained quadratic programming is np-hard,” *Operations Research Letters*, vol. 7, no. 1, pp. 33–35, Feb. 1988.
- [42] C. Jin, P. Netrapalli, and M. Jordan, “What is local optimality in nonconvex-nonconcave minimax optimization?” in *Proceedings of the 37th International Conference on Machine Learning*. PMLR, Nov. 2020, pp. 4880–4889.
- [43] C. Daskalakis, S. Skoulakis, and M. Zampetakis, “The complexity of constrained min-max optimization,” in *Proceedings of the 53rd Annual ACM SIGACT Symposium on Theory of Computing*, ser. STOC 2021. New York, NY, USA: Association for Computing Machinery, Jun. 2021, pp. 1466–1478.
- [44] B. Grimmer, H. Lu, P. Worah, and V. Mirrokni, “Limiting behaviors of nonconvex-nonconcave minimax optimization via continuous-time systems,” in *Proceedings of The 33rd International Conference on Algorithmic Learning Theory*. PMLR, Mar. 2022, pp. 465–487.
- [45] T. Fiez, B. Chasnov, and L. Ratliff, “Implicit learning dynamics in stackelberg games: Equilibria characterization, convergence analysis, and empirical study,” in *Proceedings of the 37th International Conference on Machine Learning*. PMLR, Nov. 2020, pp. 3133–3144.
- [46] K. M. Lynch and F. C. Park, *Modern Robotics: Mechanics, Planning, and Control*, 1st ed. Cambridge: Cambridge University Press, Jul. 2017.

APPENDIX A PROOF OF THEOREM 1

Proof. For $k \in \{0, \dots, T-1\}$ and any $z \in \mathcal{Z}$, we denote $W_k(z) = \sup_{\hat{p}_k \in \mathcal{F}} \inf_{\rho_k \in \mathcal{I}_k(z)} \int_{\mathcal{X}} \mathcal{L}(\hat{p}, x) + V_{k+1}^*(z') d\rho_k(x)$. Step I: $V_k^*(z) \leq W_k(z)$. According to the definitions (4),(5), and (6), we have

$$\begin{aligned}
 V_k^*(z) &= \sup_{\mathcal{F}_{k:T-1}} \inf_{\mathcal{P}_{k:T-1}} \mathbb{E}_{\mathcal{P}_{k:T-1}} \left[\sum_{t=k}^{T-1} \mathcal{L}(\mathcal{F}_t(z_t), \mathbf{x}_{t+1}) \middle| z_k = z \right] \\
 &\stackrel{(i)}{=} \sup_{\mathcal{F}_{k:T-1}} \inf_{\mathcal{P}_{k:T-1}} \mathbb{E}_{\mathcal{P}_k} \left\{ \mathcal{L}(\mathcal{F}_k(z), \mathbf{x}_{k+1}) + \right. \\
 &\quad \left. \mathbb{E}_{\mathcal{P}_{k+1:T-1}} \left[\sum_{t=k}^{T-1} \mathcal{L}(\mathcal{F}_t(z_t), \mathbf{x}_{t+1}) \right] \middle| z_k = z \right\} \\
 &\stackrel{(ii)}{=} \sup_{\mathcal{F}_{k:T-1}} \inf_{\mathcal{P}_k} \mathbb{E}_{\mathcal{P}_k} \left\{ \mathcal{L}(\mathcal{F}_k(z), \mathbf{x}_{k+1}) + \right. \\
 &\quad \left. \inf_{\mathcal{P}_{k+1:T-1}} \mathbb{E}_{\mathcal{P}_{k+1:T-1}} \left[\sum_{t=k}^{T-1} \mathcal{L}(\mathcal{F}_t(z_t), \mathbf{x}_{t+1}) \right] \middle| z_k = z \right\} \\
 &= \sup_{\mathcal{F}_{k:T-1}} \inf_{\mathcal{P}_k} \mathbb{E}_{\mathcal{P}_k} \left[\mathcal{L}(\mathcal{F}_k(z), \mathbf{x}_{k+1}) + V_{k+1}^{\mathcal{F}}(z_{k+1}) \middle| z_k = z \right]
 \end{aligned}$$

where (i) holds because $\mathcal{P}_{k+1:T-1}$ do not affect the first term $\mathcal{L}(\mathcal{F}_k(z), \mathbf{x}_{k+1})$. According to Definition 1 and 2, it follows that for all $\mathcal{F} \in \mathfrak{F}$, $\mathcal{P} \in \mathfrak{P}$, $z \in \mathcal{Z}$, and $t \in \{k, \dots, T-1\}$, there is

$$\begin{aligned}
 &\mathbb{E}_{\mathcal{P}_t} [\mathcal{L}(\mathcal{F}_t(z_t), \mathbf{x}_{t+1}) \mid z_t = z] \\
 &\leq \mathbb{E}_{\mathcal{P}_t} [C(1 + \|\mathbf{x}_{t+1}\|_2^2) \mid z_t = z] < \infty.
 \end{aligned}$$

Therefore, (ii) interchanges $\inf_{\mathcal{P}_{k+1:T-1}}$ and $\mathbb{E}_{\mathcal{P}_k}$ according to the dominated convergence theorem.

Notice that $V_{k+1}^{\mathcal{F}}(z) \leq V_{k+1}^*(z)$, $\forall z \in \mathcal{Z}$, thus one has

$$\begin{aligned}
 V_k^*(z) &\leq \sup_{\mathcal{F}_{k:T-1}} \inf_{\mathcal{P}_k} \mathbb{E}_{\mathcal{P}_k} \left[\mathcal{L}(\mathcal{F}_k(z), \mathbf{x}_{k+1}) + V_{k+1}^*(z_{k+1}) \middle| z_k = z \right] \\
 &\stackrel{(i)}{=} \sup_{\hat{p}_k \in \mathcal{F}} \inf_{\rho_k \in \mathcal{I}_k(z)} \int_{\mathcal{X}} \mathcal{L}(\hat{p}, x) + V_{k+1}^*(x, \pi_{k+1}(x)) d\rho_k(x) \\
 &= W_k(z),
 \end{aligned}$$

where (i) holds because: first, $\mathcal{F}_k$ affect the first term by $\mathcal{F}_k(z) \in \mathcal{F}$ and does not affect the second term $V_{k+1}^*(z_{k+1})$; second, given $z_k = z$, one has the conditional measure $P_k(\cdot \mid z) \in \mathcal{I}_k(z)$ according to Assumption 3.

Step II: $V_k^*(z) \geq W_k(z)$. According to the definition $V_{k+1}^*(z) = \sup_{\mathcal{F} \in \mathfrak{F}} V_{k+1}^{\mathcal{F}}(z)$, it follows that for all $\epsilon > 0$, there exists a predictor $\mathcal{F}^\epsilon \in \mathfrak{F}$ such that $V_{k+1}^{\mathcal{F}^\epsilon}(z) \geq V_{k+1}^*(z) - \epsilon$ for all $z \in \mathcal{Z}$. Therefore,

$$\begin{aligned}
 V_k^*(z) &= \sup_{\mathcal{F}_{k:T-1}} \inf_{\mathcal{P}_k} \mathbb{E}_{\mathcal{P}_k} \left[\mathcal{L}(\mathcal{F}_k(z), \mathbf{x}_{k+1}) + V_{k+1}^{\mathcal{F}}(z_{k+1}) \middle| z_k = z \right] \\
 &\geq \sup_{\mathcal{F}_k} \inf_{\mathcal{P}_k} \mathbb{E}_{\mathcal{P}_k} \left[\mathcal{L}(\mathcal{F}_k(z), \mathbf{x}_{k+1}) + V_{k+1}^{\mathcal{F}^\epsilon}(z_{k+1}) \middle| z_k = z \right] \\
 &\geq \sup_{\mathcal{F}_k} \inf_{\mathcal{P}_k} \mathbb{E}_{\mathcal{P}_k} \left[\mathcal{L}(\mathcal{F}_k(z), \mathbf{x}_{k+1}) + V_{k+1}^*(z_{k+1}) \middle| z_k = z \right] \\
 &= W_k(z) - \epsilon.
 \end{aligned}$$

Since $\epsilon > 0$ is arbitrary, it implies that $V_k^*(z) \geq W_k(z)$. Combining the results from step 1 and step 2, we have $V_k^*(z) = W_k(z)$ for all $z \in \mathcal{Z}$. $\square$

APPENDIX B DUALITY OF PROBLEM (8)

Proof. The Lagrange function for (8) is defined as

$$\begin{aligned} L(\rho_k, \nu_k, \mu_k, r, q, Q_1, Q_2, P_1, p_1, s_1, P_2, p_2, s_2) \\ := \langle \rho_k, \log \hat{p}_k + V_{k+1}^*(x, \pi_{k+1}(x)) \rangle + r [\langle \rho_k, 1 \rangle - 1] \\ + q^\top [\langle \rho_k, x \rangle - \mu_k - \nu_k] \\ + Q_1 \cdot [\langle \rho_k, (x - \nu_k - \bar{\mu}_k)(x - \nu_k - \bar{\mu}_k)^\top \rangle - \gamma_2 \bar{\Sigma}_k] \\ - Q_2 \cdot [\langle \rho_k, (x - \nu_k - \bar{\mu}_k)(x - \nu_k - \bar{\mu}_k)^\top \rangle - \gamma_3 \bar{\Sigma}_k] \\ - P_1 \cdot \bar{\Sigma}_k - 2p_1^\top (\mu_k - \bar{\mu}_k) - s_1 \gamma_1 \\ - P_2 \cdot I - 2p_2^\top (\nu_k - \bar{\nu}_k) - s_2 \gamma_0(z). \end{aligned}$$

where the Lagrange multipliers satisfy that $Q_i \succeq 0$ and $\begin{bmatrix} P_i & p_i \\ p_i^\top & s_i \end{bmatrix} \succeq 0$ for $i \in \{1, 2\}$. Minimizing over $\rho_k \in \mathcal{M}_+(\mathcal{X})$, $\mu_k \in \mathbb{R}^{d_x}$, and $\nu_k \in \mathbb{R}^{d_x}$ we have

$$\begin{aligned} & \inf_{\rho_k, \mu_k, \nu_k} L(\rho_k, \nu_k, \mu_k, r, q, Q_1, Q_2, P_1, p_1, s_1, P_2, p_2, s_2) \\ &= \inf_{\rho_k, \mu_k, \nu_k} \langle \rho_k, \log \hat{p}_k(x) + r + q^\top x + x^\top (Q_1 - Q_2)x \\ & \quad + V_{k+1}^*(x, \pi_{k+1}(x)) \rangle - (q + 2p_1)^\top \mu_k \\ & \quad - (q + 2p_2)^\top \nu_k - (\gamma_2 Q_1 - \gamma_3 Q_2) \cdot \bar{\Sigma}_k \\ & \quad - (\nu_k + 2\mu_k - \bar{\mu}_k)^\top (Q_1 - Q_2)(\nu_k + \bar{\mu}_k) - r \\ & \quad - P_1 \cdot \bar{\Sigma}_k + 2p_1^\top \bar{\mu}_k - s_1 \gamma_1 - P_2 \cdot I + 2p_2^\top \bar{\mu}_k - s_2 \gamma_0(z) \\ & \stackrel{(i)}{=} \inf_{\mu_k, \nu_k} -(q + 2p_1)^\top \mu_k - (q + 2p_2)^\top \nu_k - (\gamma_2 Q_1 - \gamma_3 Q_2) \cdot \bar{\Sigma}_k \\ & \quad - (\nu_k + 2\mu_k - \bar{\mu}_k)^\top (Q_1 - Q_2)(\nu_k + \bar{\mu}_k) - r \\ & \quad - P_1 \cdot \bar{\Sigma}_k + 2p_1^\top \bar{\mu}_k - s_1 \gamma_1 - P_2 \cdot I + 2p_2^\top \bar{\nu}_k - s_2 \gamma_0(z) \\ & \stackrel{(ii)}{=} \inf_{\nu_k} -r + \bar{\mu}_k^\top (Q_1 - Q_2) \bar{\mu}_k - (\gamma_2 Q_1 - \gamma_3 Q_2 + P_1) \cdot \bar{\Sigma}_k - P_2 \cdot I \\ & \quad + 2p_1^\top \bar{\mu}_k - s_1 \gamma_1 + 2p_2^\top \bar{\mu}_k - s_2 \gamma_0(z) - (q + 2p_2)^\top \nu_k \\ & \quad - \nu_k^\top (Q_1 - Q_2) \nu_k, \end{aligned}$$

where $q + 2p_1 + 2(Q_1 - Q_2)(\nu_k + \bar{\mu}_k) = 0$. Equation (i) holds when $\log \hat{p}_k(x) + r + q^\top x + x^\top (Q_1 - Q_2)x + V_{k+1}^*(x, \pi_{k+1}(x)) \geq 0 \ \forall x \in \mathbb{R}^{d_x}$, otherwise there is $\inf_{\rho_k \in \mathcal{M}_+(\mathcal{X})} L$ equals $-\infty$. Equation (ii) holds because

$$\begin{aligned} & \inf_{\mu_k, \nu_k} -(q + 2p_1)^\top \mu_k - (q + 2p_2)^\top \nu_k - \bar{\mu}_k^\top (Q_1 - Q_2) \bar{\mu}_k \\ & \quad - (\nu_k + 2\mu_k - \bar{\mu}_k)^\top (Q_1 - Q_2)(\nu_k + \bar{\mu}_k) \\ &= \inf_{\nu_k} \inf_{\mu_k} -[q + 2p_1 + 2(Q_1 - Q_2)(\nu_k + \bar{\mu}_k)]^\top \mu_k \\ & \quad - (q + 2p_2)^\top \nu_k - \nu_k^\top (Q_1 - Q_2) \nu_k \\ &= \inf_{\nu_k} -(q + 2p_2)^\top \nu_k - \nu_k^\top (Q_1 - Q_2) \nu_k \\ & \text{s.t. } q + 2p_1 + 2(Q_1 - Q_2)(\nu_k + \bar{\mu}_k) = 0. \end{aligned}$$

Combining those conditions and the previous multipliers' constraints, we have the dual form completed. $\square$

APPENDIX C PROOF OF THEOREM 2

Proof. Step I. We leverage the separable structure of $\mathbf{D}_1$ to reformulate the problem. For ease of notation, we use the κ to summarize all the Lagrange multipliers except for r , i.e.,

$$\kappa = (q, Q_1, Q_2, P_1, p_1, s_1, P_2, p_2, s_2).$$

First, the objective function can be separated as $G(r, \kappa) = -r + g(\kappa)$ where $g(\kappa) = \bar{\mu}_k^\top (Q_1 - Q_2) \bar{\mu}_k - (\gamma_2 Q_1 - \gamma_3 Q_2 + P_1) \cdot \bar{\Sigma}_k - P_2 \cdot I + 2p_1^\top \bar{\mu}_k - s_1 \gamma_1 + 2p_2^\top \bar{\nu}_k - s_2 \gamma_0(z) - (q + 2p_2)^\top \nu_k^* - \nu_k^{*\top} (Q_1 - Q_2) \nu_k^*$. Second, only one of the constraints concerns r . Therefore, given any group of feasible κ satisfying the other constraints in $\mathbf{D}_1$, the optimizing over $\hat{p}_k$ and r is equivalent to solving the following problem:

$$\begin{aligned} & \sup_{\hat{p}_k \in \mathcal{F}, r} -r + g(\kappa) \\ & \text{s.t. } \begin{cases} x^\top (Q_1 - Q_2)x + x^\top q + r + \log \hat{p}_k(x) \\ + V_{k+1}^*(x, \pi_{k+1}(x)) \geq 0 \ \forall x \in \mathcal{X} \end{cases} \end{aligned} \quad (20)$$

Step II. We exploit the last constraint in (20) to analyze the lower bound of r . Notice that

$$\begin{aligned} 1 &= \int_{\mathbb{R}^{d_x}} \hat{p}_k(x) dx \\ &\geq \int_{\mathbb{R}^{d_x}} \exp \{ -x^\top (Q_1 - Q_2)x - x^\top q - r \\ & \quad - V_{k+1}^*(x, \pi_{k+1}(x)) \} dx, \end{aligned} \quad (21)$$

one immediately gets that r is lower bounded by $r^* := \log[\int_{\mathbb{R}^{d_x}} \exp \{ -x^\top (Q_1 - Q_2)x - x^\top q - V_{k+1}^*(x, \pi_{k+1}(x)) \} dx]$.

Notice that the equality of (21) holds only when there is $\hat{p}_k(x) = \exp \{ -x^\top (Q_1 - Q_2)x - x^\top q - r^* \}$ for $x \in \mathbb{R}^{d_x}$ almost everywhere. Therefore, given any feasible κ , the objective function $-r + g(\kappa)$ in problem (20) can only attain its maximum when $\hat{p}_k$ belongs to the following exponential family almost everywhere on $\mathbb{R}^{d_x}$,

$$\hat{p}_k(x) \stackrel{a.s.}{\propto} \exp \{ -x^\top \theta_1 x - x^\top \theta_2 - V_{k+1}^*(x, \pi_{k+1}(x)) \},$$

where $\theta_1 \in S^{d_x}$ and $\theta_2 \in \mathbb{R}^{d_x}$.

Step III. Finally, when $k = T - 1$, there is $V_T^*(x, \pi_T(x)) = 0$. Because Q_1 and Q_2 are both semi-definite, we have $Q_1 - Q_2$ is symmetric and there exists an orthogonal matrix $U \in \mathbb{R}^{d_x \times d_x}$ and an diagonal matrix $D = \text{diag}(\lambda_1, \dots, \lambda_{d_x})$ such that $Q_1 - Q_2 = UDU^\top$. Then we have

$$\begin{aligned} & \int_{\mathbb{R}^{d_x}} \exp \{ -x^\top (Q_1 - Q_2)x - x^\top q - r \} dx \\ &= \int_{\mathbb{R}^{d_x}} \exp \{ -x^\top UDU^\top x - x^\top q - r \} dx \\ &= \int_{\mathbb{R}^{d_x}} \exp \{ -\tilde{x}^\top D \tilde{x} - \tilde{x}^\top \tilde{q} - r \} \det(U) d\tilde{x} \\ &= e^{-r} \det(U) \prod_{i=1}^{d_x} \int_{\mathbb{R}} \exp \{ -\tilde{x}_{(i)}^2 \lambda_i - \tilde{x}_{(i)} \tilde{q}_{(i)} \} d\tilde{x}_{(i)}, \end{aligned}$$

where the second equation follows by letting $\tilde{x} = U^\top x$ and $\tilde{q} = U^\top q$. If there is an index $j \in \{1, \dots, d_x\}$ such that $\lambda_j < 0$, then $\int_{\mathbb{R}} \exp \{ -\tilde{x}_{(j)}^2 \lambda_j - \tilde{x}_{(j)} \tilde{q}_{(j)} \} d\tilde{x}_{(j)} = \infty$, which contradicts (21). Therefore, $Q_1 - Q_2$ is semi-definite positive,

and a necessary optimal condition for the inner optimization is that $\hat{p}_{T-1}$ is subject to Gaussian almost everywhere. $\square$

APPENDIX D DERIVATION OF (11)

Proof. Facilitated by Theorem 2, when $k = T-1$ we can first parameterize $\hat{p}_k$ by $\hat{\mu}_k \in \mathbb{R}^{d_x}$ and $\hat{\Sigma}_k \in S_+^{d_x}$ such that

$$\hat{p}_k(x) = \frac{1}{\sqrt{(2\pi)^{d_x} \det(\hat{\Sigma}_k)}} \exp \left[-\frac{1}{2} (x - \hat{\mu}_k)^\top \hat{\Sigma}_k^{-1} (x - \hat{\mu}_k) \right].$$

Substituting it into the constraints of $\mathbf{D}_1$, there is

$$\begin{aligned} x^\top (Q_1 - Q_2)x + x^\top q + r - \frac{1}{2} \left[d_x \log(2\pi) + \log \det(\hat{\Sigma}_k) \right] \\ - \frac{1}{2} (x - \hat{\mu}_k)^\top \hat{\Sigma}_k^{-1} (x - \hat{\mu}_k) = 0 \quad \forall x \in \mathbb{R}^{d_x}. \end{aligned}$$

Then we have $Q_1 - Q_2 = \frac{1}{2} \hat{\Sigma}_k^{-1}$, $q = -\hat{\Sigma}_k^{-1} \hat{\mu}_k$, and $r = \frac{1}{2} \left[d_x \log(2\pi) + \log \det(\hat{\Sigma}_k) + \hat{\mu}_k^\top \hat{\Sigma}_k^{-1} \hat{\mu}_k \right]$. Substituting these equations and the parameterized $\hat{p}_k$ into the objective of problem $\mathbf{D}_1$, we get

$$\begin{aligned} G(r, q, Q_1, Q_2, P_1, p_1, s_1, P_2, p_2, s_2) \\ = -\frac{1}{2} \left[d_x \log(2\pi) + \log \det(\hat{\Sigma}_k) \right] \\ - P_2 \cdot I - s_1 \gamma_1 + 2p_2^\top (\bar{\nu}_k + \bar{\mu}_k - \hat{\mu}_k) - s_2 \gamma_0(z) \\ + 2(2p_2 - p_1)^\top \hat{\Sigma}_k p_1 - (\gamma_2 Q_1 - \gamma_3 Q_2 + P_1) \cdot \bar{\Sigma}_k \\ = -\frac{1}{2} d_x \log(2\pi) + \frac{1}{2} \log \det(\hat{\Sigma}_k^{-1}) - (\gamma_2 Q_1 - \gamma_3 Q_2) \cdot \bar{\Sigma}_k \\ - P_1 \cdot \bar{\Sigma}_k - P_2 \cdot I - s_1 \gamma_1 - s_2 \gamma_0(z) \\ - 2p_2^\top (\hat{\mu}_k - \bar{\mu}_k - \bar{\nu}_k) + 2p_1^\top \hat{\Sigma}_k (2p_2 - p_1). \end{aligned}$$

Substituting the above equations into the problem $\mathbf{D}_1$, we have the problem (11) derived. $\square$

APPENDIX E PROOF OF LEMMA 1

Proof. Since the objective function is linear with respect to Σ_k and the second constraint provides an upper bound for it, one immediately has the worst-case covariance should satisfy that

$$\Sigma_k^* = \gamma_2 \bar{\Sigma}_k - (\mu_k - \bar{\mu}_k)(\mu_k - \bar{\mu}_k)^\top.$$

Let $a = \mu_k - \bar{\mu}_k$, $b = \nu_k - \bar{\nu}_k$, $c = \bar{\mu}_k + \bar{\nu}_k - \hat{\mu}_k$, problem $\mathbf{P}_2$ can be further equivalent to

$$\begin{aligned} \inf_{\hat{\Sigma}_k^{-1}, c} \sup_{a, b} & -\log \det(\hat{\Sigma}_k^{-1}) + \gamma_2 \bar{\Sigma}_k \cdot \hat{\Sigma}_k^{-1} \\ & + \|a + b + c\|_{\hat{\Sigma}_k^{-1}}^2 - \|a\|_{\hat{\Sigma}_k^{-1}}^2 \\ \text{s.t. } & \|a\|_{\hat{\Sigma}_k^{-1}}^2 \leq \gamma_1, \|b\|_2^2 \leq \gamma_0(z). \end{aligned} \quad (22)$$

Suppose $(a^*, b^*) \in \arg \max_{a, b} \left\{ \|a + b\|_{\hat{\Sigma}_k^{-1}}^2 - \|a\|_{\hat{\Sigma}_k^{-1}}^2 \right\}$, we have $(-a^*, -b^*)$ can also attain the maximum. Then we have

$$\begin{aligned} & \max_{a, b} \|a + b + c\|_{\hat{\Sigma}_k^{-1}}^2 - \|a\|_{\hat{\Sigma}_k^{-1}}^2 \\ & \geq \max \left\{ \|a^* + b^* + c\|_{\hat{\Sigma}_k^{-1}}^2 - \|a^*\|_{\hat{\Sigma}_k^{-1}}^2, \right. \\ & \quad \left. \|-a^* - b^* + c\|_{\hat{\Sigma}_k^{-1}}^2 - \|-a^*\|_{\hat{\Sigma}_k^{-1}}^2 \right\} \\ & = \|a^* + b^*\|_{\hat{\Sigma}_k^{-1}}^2 - \|a^*\|_{\hat{\Sigma}_k^{-1}}^2 + \|c\|_{\hat{\Sigma}_k^{-1}}^{-1} \\ & \quad + 2 \max \left\{ \langle c, a^* + b^* \rangle_{\hat{\Sigma}_k^{-1}}, -\langle c, a^* + b^* \rangle_{\hat{\Sigma}_k^{-1}} \right\} \\ & \geq \|a^* + b^*\|_{\hat{\Sigma}_k^{-1}}^2 - \|a^*\|_{\hat{\Sigma}_k^{-1}}^2 \end{aligned}$$

where the last inequality is strict when $c \neq 0$. Therefore, we have $c^* = \arg \min_c \max_{a, b} \|a + b + c\|_{\hat{\Sigma}_k^{-1}}^2 - \|a\|_{\hat{\Sigma}_k^{-1}}^2 = 0$, which means that $\hat{\mu}_k^* = \bar{\mu}_k + \bar{\nu}_k$. $\square$

APPENDIX F DUALITY OF THE MOMENT PROBLEM (15)

Proof. The Lagrange function for (15) is defined as

$$\begin{aligned} L(\tilde{\rho}_k, \mu_k, r, q, Q_1, Q_2, P, p, s) \\ := \langle \tilde{\rho}_k, \log \hat{p}_k(w + \bar{\nu}_k) \rangle + r [\langle \tilde{\rho}_k, 1 \rangle - 1] + q^\top [\langle \tilde{\rho}_k, w \rangle - \mu_k] \\ + Q_1 \cdot [\langle \tilde{\rho}_k, (w - \bar{\mu}_k)(w - \bar{\mu}_k)^\top \rangle - \gamma_2 \bar{\Sigma}_k] \\ - Q_2 \cdot [\langle \tilde{\rho}_k, (w - \bar{\mu}_k)(w - \bar{\mu}_k)^\top \rangle - \gamma_3 \bar{\Sigma}_k] \\ - P \cdot \bar{\Sigma}_k - 2p^\top (\mu_k - \bar{\mu}_k) - s \gamma_1, \end{aligned}$$

where the Lagrange multipliers satisfy $Q_1 \succeq 0, Q_2 \succeq 0$ and $\begin{bmatrix} P & p \\ p^\top & s \end{bmatrix} \succeq 0$. Minimizing over $\tilde{\rho}_k \in \mathcal{M}_+(\mathcal{X})$ and $\mu_k \in \mathbb{R}^{d_x}$, one has

$$\begin{aligned} & \inf_{\tilde{\rho}_k \in \mathcal{M}_+(\mathcal{X}), \mu_k} L(\tilde{\rho}_k, \mu_k, r, q, Q_1, Q_2, P, p, s) \\ & = \inf_{\tilde{\rho}_k \in \mathcal{M}_+(\mathcal{X}), \mu_k} \langle \tilde{\rho}_k, \log \hat{p}_k(w + \bar{\nu}_k) + r + q^\top w + w^\top (Q_1 - Q_2) w \rangle \\ & \quad - [q + 2(Q_1 - Q_2) \bar{\mu}_k + 2p]^\top \mu_k - r + \bar{u}_k^\top (Q_1 - Q_2) \bar{u}_k \\ & \quad - (\gamma_2 Q_1 - \gamma_3 Q_2 + P) \cdot \bar{\Sigma}_k + 2p^\top \bar{\mu}_k - s \gamma_1 \\ & = -r + \bar{u}_k^\top (Q_1 - Q_2) \bar{u}_k - (\gamma_2 Q_1 - \gamma_3 Q_2 + P) \cdot \bar{\Sigma}_k + 2p^\top \bar{\mu}_k \\ & \quad - s \gamma_1, \end{aligned}$$

where the second equality holds when $\log \hat{p}_k(w + \bar{\nu}_k) + r + q^\top w + w^\top (Q_1 - Q_2) w \geq 0 \quad \forall w \in \mathbb{R}^{d_x}$ and $q + 2(Q_1 - Q_2) \bar{\mu}_k + 2p = 0$ hold simultaneously, otherwise there is $\inf_{\tilde{\rho}_k \in \mathcal{M}_+(\mathcal{X}), \mu_k} L = -\infty$. Combining these two conditions and the previous multipliers' constraints, we have attained the dual problem $\mathbf{D}_2$. $\square$

APPENDIX G PROOF OF THEOREM 3

Proof. The proof is organized in three parts: first, we use Theorem 2 to parameterize $\hat{p}_k$ by Gaussian distribution $\mathcal{N}(\hat{\mu}_k, \hat{\Sigma}_k)$ and reformulate the dual problem $\mathbf{D}_2$ to a finite-dimensional optimization; second, we solve the optimal predictive pdf $\hat{p}_k^* \sim \mathcal{N}(\hat{\mu}_k^*, \hat{\Sigma}_k^*)$; third, we substitute $\hat{p}_k^*$ into the original problem $\mathbf{P}_1$ to obtain the worst-case conditional measure ρ_k^* by solving the inner minimization.

Step I. Using the necessary optimality condition for $\hat{p}_k$, we can parameterize $\hat{p}_k$ by $\hat{\mu}_k \in \mathbb{R}^{d_x}$ and $\hat{\Sigma}_k \in S_+^{d_x}$ such that

$$\hat{p}_k(w + \bar{\nu}_k) = \frac{\exp\left[-\frac{1}{2}(w - \hat{\mu}_k)^\top \hat{\Sigma}_k^{-1}(w - \hat{\mu}_k)\right]}{\sqrt{(2\pi)^{d_x} \det(\hat{\Sigma}_k)}}.$$

Moreover, from the constraint $w^\top(Q_1 - Q_2)w + w^\top q + r + \log \hat{p}_k(w + \bar{\nu}_k) \geq 0 \ \forall w \in \mathbb{R}^{d_x}$, we get

$$w^\top(Q_1 - Q_2)w + w^\top q + r - \frac{1}{2} \left[d_x \log(2\pi) + \log \det(\hat{\Sigma}_k) \right] - \frac{1}{2}(w - \hat{\mu}_k)^\top \hat{\Sigma}_k^{-1}(w - \hat{\mu}_k) = 0 \quad \forall w \in \mathbb{R}^{d_x},$$

which indicates that $Q_1 - Q_2 = \frac{1}{2}\hat{\Sigma}_k^{-1}$, $q = -\hat{\Sigma}_k^{-1}\hat{\mu}_k$, and $r = \frac{1}{2} \left[d_x \log(2\pi) + \log \det(\hat{\Sigma}_k) + \hat{\mu}_k^\top \hat{\Sigma}_k^{-1} \hat{\mu}_k \right]$. Next, from the constraint $q + 2(Q_1 - Q_2)\bar{\mu}_k + 2p = 0$ one can get $p = \frac{1}{2}\hat{\Sigma}_k^{-1}(\hat{\mu}_k - \bar{\mu}_k)$. Substituting the parameterized $\hat{p}_k$ and the above equations into problem **D**₂, there is

$$\begin{aligned} G(r, q, Q_1, Q_2, P, p, s) &= -\frac{1}{2} \left[d_x \log(2\pi) + \log \det(\hat{\Sigma}_k) \right] \\ &\quad - \frac{1}{2} \hat{\mu}_k^\top \hat{\Sigma}_k^{-1} \hat{\mu}_k + \bar{\mu}_k^\top (Q_1 - Q_2) \bar{\mu}_k + \bar{\mu}_k^\top \hat{\Sigma}_k^{-1} (\hat{\mu}_k - \bar{\mu}_k) \\ &\quad - (\gamma_2 Q_1 - \gamma_3 Q_2 + P) \cdot \bar{\Sigma}_k - s\gamma_1 \\ &= -\frac{1}{2} \left[d_x \log(2\pi) - \log \det(\hat{\Sigma}_k^{-1}) + \|\hat{\mu}_k - \bar{\mu}_k\|_{\hat{\Sigma}_k^{-1}}^2 \right] \\ &\quad - (\gamma_2 Q_1 - \gamma_3 Q_2 + P) \cdot \bar{\Sigma}_k - s\gamma_1. \end{aligned}$$

As a result, **D**₂ can be transformed into a finite-dimensional optimization problem as follows:

$$\begin{aligned} \sup_{\hat{\mu}_k, \hat{\Sigma}_k, Q_1, Q_2, P, p, s} & -\frac{1}{2} d_x \log(2\pi) + \frac{1}{2} \log \det(\hat{\Sigma}_k^{-1}) - s\gamma_1 \\ & - \frac{1}{2} \|\hat{\mu}_k - \bar{\mu}_k\|_{\hat{\Sigma}_k^{-1}}^2 - (\gamma_2 Q_1 - \gamma_3 Q_2 + P) \cdot \bar{\Sigma}_k \\ \text{s.t.} & \begin{cases} Q_1 \succeq 0, Q_2 \succeq 0, Q_1 - Q_2 = \frac{1}{2}\hat{\Sigma}_k^{-1}, \hat{\Sigma}_k \succ 0 \\ \begin{bmatrix} P & p \\ p^\top & s \end{bmatrix} \succeq 0, p = \frac{1}{2}\hat{\Sigma}_k^{-1}(\hat{\mu}_k - \bar{\mu}_k). \end{cases} \end{aligned} \quad (23)$$

Step II. Notice that the constraint $\hat{\Sigma}_k \succ 0$ and the relationship $p = \frac{1}{2}\hat{\Sigma}_k^{-1}(\hat{\mu}_k - \bar{\mu}_k)$ reveals that p and $\hat{\mu}_k$ are one-on-one, one can optimize on either of them. In this proof, we choose to maximize over p and substitute $\hat{\mu}_k$ in the objective by p . Next, notice that the constraints for Q_1, Q_2 only depend on $\hat{\Sigma}_k$, and it can be separated from the constraints for P, p, s , we can first maximize over Q_1 and Q_2 , then over P, p, s , finally over $\hat{\Sigma}_k$. In other words, the objective of problem (23) can be reformulated as

$$\begin{aligned} \sup_{\hat{\Sigma}_k} \sup_{P, p, s} \sup_{Q_1, Q_2} & -\frac{1}{2} d_x \log(2\pi) + \frac{1}{2} \log \det(\hat{\Sigma}_k^{-1}) \\ & - 2p^\top \hat{\Sigma}_k p - (\gamma_2 Q_1 - \gamma_3 Q_2 + P) \cdot \bar{\Sigma}_k - s\gamma_1. \end{aligned} \quad (24)$$

The solution of the first maximization is $Q_1^* = \frac{1}{2}\hat{\Sigma}_k^{-1}$, $Q_2^* = 0$ because

$$\begin{aligned} Q_1^* &= \arg \max_{Q_1 \succeq \frac{1}{2}\hat{\Sigma}_k^{-1}} -((\gamma_2 - \gamma_3)Q_1 + \frac{\gamma_3}{2}\hat{\Sigma}_k^{-1}) \cdot \bar{\Sigma}_k \\ &= \arg \max_{Q_1 \succeq \frac{1}{2}\hat{\Sigma}_k^{-1}} -(\gamma_2 - \gamma_3)Q_1 \cdot \bar{\Sigma}_k = \frac{1}{2}\hat{\Sigma}_k^{-1}, \end{aligned}$$

and $Q_2^* = Q_1^* - \frac{1}{2}\hat{\Sigma}_k^{-1} = 0$.

The solution of the second maximization is $P^* = 0, p^* = 0, s^* = 0$. Notice that $\begin{bmatrix} P & p \\ p^\top & s \end{bmatrix} \succeq 0$ is equivalent to $s \geq p^\top P^{-1}p$, therefore $s \geq 0$ and $\arg \max_s -s\gamma_1 = 0$. Again, $\arg \max_p -2p^\top \hat{\Sigma}_k p = 0$ and $\arg \max_P P \cdot \bar{\Sigma}_k = 0$. Additionally, $p^* = 0$ indicates that $\hat{\mu}_k^* = 2\hat{\Sigma}_k p + \bar{\mu}_k = \bar{\mu}_k$.

The solution of the third maximization is $\hat{\Sigma}_k^* = \gamma_2 \bar{\Sigma}_k$. This is because $\hat{\Sigma}_k^* = [\arg \max_{\hat{\Sigma}_k^{-1} \succ 0} F(\hat{\Sigma}_k^{-1})]^{-1}$ where $F(X) = \log \det(X) - X \cdot \gamma_2 \bar{\Sigma}_k$. Notice that $\frac{d}{dX} F(X) = X^{-1} - \gamma_2 \bar{\Sigma}_k$ and $\frac{d}{dX} F$ monotonously decrease on X . Therefore, F is concave, which can be maximized at X^* where $\frac{d}{dX} F(X^*) = 0$, i.e., $X^* = (\gamma_2 \bar{\Sigma}_k)^{-1}$. Finally, we have $\hat{\Sigma}_k^*$ equals the inverse of X^* , which is $\gamma_2 \bar{\Sigma}_k$. Therefore, the optimal predictive pdf is $\hat{p}_k^* \sim \mathcal{N}(\bar{\nu}_k + \bar{\mu}_k, \gamma_2 \bar{\Sigma}_k)$.

Step III. Substituting the optimal $\hat{p}_k^*$ into problem **P**₁, we have the inner minimization problem as

$$\begin{aligned} \inf_{\bar{p}_k, \mu_k} & \mathbb{E}_{w \sim \bar{p}_k} -\frac{1}{2} \left[d_x \log(2\pi) + \log \det(\gamma_2 \bar{\Sigma}_k) \right. \\ & \quad \left. + (w - \bar{\mu}_k)^\top (\gamma_2 \bar{\Sigma}_k)^{-1} (w - \bar{\mu}_k) \right] \\ \text{s.t.} & \begin{cases} \mathbb{E}_{w \sim \bar{p}_k}[1] = 1, \mathbb{E}_{w \sim \bar{p}_k}[w] = \mu_k \\ \gamma_3 \bar{\Sigma}_k \preceq \mathbb{E}_{w \sim \bar{p}_k} \left[(w - \bar{\mu}_k)(w - \bar{\mu}_k)^\top \right] \preceq \gamma_2 \bar{\Sigma}_k \\ \begin{bmatrix} \bar{\Sigma}_k & (\mu_k - \bar{\mu}_k) \\ (\mu_k - \bar{\mu}_k)^\top & \gamma_1 \end{bmatrix} \succeq 0. \end{cases} \end{aligned} \quad (25)$$

The objective is minimized when $\mathbb{E}[(w_k - \bar{\mu}_k)(w_k - \bar{\mu}_k)^\top]$ equals $\gamma_2 \bar{\Sigma}_k$, and the proof is completed. $\square$



Tao Xu (S'22) received a B.S. degree in the School of Mathematical Sciences from Shanghai Jiao Tong University (SJTU), Shanghai, China. He is currently working toward the Ph.D. degree with the Department of Automation, SJTU. He is a member of Intelligent of Wireless Networking and Cooperative Control Group. His research interests include probabilistic prediction, distributionally robust optimization, dynamic games, and robotics.



Jianping He (SM'19) is currently a full professor in the School of Automation and Intelligent Sensing at Shanghai Jiao Tong University, Shanghai, China. He received the Ph.D. degree in control science and engineering from Zhejiang University, Hangzhou, China, in 2013, and had been a research fellow in the Department of Electrical and Computer Engineering at University of Victoria, Canada, from Dec. 2013 to Mar. 2017. His research interests mainly include the distributed learning, control and optimization, security and privacy in network systems.