

From data to decisions

Distributionally robust optimization is optimal

van Parys, Bart P.G.; Mohajerin Esfahani, P.; Kuhn, Daniel

DOI

[10.1287/mnsc.2020.3678](https://doi.org/10.1287/mnsc.2020.3678)

Publication date

2021

Document Version

Final published version

Published in

Management Science

Citation (APA)

van Parys, B. P. G., Mohajerin Esfahani, P., & Kuhn, D. (2021). From data to decisions: Distributionally robust optimization is optimal. *Management Science*, 67(6), 3387-3402.
<https://doi.org/10.1287/mnsc.2020.3678>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.



Management Science

Publication details, including instructions for authors and subscription information:
<http://pubsonline.informs.org>

From Data to Decisions: Distributionally Robust Optimization Is Optimal

Bart P. G. Van Parys, Peyman Mohajerin Esfahani, Daniel Kuhn

To cite this article:

Bart P. G. Van Parys, Peyman Mohajerin Esfahani, Daniel Kuhn (2021) From Data to Decisions: Distributionally Robust Optimization Is Optimal. Management Science 67(6):3387-3402. <https://doi.org/10.1287/mnsc.2020.3678>

Full terms and conditions of use: <https://pubsonline.informs.org/Publications/Librarians-Portal/PubsOnLine-Terms-and-Conditions>

This article may be used only for the purposes of research, teaching, and/or private study. Commercial use or systematic downloading (by robots or other automatic processes) is prohibited without explicit Publisher approval, unless otherwise noted. For more information, contact permissions@informs.org.

The Publisher does not warrant or guarantee the article's accuracy, completeness, merchantability, fitness for a particular purpose, or non-infringement. Descriptions of, or references to, products or publications, or inclusion of an advertisement in this article, neither constitutes nor implies a guarantee, endorsement, or support of claims made of that product, publication, or service.

Copyright © 2020, INFORMS

Please scroll down for article—it is on subsequent pages



With 12,500 members from nearly 90 countries, INFORMS is the largest international association of operations research (O.R.) and analytics professionals and students. INFORMS provides unique networking and learning opportunities for individual professionals, and organizations of all types and sizes, to better understand and use O.R. and analytics tools and methods to transform strategic visions and achieve better outcomes.

For more information on INFORMS, its publications, membership, or meetings visit <http://www.informs.org>

From Data to Decisions: Distributionally Robust Optimization Is Optimal

Bart P. G. Van Parys,^a Peyman Mohajerin Esfahani,^b Daniel Kuhn^c

^aOperations Research Center, Massachusetts Institute of Technology, Cambridge, Massachusetts 02139; ^bDelft Center for Systems and Control, Technische Universiteit Delft, 2628 CD Delft, Netherlands; ^cRisk Analytics and Optimization Chair, Ecole Polytechnique Fédérale de Lausanne, CH-1015 Lausanne, Switzerland

Contact: vanparys@mit.edu,  <https://orcid.org/0000-0003-4177-4849> (BPGVP); p.mohajerinesfahani@tudelft.nl,  <https://orcid.org/0000-0003-1286-8782> (PME); daniel.kuhn@epfl.ch,  <https://orcid.org/0000-0003-2697-8886> (DK)

Received: April 13, 2017

Revised: June 13, 2019; December 22, 2019

Accepted: April 17, 2020

Published Online in Articles in Advance:
November 23, 2020

<https://doi.org/10.1287/mnsc.2020.3678>

Copyright: © 2020 INFORMS

Abstract. We study stochastic programs where the decision maker cannot observe the distribution of the exogenous uncertainties but has access to a finite set of independent samples from this distribution. In this setting, the goal is to find a procedure that transforms the data to an estimate of the expected cost function under the unknown data-generating distribution, that is, a *predictor*, and an optimizer of the estimated cost function that serves as a near-optimal candidate decision, that is, a *prescriptor*. As functions of the data, predictors and prescriptors constitute statistical estimators. We propose a meta-optimization problem to find the least conservative predictors and prescriptors subject to constraints on their out-of-sample disappointment. The out-of-sample disappointment quantifies the probability that the actual expected cost of the candidate decision under the unknown true distribution exceeds its predicted cost. Leveraging tools from large deviations theory, we prove that this meta-optimization problem admits a unique solution: The best predictor-prescriptor-pair is obtained by solving a distributionally robust optimization problem over all distributions within a given relative entropy distance from the empirical distribution of the data.

History: Accepted by Chung Piaw Teo, optimization.

Funding: Funding from the Swiss National Science Foundation for the first author [Early Postdoc Mobility Grant 165226] and the third author [Grant BSCGIO_157733] is gratefully acknowledged.

Supplemental Material: The online appendix is available at <https://doi.org/10.1287/mnsc.2020.3678>

Keywords: data-driven optimization • distributionally robust optimization • large deviations theory • relative entropy • convex optimization

1. Introduction

We study static decision problems under uncertainty, where the decision maker cannot observe the probability distribution of the uncertain problem parameters but has access to a finite number of independent samples from this distribution. Classical stochastic programming uses this data only indirectly. The data serves as the input for a statistical estimation problem that aims to infer the distribution of the uncertain problem parameters. The estimated distribution then serves as an input for an optimization problem that outputs a near-optimal decision as well as an estimate of the expected cost incurred by this decision. Thus, classical stochastic programming separates the decision-making process into an estimation phase and a subsequent optimization phase. The estimation method is typically selected with the goal to achieve maximum prediction accuracy but without tailoring it to the optimization problem at hand.

In this paper, we develop a method of data-driven stochastic programming that avoids the artificial decoupling of estimation and optimization and that

chooses an estimator that adapts to the underlying optimization problem. Specifically, we model data-driven solutions to a stochastic program through a *predictor* and its corresponding *prescriptor*. For any fixed feasible decision, the predictor maps the observable data to an estimate of the decision's expected cost. The prescriptor, on the other hand, computes a decision that minimizes the cost estimated by the predictor.

The set of all possible predictors and their induced prescriptors is vast. Indeed, there are countless possibilities to estimate the expected costs of a fixed decision from data, for example, via the popular sample average approximation (Shapiro et al. 2014, chapter 5), by postulating a parametric model for the exogenous uncertainties and estimating its parameters via maximum likelihood estimation (Dupačová and Wets 1988), or through kernel density estimation (Parpas et al. 2015). Recently, it has become fashionable to construct conservative (pessimistic) estimates of the expected costs via methods of distributionally robust optimization. In this setting, the available data are used to generate an ambiguity set

that represents a confidence region in the space of probability distributions and contains the unknown data-generating distribution with high probability. The expected cost of a fixed decision under the unknown true distribution is then estimated by the worst-case expectation over all distributions in the ambiguity set. Since the ambiguity set constitutes a confidence region for the unknown true distribution, the worst-case expectation represents an upper confidence bound on the true expected cost. The ambiguity set can be defined, for example, through confidence intervals for the distribution's moments (Delage and Ye 2010). Alternatively, the ambiguity set may contain all distributions that achieve a prescribed level of likelihood (Wang et al. 2016), that pass a statistical hypothesis test (Bertsimas et al. 2018a), or that are sufficiently close to a reference distribution with respect to a probability metric such as the Prokhorov metric (Erdoğan and Iyengar 2006), the Wasserstein distance (Pflug and Wozabal 2007, Mohajerin Esfahani and Kuhn 2018, Zhao and Guan 2018), the total variation distance (Sun and Xu 2016), or the L^1 -norm (Jiang and Guan 2018). Ben-Tal et al. (2013) have shown that confidence sets for distributions can also be constructed using ϕ -divergences such as the Pearson divergence, the Burg entropy, or the Kullback-Leibler divergence. More recently, Bayraksan and Love (2015) provide a systematic classification of ϕ -divergences and investigate the richness of the corresponding ambiguity sets.

Given the numerous possibilities for constructing predictors from a given data set, it is easy to lose oversight. In practice, predictors are often selected manually from within a small menu with the goal to meet certain statistical and/or computational requirements. However, there are typically many different predictors that exhibit the desired properties, and there always remains some doubt as to whether the chosen predictor is best suited for the particular decision problem at hand. In this paper, we propose a principled approach to data-driven stochastic programming by solving a meta-optimization problem over a rich class of predictor-prescriptor pairs including, among others, all examples reviewed in this section. This meta-optimization problem aims to find the least conservative (i.e., pointwise smallest) prescriptor whose out-of-sample disappointment decays at a prescribed exponential rate r as the sample size tends to infinity—irrespective of the true data-generating distribution. The out-of-sample disappointment quantifies the probability that the actual expected cost of the prescriptor exceeds its predicted cost. Put differently, it represents the probability that the predicted cost of a candidate decision is over-optimistic and leads to disappointment in out-of-sample tests. Thus, the proposed meta-optimization problem tries to

identify the predictor-prescriptor-pairs that overestimate the expected out-of-sample costs by the least amount possible without risking disappointment under any thinkable data-generating distribution.

Our main results can be summarized as follows:

- By leveraging Sanov's theorem from large deviations theory, we prove that the meta-optimization problem admits a unique optimal solution for any given stochastic program.

- We show that the optimal data-driven predictor estimates the expected costs under the unknown true distribution by a worst-case expectation over all distributions within a given relative entropy distance from the empirical distribution of the data. This suggests that, among all possible data-driven solutions, a distributionally robust approach based on a relative entropy ambiguity set is optimal. This is perhaps surprising because the meta-optimization problem does not impose any structure on the predictors, which are generic functions of the data. In particular, there is no requirement forcing predictors to admit a distributionally robust interpretation.

- In contrast to most of the existing work on data-driven distributionally robust optimization, our relative entropy ambiguity set does not play the role of a confidence region that contains the unknown data-generating distribution with a prescribed level of probability (see the following discussions of Lam 2016, Gupta 2019 for exceptions). Instead, the radius of the relative entropy ambiguity set coincides with the desired exponential decay rate r of the out-of-sample disappointment imposed by the meta-optimization problem.

- We prove that the optimal (distributionally robust) predictor admits a dual representation as the optimal value of a one-dimensional convex optimization problem that can be solved highly efficiently. For continuously distributed problem parameters, this representation seems to be new.

To our best knowledge, we are the first to recognize the optimality of distributionally robust optimization in its ability to transform data to predictors and prescriptors. The optimal distributionally robust predictor identified in this paper can be evaluated by solving a tractable convex optimization problem. Under standard convexity assumptions about the feasible set and the cost function of the stochastic program, the corresponding optimal prescriptor can also be evaluated in polynomial time. Although perhaps desirable, the tractability and distributionally robust nature of the optimal predictor-prescriptor-pair are not dictated *ex ante* but emerge naturally.

Relative entropy ambiguity sets have already attracted considerable interest in distributionally robust optimization (Calafiore 2007, Ben-Tal et al. 2013, Hu and Hong 2013, Lam 2016, Wang et al. 2016).

Note, however, that the relative entropy constitutes an asymmetric distance measure between two distributions. The asymmetry implies, among others, that the first distribution must be absolutely continuous to the second one but not vice versa. Thus, ambiguity sets can be constructed in two different ways by designating the reference distribution either as the first or as the second argument of the relative entropy. All papers listed previously favor the second option, and thus the emerging ambiguity sets contain only distributions that are absolutely continuous to the reference distribution. Maybe surprisingly, the optimal predictor resulting from our meta-optimization problem uses the reference distribution as the first argument of the relative entropy instead. Thus, the reference distribution is absolutely continuous to every distribution in the emerging ambiguity set. Relative entropy balls of this kind have previously been studied by Lam (2016), Bertsimas et al. (2018b), and Gupta (2019).

Adopting a Bayesian perspective, Gupta (2019) determines the smallest ambiguity sets that contain the unknown data-generating distribution with a prescribed level of confidence as the sample size tends to infinity. Both Pearson divergence and relative entropy ambiguity sets with properly scaled radii are optimal in this setting. In the terminology of the present paper, Gupta (2019) thus restricts attention to the subclass of distributionally robust predictors and operates with an asymptotic notion of optimality. The meta-optimization problem proposed here entails a stronger notion of optimality, under which the distributionally robust predictor with relative entropy ambiguity set emerges as the unique optimizer. Lam (2016) also seeks distributionally robust predictors that trade conservatism for out-of-sample performance. He studies the probability that the estimated expected cost function dominates the actual expected cost function uniformly across all decisions, and he calls a predictor optimal if this probability is asymptotically equal to a prescribed confidence level. Using the empirical likelihood theorem of Owen (1988), he shows that Pearson divergence and relative entropy ambiguity sets with properly scaled radii are optimal in this sense. This notion of optimality has again an asymptotic flavor in the sense that it refers to sequences of ambiguity sets that converge to a singleton, and it admits multiple optimizers.

The rest of the paper unfolds as follows. Section 2 provides a formal introduction to data-driven stochastic programming on finite state spaces and develops the meta-optimization problem for identifying the best predictor-prescriptor-pair. Section 3 reviews weak and strong large deviation principles, which are then used in Section 4 to determine the unique optimal solution of the meta-optimization problem.

An extension to continuous state spaces is discussed in Section 5.

1.1. Notation

The natural logarithm of $p \in \mathbb{R}_+$ is denoted by $\log(p)$, where we use the conventions $0 \log(0/p) = 0$ for any $p \geq 0$ and $p' \log(p'/0) = \infty$ for any $p' > 0$. A function $f: \mathcal{P} \rightarrow X$ from $\mathcal{P} \subseteq \mathbb{R}^d$ to $X \subseteq \mathbb{R}^n$ is called quasi-continuous at $\mathbb{P} \in \mathcal{P}$ if for every $\epsilon > 0$ and neighborhood $U \subseteq \mathcal{P}$ of $\mathbb{P}$ there is a nonempty open set $V \subseteq U$ with $|f(\mathbb{P}) - f(\mathbb{Q})| \leq \epsilon$ for all $\mathbb{Q} \in V$. Note that V does not necessarily contain $\mathbb{P}$. For any logical statement $\mathcal{E}$, the indicator function $\mathbb{I}_{\mathcal{E}}$ evaluates to 1 if $\mathcal{E}$ is true and to 0 otherwise.

2. Data-Driven Stochastic Programming

Stochastic programming is a powerful modeling paradigm for taking informed decisions in an uncertain environment. A generic single-stage stochastic program can be represented as

$$\underset{x \in X}{\text{minimize}} \mathbb{E}_{\mathbb{P}^*}[\gamma(x, \xi)]. \quad (1)$$

Here, the goal is to minimize the expected value of a cost function $\gamma(x, \xi) \in \mathbb{R}$, which depends both on a decision variable $x \in X$ and a random parameter $\xi \in \Xi$ governed by a probability distribution $\mathbb{P}^*$. We will assume that the cost $\gamma(x, \xi)$ is continuous in x for every fixed $\xi \in \Xi$, the feasible set $X \subseteq \mathbb{R}^n$ is compact, and $\Xi = \{1, \dots, d\}$ is finite. Thus, ξ has d distinct scenarios that are represented—without loss of generality—by the integers $1, \dots, d$. We will relax this requirement in Section 5, where Ξ will be modeled as an arbitrary compact subset of $\mathbb{R}^d$. A wide spectrum of decision problems can be cast as instances of (1). Shapiro et al. (2014) point out, for example, that (1) can be viewed as the first stage of a two-stage stochastic program, where the cost function $\gamma(x, \xi)$ embodies the optimal value of a subordinate second-stage problem. Alternatively, problem (1) may also be interpreted as a generic learning problem in the spirit of statistical learning theory.

In the following, we distinguish the prediction problem, which merely aims to predict the expected cost associated with a fixed decision x , and the prescription problem, which seeks to identify a decision x^* that minimizes the expected cost across all $x \in X$.

Any attempt to solve the prescription problem seems futile unless there is a procedure for solving the corresponding prediction problem. The generic prediction problem is closely related to what Le Maître and Knio (2010) call an uncertainty quantification problem and is therefore of prime interest in its own right. Throughout the rest of the paper, we thus analyze prediction and prescription problems on equal footing.

In the what follows we formalize the notion of a data-driven solution to the prescription and prediction problems, respectively. Furthermore, we introduce the basic assumptions as well as the notation used throughout the remainder of the paper.

2.1. Data-Driven Predictors and Prescriptors

If the distribution $\mathbb{P}^*$ of ξ is unobservable and must be estimated from a training data set consisting of finitely many independent samples from $\mathbb{P}^*$, we lack essential information to evaluate the expected cost of any fixed decision and to solve the stochastic program (1). The standard approach to overcome this deficiency is to approximate $\mathbb{P}^*$ with a parametric or nonparametric estimate $\hat{\mathbb{P}}$ inferred from the samples and to minimize the expected cost under $\hat{\mathbb{P}}$ instead of the true expected cost under $\mathbb{P}^*$. However, if we calibrate a stochastic program to a training data set and evaluate its optimal decision on a test data set, then the resulting test performance is often disappointing—even if the two data sets are sampled independently from $\mathbb{P}^*$. This phenomenon has been observed in many different contexts. It is particularly pronounced in finance, where Michaud (1989) refers to it as the “error maximization effect” of portfolio optimization, and in statistics or machine learning, where it is known as “overfitting.” In decision analysis, Smith and Winkler (2006) refer to it as the “optimizer’s curse.” Thus, when working with data instead of exact probability distributions, one should safeguard against solutions that display promising in-sample performance but lead to out-of-sample disappointment.

Initially, the distribution $\mathbb{P}^*$ is only known to belong to the probability simplex $\mathcal{P} = \{\mathbb{P} \in \mathfrak{N}_+^d : \sum_{i \in \Xi} \mathbb{P}(i) = 1\}$. Over time, however, independent samples ξ_t , $t \in \mathbb{N}$, from $\mathbb{P}^*$ are revealed to the decision maker that provide increasingly reliable statistical information about $\mathbb{P}^*$.

Any $\mathbb{P} \in \mathcal{P}$ encodes a possible probabilistic model for the data process. Thus, by slight abuse of terminology, we will henceforth refer to the distributions $\mathbb{P} \in \mathcal{P}$ as models and to $\mathcal{P}$ as the model class. Evidently, the true model $\mathbb{P}^*$ is an (albeit unknown) element of $\mathcal{P}$. Next, we introduce model-based predictors and prescriptors corresponding to the stochastic program (1), where the true unknown distribution $\mathbb{P}^*$ is replaced with a hypothetical model $\mathbb{P} \in \mathcal{P}$.

Definition 1 (Model-Based Predictors and Prescriptors). For any fixed model $\mathbb{P} \in \mathcal{P}$, we define the model-based predictor $c(x, \mathbb{P}) = \mathbb{E}_{\mathbb{P}}[\gamma(x, \xi)] = \sum_{i \in \Xi} \mathbb{P}(i) \gamma(x, i)$ as the expected cost of a given decision $x \in X$ and the model-based prescriptor $x^*(\mathbb{P}) \in \arg \min_{x \in X} c(x, \mathbb{P})$ as a decision that minimizes $c(x, \mathbb{P})$ over $x \in X$.

Note that the model-based predictor $c(x, \mathbb{P})$ is jointly continuous in x and $\mathbb{P}$ because Ξ is finite and $\gamma(x, \xi)$ is continuous in x for every fixed $\xi \in \Xi$. The continuity of $c(x, \mathbb{P})$ then guarantees via the compactness of X that the model-based prescriptor $x^*(\mathbb{P})$ exists for every model $\mathbb{P} \in \mathcal{P}$. In view of Definition 1, the stochastic program (1) can be identified with the prescription problem of computing $x^*(\mathbb{P}^*)$. Similarly, the evaluation of the expected cost of a given decision $x \in X$ in (1) can be identified with the prediction problem of computing $c(x, \mathbb{P}^*)$. These prediction and prescription problems cannot be solved, however, as they depend on the unknown true model $\mathbb{P}^*$.

If one has only access to a finite set $\{\xi_t\}_{t=1}^T$ of independent samples from $\mathbb{P}^*$ instead of $\mathbb{P}^*$ itself, then it may be useful to construct an empirical estimator for $\mathbb{P}^*$.

Definition 2 (Empirical Distribution). The empirical distribution $\hat{\mathbb{P}}_T$ corresponding to the sample path $\{\xi_t\}_{t=1}^T$ of length T is defined through

$$\hat{\mathbb{P}}_T(i) = \frac{1}{T} \sum_{t=1}^T \mathbb{1}_{\xi_t=i} \quad \forall i \in \Xi.$$

Note that $\hat{\mathbb{P}}_T$ can be viewed as the vector of empirical state frequencies. Indeed, its i th entry records the proportion of time that the sample path spends in state i . As the samples are drawn independently, the state frequencies capture all useful statistical information about $\mathbb{P}^*$ that can possibly be extracted from a given sample path. Note also that $\hat{\mathbb{P}}_T$ is in fact the maximum likelihood estimator of $\mathbb{P}^*$. In the following, we will therefore approximate the unknown predictor $c(x, \mathbb{P}^*)$ as well as the unknown prescriptor $x^*(\mathbb{P}^*)$ by suitable functions of the empirical distribution $\hat{\mathbb{P}}_T$.

Definition 3 (Data-Driven Predictors and Prescriptors). A continuous function $\hat{c} : X \times \mathcal{P} \rightarrow \mathfrak{N}$ is called a data-driven predictor if $\hat{c}(x, \hat{\mathbb{P}}_T)$ is used as an approximation for $c(x, \mathbb{P}^*)$. A quasi-continuous function $\hat{x} : \mathcal{P} \rightarrow X$ is called a data-driven prescriptor if there exists a data-driven predictor $\hat{c}$ with

$$\hat{x}(\mathbb{P}') \in \arg \min_{x \in X} \hat{c}(x, \mathbb{P}')$$

for all possible estimator realizations $\mathbb{P}' \in \mathcal{P}$, and $\hat{x}(\hat{\mathbb{P}}_T)$ is used as an approximation for $x^*(\mathbb{P}^*)$.

Every data-driven predictor $\hat{c}$ induces a data-driven prescriptor $\hat{x}$. To see this, note that the argmin mapping is non-empty-valued and upper semicontinuous due to Berge’s maximum theorem (Berge 1963), which applies because $\hat{c}$ is continuous and X is both compact and independent of $\mathbb{P}'$. Corollary 4 in Matejdes (1987), which applies because $\mathcal{P}$ is a Baire space and X is a

metric space, thus ensures that the argmin mapping admits a quasi-continuous selector, which serves as a valid data-driven prescriptor. One can show that the set of points where this quasi-continuous prescriptor is discontinuous is a meagre subset of $\mathcal{P}$ (Bledsoe 1952). By the Baire category theorem, the points of continuity of the data-driven prescriptor at hand are thus dense in $\mathcal{P}$ (Baire 1899). Thus, data-driven prescriptors in the sense of Definition 3 are mostly continuous.

Example 1 (Sample Average Predictor). The model-based predictor c introduced in Definition 1 constitutes a simple data-driven predictor $\hat{c} = c$, that is, $c(x, \hat{\mathbb{P}}_T)$ can readily be used as a naïve approximation for $c(x, \mathbb{P}^*)$. Note that the model-based predictor c is indeed continuous as desired. By the definition of the empirical estimator, this naïve predictor approximates $c(x, \mathbb{P}^*)$ with

$$c(x, \hat{\mathbb{P}}_T) = \frac{1}{T} \sum_{t=1}^T \gamma(x, \xi_t),$$

which is readily recognized as the popular sample average approximation.

2.2. Optimizing over All Data-Driven Predictors and Prescriptors

The estimates $\hat{c}(x, \hat{\mathbb{P}}_T)$ and $\hat{x}(\hat{\mathbb{P}}_T)$ inherit the randomness from the empirical estimator $\hat{\mathbb{P}}_T$, which is constructed from the (random) samples $\{\xi_t\}_{t=1}^T$. Note that the prediction and prescription problems are naturally interpreted as instances of statistical estimation problems. Indeed, data-driven prediction aims to estimate the expected cost $c(x, \mathbb{P}^*)$ from data. Standard statistical estimation theory would typically endeavor to find a data-driven predictor $\hat{c}$ that (approximately) minimizes the mean squared error,

$$\mathbb{E} \left[\left| c(x, \mathbb{P}^*) - \hat{c}(x, \hat{\mathbb{P}}_T) \right|^2 \right],$$

over some appropriately chosen class of predictors $\hat{c}$, where the expectation is taken with respect to the distribution $(\mathbb{P}^*)^\infty$ governing the sample path and the empirical estimator. The mean squared error penalizes the mismatch between the actual cost $c(x, \mathbb{P}^*)$ and its estimator $\hat{c}(x, \hat{\mathbb{P}}_T)$. Events in which we are left disappointed ($c(x, \mathbb{P}^*) > \hat{c}(x, \hat{\mathbb{P}}_T)$) are not treated differently from positive surprises ($c(x, \mathbb{P}^*) < \hat{c}(x, \hat{\mathbb{P}}_T)$). In a decision-making context where the goal is to minimize costs, however, disappointments (underestimated costs) are more harmful than positive surprises (overestimated costs). Although statisticians strive for accuracy by minimizing a symmetric estimation error, decision makers endeavor to limit the one-sided prediction disappointment.

Definition 4 (Out-of-Sample Disappointment). For any data-driven predictor $\hat{c}$, the probability

$$\mathbb{P}^\infty(c(x, \mathbb{P}) > \hat{c}(x, \hat{\mathbb{P}}_T)) \quad (2a)$$

is referred to as the out-of-sample prediction disappointment of $x \in X$ under model $\mathbb{P} \in \mathcal{P}$. Similarly, for any data-driven prescriptor $\hat{x}$ induced by a data-driven predictor $\hat{c}$, the probability

$$\mathbb{P}^\infty(c(\hat{x}(\hat{\mathbb{P}}_T), \mathbb{P}) > \hat{c}(\hat{x}(\hat{\mathbb{P}}_T), \hat{\mathbb{P}}_T)) \quad (2b)$$

is termed the out-of-sample prescription disappointment under model $\mathbb{P} \in \mathcal{P}$.

The out-of-sample prediction disappointment quantifies the probability (with respect to the sample path distribution $\mathbb{P}^\infty$ under some model $\mathbb{P} \in \mathcal{P}$) that the expected cost $c(x, \mathbb{P})$ of a fixed decision x exceeds the predicted cost $\hat{c}(x, \hat{\mathbb{P}}_T)$. Thus, the out-of-sample prediction disappointment is independent of the actual realization of the empirical estimator $\hat{\mathbb{P}}_T$, but depends on the hypothesized model $\mathbb{P}$. A similar statement holds for the out-of-sample prescription disappointment.

The main objective of this paper is to construct attractive data-driven predictors and prescriptors, which are optimal in a sense to be made precise later. We first develop a notion of optimality for data-driven predictors and extend it later to data-driven prescriptors. As indicated earlier, a crucial requirement for any data-driven predictor is that it must limit the out-of-sample disappointment. This informal requirement can be operationalized either in an asymptotic sense or in a finite sample sense:

i. Asymptotic guarantee: As T grows, the out-of-sample prediction disappointment (2a) decays exponentially at a rate at least equal to $r > 0$ up to first order in the exponent, that is,

$$\limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(c(x, \mathbb{P}) > \hat{c}(x, \hat{\mathbb{P}}_T)) \leq -r \quad \forall x \in X, \mathbb{P} \in \mathcal{P}. \quad (3)$$

ii. Finite sample guarantee: For every fixed T , the out-of-sample prediction disappointment (2a) is bounded above by a known function $g(T)$ that decays exponentially at rate at least equal to $r > 0$ to first order in the exponent, that is,

$$\mathbb{P}^\infty(c(x, \mathbb{P}) > \hat{c}(x, \hat{\mathbb{P}}_T)) \leq g(T) \quad \forall x \in X, \mathbb{P} \in \mathcal{P}, T \in \mathbb{N}, \quad (4)$$

where $\limsup_{T \rightarrow \infty} \frac{1}{T} \log g(T) \leq -r$.

The inequalities (3) and (4) are imposed across all models $\mathbb{P} \in \mathcal{P}$. This ensures that they are satisfied under the true model $\mathbb{P}^*$, which is only known to reside within $\mathcal{P}$. By requiring the inequalities to hold for all $x \in X$, we further ensure that the out-of-sample prediction disappointment is eventually small irrespective of

the chosen decision. Note that the finite sample guarantee (4) is sufficient but not necessary for the asymptotic guarantee (3). Knowing the finite sample bounds $g(T)$ has the advantage, among others, that one can determine the sample complexity,

$$\min\{T_0 \in \mathbb{N} : g(T) \leq \beta, \forall T \geq T_0\},$$

that is, the minimum number of samples needed to certify that the out-of-sample prediction disappointment does not exceed a prescribed significance level $\beta \in [0, 1]$.

At first sight, the requirements (3) and (4) may seem restrictive, and the existence of data-driven predictors with exponentially decaying out-of-sample disappointment may be questioned. Later we will argue, however, that these requirements are in fact natural and satisfied by all reasonable predictors. To see this, note that if the training data are generated by $\mathbb{P}$, then the empirical distribution $\hat{\mathbb{P}}_T$ converges $\mathbb{P}^\infty$ -almost surely to $\mathbb{P}$ by virtue of the strong law of large numbers. Thus, the out-of-sample disappointment of a predictor $\hat{c}$ with $\hat{c}(x, \mathbb{P}) > c(x, \mathbb{P})$ must decay to 0 as T grows. Conversely, if $\hat{c}(x, \mathbb{P}) < c(x, \mathbb{P})$, then the out-of-sample disappointment of $\hat{c}$ must approach 1 as T tends to infinity. The following example shows that the out-of-sample disappointment generically fails to vanish asymptotically in the limiting case when $\hat{c}(x, \mathbb{P}) = c(x, \mathbb{P})$.

Example 2 (Large Out-of-Sample Disappointment). Set the cost function to $\gamma(x, \xi) = \xi$. In this case, the sample average predictor approximates the expected cost $c(x, \mathbb{P}) = \sum_{i \in \Xi} i \mathbb{P}(i)$ by its sample mean $c(x, \hat{\mathbb{P}}_T) = \frac{1}{T} \sum_{t=1}^T \xi_t$. As the sample size T tends to infinity, the central limit theorem implies that

$$\sqrt{T} [c(x, \hat{\mathbb{P}}_T) - c(x, \mathbb{P})]$$

converges in law to a normal distribution with mean 0 and variance $\mathbb{E}_{\mathbb{P}}[(\xi - \mathbb{E}_{\mathbb{P}}[\xi])^2]$. Thus,

$$\begin{aligned} \lim_{T \rightarrow \infty} \mathbb{P}^\infty(c(x, \mathbb{P}) > \hat{c}(x, \hat{\mathbb{P}}_T)) \\ = \lim_{T \rightarrow \infty} \mathbb{P}^\infty(\sqrt{T}(\hat{c}(x, \hat{\mathbb{P}}_T) - c(x, \mathbb{P})) < 0) = \frac{1}{2}, \end{aligned}$$

which means that the out-of-sample prediction disappointment remains large for all sample sizes. The sample average predictor hence violates the asymptotic guarantee (3) and the stronger finite sample guarantee (4). Note that by adding any positive constant to the sample average predictor, we recover a predictor with exponentially decaying out-of-sample disappointment.

In the following we call a predictor $\hat{c}$ conservative if $\hat{c}(x, \mathbb{P}') > c(x, \mathbb{P}')$ for all decisions $x \in X$ and estimator realizations $\mathbb{P}' \in \mathcal{P}$. The previous discussion shows that if we require the out-of-sample disappointment

to decay asymptotically, we must focus on conservative predictors. Basic results from large deviations theory further ensure that the out-of-sample disappointment of any conservative predictor necessarily decays at an exponential rate. Specifically, asymptotic guarantees of the type (3) hold whenever the empirical distribution $\hat{\mathbb{P}}_T$ satisfies a weak large deviation principle, while finite sample guarantees of the type (4) hold when $\hat{\mathbb{P}}_T$ satisfies a strong large deviation principle. As will be shown in Section 3, the empirical distribution does satisfy weak and strong large deviation principles. One predictor that fails to be conservative is the sample average predictor.

For ease of exposition, we henceforth denote by $\mathcal{C}$ the set of all data-driven predictors, that is, all continuous functions that map $X \times \mathcal{P}$ to the reals. Moreover, we introduce a partial order $\leq_{\mathcal{C}}$ on $\mathcal{C}$ defined through

$$\hat{c}_1 \leq_{\mathcal{C}} \hat{c}_2 \iff \hat{c}_1(x, \mathbb{P}') \leq \hat{c}_2(x, \mathbb{P}') \quad \forall x \in X, \mathbb{P}' \in \mathcal{P},$$

for any $\hat{c}_1, \hat{c}_2 \in \mathcal{C}$. Thus, $\hat{c}_1 \leq_{\mathcal{C}} \hat{c}_2$ means that $\hat{c}_1$ is (weakly) less conservative than $\hat{c}_2$. The problem of finding the least conservative predictor among all data-driven predictors whose out-of-sample disappointment decays at rate at least $r > 0$ can thus be formalized as the following vector optimization problem:

$$\begin{aligned} & \underset{\hat{c} \in \mathcal{C}}{\text{minimize}} \quad \hat{c} \\ & \text{subject to} \quad \limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(c(x, \mathbb{P}) > \hat{c}(x, \hat{\mathbb{P}}_T)) \leq -r \\ & \quad \quad \quad \forall x \in X, \mathbb{P} \in \mathcal{P}. \end{aligned} \quad (5)$$

We highlight that the minimization in (5) is understood with respect to the partial order $\leq_{\mathcal{C}}$. Thus, the relation $\hat{c}_1 \leq_{\mathcal{C}} \hat{c}_2$ between two feasible decision means that $\hat{c}_1$ is weakly preferred to $\hat{c}_2$. However, not all pairs of feasible decisions are comparable, that is, it is possible that both $\hat{c}_1 \not\leq_{\mathcal{C}} \hat{c}_2$ and $\hat{c}_2 \not\leq_{\mathcal{C}} \hat{c}_1$. A predictor $\hat{c}^*$ is a strongly optimal solution for (5) if it is feasible and weakly preferred to every other feasible solution (i.e., every $\hat{c} \neq \hat{c}^*$ feasible in (5) satisfies $\hat{c}^* \leq_{\mathcal{C}} \hat{c}$). Similarly, $\hat{c}^*$ is a weakly optimal solution for (5) if it is feasible and if every other solution preferred to $\hat{c}^*$ is infeasible (i.e., every $\hat{c} \neq \hat{c}^*$ with $\hat{c} \leq_{\mathcal{C}} \hat{c}^*$ is infeasible in (5)). Although vector optimization problems can have many weak solutions, we point out that strong solutions are necessarily unique. To see this, assume for the sake of contradiction that $\hat{c}_1^*$ and $\hat{c}_2^*$ are two strong solutions of (5). In this case, the strong optimality of $\hat{c}_1^*$ implies that $\hat{c}_2^* \leq_{\mathcal{C}} \hat{c}_1^*$, whereas the strong optimality of $\hat{c}_2^*$ implies that $\hat{c}_1^* \leq_{\mathcal{C}} \hat{c}_2^*$. These two relations imply that $\hat{c}_1^* = \hat{c}_2^*$, that is, there cannot be two different strongly optimal solutions.

We are now ready to construct a meta-optimization problem akin to (5), which enables us to identify the best prescriptor. To this end, we henceforth denote by $\mathcal{X}$ the set of all data-driven predictor-prescriptor pairs $(\hat{c}, \hat{x})$, where $\hat{c} \in \mathcal{C}$, and $\hat{x}$ is a prescriptor induced by $\hat{c}$ as per Definition 3. Moreover, we equip $\mathcal{X}$ with a partial order $\leq_{\mathcal{X}}$, which is defined through the following:

$$(\hat{c}_1, \hat{x}_1) \leq_{\mathcal{X}} (\hat{c}_2, \hat{x}_2) \iff \hat{c}_1(\hat{x}_1(\mathbb{P}'), \mathbb{P}') \leq \hat{c}_2(\hat{x}_2(\mathbb{P}'), \mathbb{P}') \quad \forall \mathbb{P}' \in \mathcal{P}.$$

Note that $\hat{c}_1 \leq_{\mathcal{C}} \hat{c}_2$ actually implies $(\hat{c}_1, \hat{x}_1) \leq_{\mathcal{X}} (\hat{c}_2, \hat{x}_2)$, but not vice versa. The problem of finding the least conservative predictor-prescriptor pair whose out-of-sample prescription disappointment decays at rate at least $r > 0$ can now be formalized as the following vector optimization problem:

$$\begin{aligned} & \underset{(\hat{c}, \hat{x}) \in \mathcal{X}}{\text{minimize}} \quad (\hat{c}, \hat{x}) \\ & \text{subject to} \quad \limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^{\infty}(c(\hat{x}(\hat{\mathbb{P}}_T), \mathbb{P}) > \hat{c}(\hat{x}(\hat{\mathbb{P}}_T), \hat{\mathbb{P}}_T)) \\ & \quad \leq -r \quad \forall \mathbb{P} \in \mathcal{P}. \end{aligned} \quad (6)$$

Generic vector optimization problems typically only admit weak solutions. In Section 4 we will show, however, that (5) as well as (6) admit (unique) strong solutions in closed form. In fact, we will show that these closed-form solutions have a natural interpretation as the solutions of convex distributionally robust optimization problems.

Remark 1 (Out-of-Sample and In-Sample Performance). The natural performance measure to quantify the goodness of a data-driven prescriptor $\hat{x}$ is its out-of-sample performance $c(\hat{x}(\hat{\mathbb{P}}_T), \mathbb{P}^*)$ under the true model $\mathbb{P}^*$. As $\mathbb{P}^*$ is unknown, however, the out-of-sample performance cannot be optimized directly. A naïve remedy would be to formulate a meta-optimization problem that minimizes the worst case (or some average) of the out-of-sample performance of $\hat{x}$ across all models $\mathbb{P} \in \mathcal{P}$. The approach proposed here optimizes the out-of-sample performance implicitly. Indeed, the meta-optimization problem (6) represents $\hat{x}$ as a minimizer of some predictor $\hat{c}$, where $\hat{c}(\hat{x}(\hat{\mathbb{P}}_T), \hat{\mathbb{P}}_T)$ should be interpreted as the in-sample performance of $\hat{x}$. Instead of minimizing the out-of-sample performance of $\hat{x}$, problem (6) minimizes the in-sample performance of $\hat{x}$ but ensures through the constraints on the disappointment that the out-of-sample performance is smaller than the in-sample performance with increasingly high confidence as the sample size grows. In this sense, problem (6) minimizes a tight upper bound on the out-of-sample performance of $\hat{x}$.

3. Large Deviation Principles

Large deviations theory provides bounds on the exact exponential rate at which the probabilities of atypical

estimator realizations decay under a model $\mathbb{P}$ as the sample size T tends to infinity. These bounds are expressed in terms of the relative entropy of $\hat{\mathbb{P}}_T$ with respect to $\mathbb{P}$.

Definition 5 (Relative Entropy). The relative entropy of an estimator realization $\mathbb{P}' \in \mathcal{P}$ with respect to a model $\mathbb{P} \in \mathcal{P}$ is defined as:

$$I(\mathbb{P}', \mathbb{P}) = \sum_{i \in \Xi} \mathbb{P}'(i) \log \left(\frac{\mathbb{P}'(i)}{\mathbb{P}(i)} \right),$$

where we use the conventions $0 \log(0/p) = 0$ for any $p \geq 0$ and $p' \log(p'/0) = \infty$ for any $p' > 0$.

The relative entropy is also known as information for discrimination, cross-entropy, information gain, or Kullback-Leibler divergence (Kullback and Leibler 1951). The following proposition summarizes the key properties of the relative entropy relevant for this paper.

Proposition 1 (Relative Entropy). The relative entropy enjoys the following properties:

- Information inequality:* $I(\mathbb{P}', \mathbb{P}) \geq 0$ for all $\mathbb{P}, \mathbb{P}' \in \mathcal{P}$, while $I(\mathbb{P}', \mathbb{P}) = 0$ if and only if $\mathbb{P}' = \mathbb{P}$.
- Convexity:* For all pairs $(\mathbb{P}'_1, \mathbb{P}_1), (\mathbb{P}'_2, \mathbb{P}_2) \in \mathcal{P} \times \mathcal{P}$, and $\lambda \in [0, 1]$ we have

$$I((1 - \lambda)\mathbb{P}'_1 + \lambda\mathbb{P}'_2, (1 - \lambda)\mathbb{P}_1 + \lambda\mathbb{P}_2).$$

- Lower semicontinuity:* $I(\mathbb{P}', \mathbb{P}) \geq 0$ is lower semicontinuous in $(\mathbb{P}', \mathbb{P}) \in \mathcal{P} \times \mathcal{P}$.

Proof. Assertions (i) and (ii) follow from theorems 2.6.3 and 2.7.2 in Cover and Thomas (2006), respectively, whereas assertion (iii) follows directly from the definition of the relative entropy and our standard conventions regarding the natural logarithm. $\square$

We now show that the empirical estimators satisfy a weak large deviation principle (LDP). This result follows immediately from a finite version of Sanov's classical theorem. A textbook proof using the so-called method of types can be found in Cover and Thomas (2006, theorem 11.4.1). As the proof is illuminating and to keep this paper self-contained, we sketch the proof in Online Appendix A.

Theorem 1 (Weak LDP). If the samples $\{\xi_t\}_{t \in \mathbb{N}}$ are drawn independently from some $\mathbb{P} \in \mathcal{P}$, then for every Borel set $\mathcal{D} \subseteq \mathcal{P}$, the sequence of empirical distributions $\{\hat{\mathbb{P}}_T\}_{T \in \mathbb{N}}$ satisfies

$$\limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^{\infty}(\hat{\mathbb{P}}_T \in \mathcal{D}) \leq - \inf_{\mathbb{P}' \in \mathcal{D}} I(\mathbb{P}', \mathbb{P}). \quad (7a)$$

If additionally $\mathbb{P} > 0$, then for every Borel set $\mathcal{D} \subseteq \mathcal{P}$ we have¹

$$\liminf_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^{\infty}(\hat{\mathbb{P}}_T \in \mathcal{D}) \geq - \inf_{\mathbb{P}' \in \text{int} \mathcal{D}} I(\mathbb{P}', \mathbb{P}). \quad (7b)$$

Note that the inequality (7a) provides an upper LDP bound on the exponential rate at which the probability of the event $\hat{\mathbb{P}}_T \in \mathcal{D}$ decays under model $\mathbb{P}$. This upper bound is expressed in terms of a convex optimization problem that minimizes the relative entropy of $\mathbb{P}'$ with respect to $\mathbb{P}$ across all estimator realizations $\mathbb{P}'$ within $\mathcal{D}$. Similarly, (7b) offers a lower LDP bound on the decay rate. Note that in (7b), the relative entropy is minimized over the interior of $\mathcal{D}$ instead of $\mathcal{D}$.

If the data-generating model $\mathbb{P}$ itself belongs to $\mathcal{D}$, then $\inf_{\mathbb{P}' \in \mathcal{D}} I(\mathbb{P}', \mathbb{P}) = I(\mathbb{P}, \mathbb{P}) = 0$, which leads to the trivial upper bound $\mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}) \leq 1$. On the other hand, if $\mathcal{D}$ has empty interior (e.g., if $\mathcal{D} = \{\mathbb{P}\}$ is a singleton containing only the true model), then $\inf_{\mathbb{P}' \in \text{int}(\mathcal{D})} I(\mathbb{P}', \mathbb{P}) = \infty$, which leads to the trivial lower bound $\mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}) \geq 0$. Nontrivial bounds are obtained if $\mathbb{P} \notin \mathcal{D}$ and $\text{int}(\mathcal{D}) \neq \emptyset$. In these cases, the relative entropy bounds the exponential rate at which the probability of the atypical event $\hat{\mathbb{P}}_T \in \mathcal{D}$ decays with T . For some sets $\mathcal{D}$, this rate of decay is precisely determined by the relative entropy. Specifically, a Borel set $\mathcal{D} \subseteq \mathcal{P}$ is called I-continuous under model $\mathbb{P}$ if

$$\inf_{\mathbb{P}' \in \text{int}(\mathcal{D})} I(\mathbb{P}', \mathbb{P}) = \inf_{\mathbb{P}' \in \mathcal{D}} I(\mathbb{P}', \mathbb{P}).$$

Clearly, every open set $\mathcal{D} \subseteq \mathcal{P}$ is I-continuous under any model $\mathbb{P}$. Moreover, as the relative entropy is continuous in $\mathbb{P}'$ for any fixed $\mathbb{P} > 0$, every Borel set $\mathcal{D} \subseteq \mathcal{P}$ with $\mathcal{D} \subseteq \text{cl}(\text{int}(\mathcal{D}))$ is I-continuous under $\mathbb{P}$ whenever $\mathbb{P} > 0$. The LDP (7) implies that for large T , the probability of an I-continuous set $\mathcal{D}$ decays at rate

$\inf_{\mathbb{P}' \in \mathcal{D}} I(\mathbb{P}', \mathbb{P})$ under model $\mathbb{P}$ to first order in the exponent, that is, we have

$$\mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}) = e^{-T \inf_{\mathbb{P}' \in \mathcal{D}} I(\mathbb{P}', \mathbb{P}) + o(T)}. \quad (8)$$

If we interpret the relative entropy $I(\mathbb{P}', \mathbb{P})$ as the distance of $\mathbb{P}$ from $\mathbb{P}'$, then the decay rate of $\mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D})$ coincides with the distance of the model $\mathbb{P}$ from the atypical event set $\mathcal{D}$ (see Figure 1). Moreover, if $\mathcal{D}$ is I-continuous under $\mathbb{P}$, then (8) implies that $\mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}) \leq \beta$ whenever

$$T \gtrsim \frac{1}{r} \cdot \log\left(\frac{1}{\beta}\right),$$

where $r = \inf_{\mathbb{P}' \in \mathcal{D}} I(\mathbb{P}', \mathbb{P})$ is the I-distance from $\mathbb{P}$ to the set $\mathcal{D}$, and $\beta \in (0, 1)$ is a prescribed significance level.

The weak LDP of Theorem 1 provides only asymptotic bounds on the decay rates of atypical events. However, one can also establish a strong LDP, which offers finite sample guarantees. Most results of this paper, however, are based on the weak LDP of Theorem 1.

Theorem 2 (Strong LDP). *If the samples $\{\xi_t\}_{t \in \mathbb{N}}$ are drawn independently from some $\mathbb{P} \in \mathcal{P}$, then for every Borel set $\mathcal{D} \subseteq \mathcal{P}$, the sequence of empirical distributions $\{\hat{\mathbb{P}}_T\}_{T \in \mathbb{N}}$ satisfies*

$$\mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}) \leq (T + 1)^d e^{-T \inf_{\mathbb{P}' \in \mathcal{D}} I(\mathbb{P}', \mathbb{P})} \quad \forall T \in \mathbb{N}. \quad (9)$$

Proof. The claim follows immediately from inequality (27) in the proof of Theorem 1 in Online Appendix A. Note that (27) does not rely on the assumption that $\mathbb{P} > 0$. $\square$

4. Distributionally Robust Predictors and Prescriptors Are Optimal

Armed with the fundamental results of large deviations theory, we now endeavor to identify the least conservative data-driven predictors and prescriptors whose out-of-sample disappointment decays at a rate no less than some prescribed threshold $r > 0$ under any model $\mathbb{P} \in \mathcal{P}$, that is, we aim to solve the vector optimization problems (5) and (6).

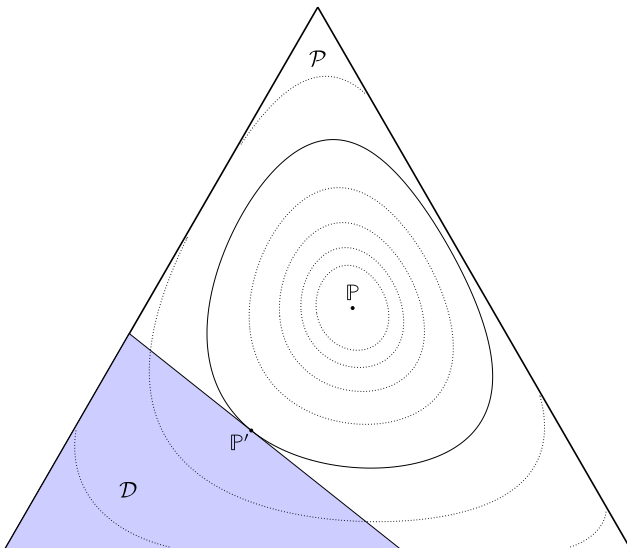
4.1. Distributionally Robust Predictors

The relative entropy lends itself to constructing a data-driven predictor in the sense of Definition 3. We will now show that this predictor is strongly optimal in (5).

Definition 6 (Distributionally Robust Predictors). For any fixed threshold $r \geq 0$, we define the data-driven predictor $\hat{c}_r : X \times \mathcal{P} \rightarrow \mathfrak{N}$ through

$$\hat{c}_r(x, \mathbb{P}') = \sup_{\mathbb{P} \in \mathcal{P}} \{c(x, \mathbb{P}) : I(\mathbb{P}', \mathbb{P}) \leq r\} \quad \forall x \in X, \mathbb{P}' \in \mathcal{P}. \quad (10)$$

Figure 1. Visualization of the LDP (7)



Notes: If $\mathcal{D} \subseteq \mathcal{P}$ is I-continuous and $\mathbb{P} \notin \mathcal{D}$, then the probability $\mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D})$ decays at the exponential rate $\inf_{\mathbb{P}' \in \mathcal{D}} I(\mathbb{P}', \mathbb{P})$, which can be viewed as the relative entropy distance of $\mathbb{P}$ from $\mathcal{D}$.

The data-driven predictor $\hat{c}_r$ admits a distributionally robust interpretation. In fact, $\hat{c}_r(x, \mathbb{P}')$ represents the worst-case expected cost associated with the decision x , where the worst case is taken across all models $\mathbb{P} \in \mathcal{P}$ whose relative entropy distance to $\mathbb{P}'$ is at most r . Observe that the supremum in (10) is always attained because $c(x, \mathbb{P})$ is linear in $\mathbb{P}$ and the feasible set of (10) is compact, which follows from the compactness of $\mathcal{P}$ and the lower semicontinuity of the relative entropy in $\mathbb{P}$ for any fixed $\mathbb{P}'$ (see Proposition 1(iii)). Note also that $\hat{c}_r(x, \mathbb{P}')$ can be evaluated efficiently because (10) constitutes a convex conic optimization problem with d decision variables. A particularly simple and efficient method to evaluate $\hat{c}_r(x, \mathbb{P}')$ is to solve the one-dimensional convex minimization problem dual to (10) by using bisection or another line search method.

Proposition 2 (Dual Representation of $\hat{c}_r$). *If $r > 0$ and $\bar{\gamma}(x) = \max_{i \in \Xi} \gamma(x, i)$ denotes the worst-case cost function, then the distributionally robust predictor admits the dual representation*

$$\hat{c}_r(x, \mathbb{P}') = \min_{\alpha \geq \bar{\gamma}(x)} \alpha - e^{-r} \prod_{i \in \Xi} (\alpha - \gamma(x, i))^{\mathbb{P}'(i)}. \quad (11)$$

Problem (11) has a minimizer α^ that satisfies $\bar{\gamma}(x) \leq \alpha^* \leq \frac{\bar{\gamma}(x) - e^{-r} c(x, \mathbb{P}')}{1 - e^{-r}}$.*

Proof. See Online Appendix A. $\square$

Remark 2 (Sample Average Predictor). For $r = 0$, the distributionally robust predictor $\hat{c}_r$ collapses to the sample average predictor of Example 1. Indeed, because of the strict positivity of the relative entropy $I(\mathbb{P}', \mathbb{P}) > 0$ for $\mathbb{P}' \neq \mathbb{P}$ (see Proposition 1(i)), we have that

$$\hat{c}_0(x, \mathbb{P}') = c(x, \mathbb{P}').$$

As shown in Example 2, the sample average predictor fails to offer asymptotic or finite sample guarantees of the form (3) and (4), respectively.

Remark 3 (Alternative Distributionally Robust Predictors). The relative entropy can also be used to construct a reverse distributionally robust predictor $\check{c}_r \in \mathcal{C}$ defined through

$$\check{c}_r(x, \mathbb{P}') = \sup_{\mathbb{P} \in \mathcal{P}} \{c(x, \mathbb{P}) : \mathbb{P} \ll \mathbb{P}', I(\mathbb{P}, \mathbb{P}') \leq r\} \quad \forall x \in X, \mathbb{P}' \in \mathcal{P}, \quad (12)$$

In contrast to $\hat{c}_r$, the reverse distributionally robust predictor $\check{c}_r$ fixes the second argument of the relative entropy and maximizes over the first argument. Note that $\check{c}_r$ can be viewed as the entropic value-at-risk of the uncertain cost $\gamma(x, \xi)$ (see Ahmadi-Javid 2012, theorem 3.3). Another predictor related to $\hat{c}$ is the

restricted distributionally robust predictor $\bar{c}_r \in \mathcal{C}$ defined through

$$\bar{c}_r(x, \mathbb{P}') = \sup_{\mathbb{P} \in \mathcal{P}} \{c(x, \mathbb{P}) : \mathbb{P} \ll \mathbb{P}', I(\mathbb{P}', \mathbb{P}) \leq r\} \quad \forall x \in X, \mathbb{P}' \in \mathcal{P}, \quad (13)$$

where $\mathbb{P} \ll \mathbb{P}'$ expresses the requirement that $\mathbb{P}$ must be absolutely continuous with respect to $\mathbb{P}'$. Formally, $\mathbb{P} \ll \mathbb{P}'$ means that $\mathbb{P}(i) = 0$ for all outcomes $i \in \Xi$ with $\mathbb{P}'(i) = 0$. By Ahmadi-Javid (2012, definition 5.1), $\bar{c}_r$ can be interpreted as the negative log-entropic risk of $\gamma(x, \xi)$.

The predictors $\hat{c}_r$ and $\check{c}_r$ differ because the relative entropy fails to be symmetric. We emphasize that the reverse predictor $\check{c}_r$ has appeared often in the literature on distributionally robust optimization (see, e.g., Calafiore 2007, Ben-Tal et al. 2013, Hu and Hong 2013, Lam 2016, Wang et al. 2016). The predictors $\hat{c}_r$ and $\bar{c}_r$ differ, too, because of the additional constraint $\mathbb{P} \ll \mathbb{P}'$, which is significant when not all outcomes in Ξ have been observed. The statistical properties of the predictor $\bar{c}_r$ have been analyzed by Lam (2016) and more recently by Duchi et al. (2016) from the perspective of the empirical likelihood theory introduced by Owen (1988). The predictor $\hat{c}_r$ suggested here has not yet been studied extensively even though—as we will demonstrate later—it displays attractive theoretical properties that are not shared by either $\check{c}_r$ or $\bar{c}_r$. The difference between $\hat{c}_r$ and $\check{c}_r$ or $\bar{c}_r$ is significant. Indeed, both $\check{c}_r$ and $\bar{c}_r$ hedge only against models $\mathbb{P}$ that are absolutely continuous with respect to the (observed realization of the) empirical distribution $\mathbb{P}'$. Although it is clear that the empirical distribution must be absolutely continuous with respect to the data-generating distribution, however, the converse implication is generally false. Indeed, an outcome can have positive probability even if it does not show up in a given finite time series. By taking the worst case only over models that are absolutely continuous with respect to $\mathbb{P}'$, both predictors $\check{c}_r$ and $\bar{c}_r$ potentially ignore many models that could have generated the observed data.

We first establish that $\hat{c}_r$ indeed belongs to the set $\mathcal{C}$ of all data-driven predictors, that is, the family of continuous functions mapping $X \times \mathcal{P}$ to the reals.

Proposition 3 (Continuity of $\hat{c}_r$). *If $r \geq 0$, then the distributionally robust predictor $\hat{c}_r$ is continuous on $X \times \mathcal{P}$.*

Proof. By Proposition 2, the distributionally robust predictor $\hat{c}_r$ admits the dual representation (11). Note that the objective function of (11) is manifestly continuous in $(\alpha, x, \mathbb{P}')$ and that (11) is guaranteed to have a minimizer in the compact interval $[\bar{\gamma}(x), \frac{\bar{\gamma}(x) - e^{-r} c(x, \mathbb{P}')}{1 - e^{-r}}]$, whose boundaries depend continuously on $(x, \mathbb{P}')$.

Consequently, the predictor $\hat{c}_r$ is continuous by Berge's celebrated maximum theorem (Berge 1963). $\square$

We now analyze the performance of the distributionally robust data-driven predictor $\hat{c}_r$ using arguments from large deviations theory. The parameter r encoding the predictor $\hat{c}_r$ captures the fundamental trade-off between out-of-sample disappointment and accuracy, which is inherent to any approach to data-driven prediction. Indeed, as r increases, the predictor $\hat{c}_r$ becomes more reliable in the sense that its out-of-sample disappointment decreases. However, increasing r also results in more conservative (pessimistically biased) predictions. In the following, we will demonstrate that $\hat{c}_r$ strikes indeed an optimal balance between reliability and conservatism.

Theorem 3 (Feasibility of $\hat{c}_r$). *If $r \geq 0$, then the predictor $\hat{c}_r$ is feasible in (5).*

Proof. From Proposition 3, we already know that $\hat{c}_r \in \mathcal{C}$. It remains to be shown that the out-of-sample disappointment of $\hat{c}_r$ decays at a rate of at least r . We have $c(x, \mathbb{P}) > \hat{c}_r(x, \hat{\mathbb{P}}_T)$ if and only if the estimator $\hat{\mathbb{P}}_T$ falls within the following disappointment set:

$$\mathcal{D}(x, \mathbb{P}) = \{\mathbb{P}' \in \mathcal{P} : c(x, \mathbb{P}) > \hat{c}_r(x, \mathbb{P}')\}.$$

Note that by the definition of $\hat{c}_r$, we have

$$\begin{aligned} I(\mathbb{P}', \mathbb{P}) \leq r &\Rightarrow \hat{c}_r(x, \mathbb{P}') \\ &= \sup_{\mathbb{P}'' \in \mathcal{P}} \{c(x, \mathbb{P}'') : I(\mathbb{P}', \mathbb{P}'') \leq r\} \geq c(x, \mathbb{P}). \end{aligned}$$

By contraposition, this implication is equivalent to

$$c(x, \mathbb{P}) > \hat{c}_r(x, \mathbb{P}') \Rightarrow I(\mathbb{P}', \mathbb{P}) > r.$$

Therefore, $\mathcal{D}(x, \mathbb{P})$ is a subset of

$$\mathcal{J}(\mathbb{P}) = \{\mathbb{P}' \in \mathcal{P} : I(\mathbb{P}', \mathbb{P}) > r\},$$

irrespective of $x \in X$. We thus have

$$\begin{aligned} \limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}(x, \mathbb{P})) \\ \leq \limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{J}(\mathbb{P})) \\ \leq - \inf_{\mathbb{P}' \in \mathcal{J}(\mathbb{P})} I(\mathbb{P}', \mathbb{P}) \leq -r, \end{aligned}$$

where the first inequality holds because $\mathcal{D}(x, \mathbb{P}) \subseteq \mathcal{J}(\mathbb{P})$, whereas the second inequality exploits the weak LDP upper bound (7a). Thus, $\hat{c}_r$ is feasible in (5). $\square$

Note that any predictor $\hat{c}$ with $\hat{c}_r \leq_{\mathcal{C}} \hat{c}$ has a smaller disappointment set than $\hat{c}_r$, and thus the out-of-sample disappointment of $\hat{c}$ decays at least as fast

as that of $\hat{c}_r$. Hence, $\hat{c}$ is also feasible in (5). In particular, this immediately implies that if we inflate the relative entropy ball of the distributionally robust predictor $\hat{c}_r$ to any larger ambiguity set, we obtain another predictor that is feasible in (5). As an example, consider the total variation predictor

$$\begin{aligned} \hat{c}_r^{\text{tv}}(x, \mathbb{P}') = \sup_{\mathbb{P} \in \mathcal{P}} \left\{ c(x, \mathbb{P}) : \|\mathbb{P} - \mathbb{P}'\|_{\text{tv}} \leq \sqrt{2r} \right\} \\ \forall x \in X, \mathbb{P}' \in \mathcal{P}, \end{aligned}$$

where $\|\mathbb{P} - \mathbb{P}'\|_{\text{tv}}$ denotes the total variation distance between $\mathbb{P}$ and $\mathbb{P}'$. Pinsker's classical inequality asserts that $\|\mathbb{P} - \mathbb{P}'\|_{\text{tv}} \leq \sqrt{2I(\mathbb{P}', \mathbb{P})}$ for all $\mathbb{P}$ and $\mathbb{P}'$ in $\mathcal{P}$. Thus, we have $\hat{c}_r \leq_{\mathcal{C}} \hat{c}_r^{\text{tv}}$, which implies that the total variation predictor is feasible in (5). This suggests that (5) has a rich feasible set.

The following main theorem establishes that $\hat{c}_r$ is not only a feasible but even a strongly optimal solution for the vector optimization problem (5). This means that if an arbitrary data-driven predictor $\hat{c}$ predicts a lower expected cost than $\hat{c}_r$, even for a single estimator realization $\mathbb{P}' \in \mathcal{P}$, then $\hat{c}$ must suffer from a higher out-of-sample disappointment than $\hat{c}_r$ to first order in the exponent.

Theorem 4 (Optimality of $\hat{c}_r$). *If $r > 0$, then $\hat{c}_r$ is strongly optimal in (5).*

Proof. Assume for the sake of argument that $\hat{c}_r$ fails to be a strong solution for (5). Thus, there exists a data-driven predictor $\hat{c} \in \mathcal{C}$ that is feasible in (5) but not dominated by $\hat{c}_r$, that is, $\hat{c}_r \not\leq_{\mathcal{C}} \hat{c}$. This means that there exists $x \in X$ and $\mathbb{P}'_0 \in \mathcal{P}$ with $\hat{c}_r(x, \mathbb{P}'_0) > \hat{c}(x, \mathbb{P}'_0)$. For later reference we set $\epsilon = \hat{c}_r(x, \mathbb{P}'_0) - \hat{c}(x, \mathbb{P}'_0) > 0$. In the remainder of the proof we will demonstrate that $\hat{c}$ cannot be feasible in (5), which contradicts our initial assumption.

Let $\mathbb{P}_0 \in \mathcal{P}$ be an optimal solution of problem (10) at $\mathbb{P}' = \mathbb{P}'_0$. Thus, we have $I(\mathbb{P}'_0, \mathbb{P}_0) \leq r$ and

$$\hat{c}_r(x, \mathbb{P}'_0) = c(x, \mathbb{P}_0). \quad (14)$$

In the following, we will first perturb $\mathbb{P}_0$ to obtain a model $\mathbb{P}_1$ that is $\frac{\epsilon}{2}$ -suboptimal in (10) but satisfies $I(\mathbb{P}'_0, \mathbb{P}_1) < r$. Subsequently, we will perturb $\mathbb{P}_1$ to obtain a model $\mathbb{P}_2$ that is ϵ -suboptimal in (10) but satisfies $I(\mathbb{P}'_0, \mathbb{P}_2) < r$ as well as $\mathbb{P}_2 > 0$.

To construct $\mathbb{P}_1$, consider all models $\mathbb{P}(\lambda) = \lambda \mathbb{P}'_0 + (1 - \lambda) \mathbb{P}_0$, $\lambda \in [0, 1]$, on the line segment between $\mathbb{P}'_0$ and $\mathbb{P}_0$. As r is strictly positive, the convexity of the relative entropy implies that

$$\begin{aligned} I(\mathbb{P}'_0, \mathbb{P}(\lambda)) &\leq \lambda I(\mathbb{P}'_0, \mathbb{P}'_0) + (1 - \lambda) I(\mathbb{P}'_0, \mathbb{P}_0) \\ &\leq (1 - \lambda)r < r \quad \forall \lambda \in (0, 1]. \end{aligned}$$

Moreover, as the expected cost $c(x, \mathbb{P}(\lambda))$ changes continuously in λ , there exists a sufficiently small $\lambda_1 \in (0, 1]$ such that $\mathbb{P}_1 = \mathbb{P}(\lambda_1)$ and $r_1 = I(\mathbb{P}'_0, \mathbb{P}_1)$ satisfy $0 < r_1 < r$ and

$$c(x, \mathbb{P}_0) < c(x, \mathbb{P}_1) + \frac{\epsilon}{2}.$$

To construct $\mathbb{P}_2$, we consider all models $\mathbb{P}(\lambda) = \lambda\mathbb{U} + (1 - \lambda)\mathbb{P}_1$, $\lambda \in [0, 1]$, on the line segment between the uniform distribution $\mathbb{U}$ on Ξ and $\mathbb{P}_1$. By the convexity of the relative entropy we have

$$\begin{aligned} I(\mathbb{P}'_0, \mathbb{P}(\lambda)) &\leq \lambda I(\mathbb{P}'_0, \mathbb{U}) + (1 - \lambda)I(\mathbb{P}'_0, \mathbb{P}_1) \\ &\leq \lambda I(\mathbb{P}'_0, \mathbb{U}) + (1 - \lambda)r_1 \quad \forall \lambda \in [0, 1]. \end{aligned}$$

As $r_1 < r$ and the expected cost $c(x, \mathbb{P}(\lambda))$ change continuously in λ , there exists a sufficiently small $\lambda_2 \in (0, 1]$ such that $\mathbb{P}_2 = \mathbb{P}(\lambda_2)$ and $r_2 = I(\mathbb{P}'_0, \mathbb{P}_2)$ satisfy $0 < r_2 < r$, $\mathbb{P}_2 > 0$, and

$$c(x, \mathbb{P}_0) < c(x, \mathbb{P}_2) + \epsilon. \quad (15)$$

In summary, we thus have

$$\begin{aligned} \hat{c}(x, \mathbb{P}'_0) &= \hat{c}_r(x, \mathbb{P}'_0) - \epsilon = c(x, \mathbb{P}_0) - \epsilon < c(x, \mathbb{P}_2) \\ &\leq \hat{c}_r(x, \mathbb{P}'_0), \end{aligned} \quad (16)$$

where the first equality follows from the definition of ϵ , and the second equality exploits (14). Moreover, the strict inequality holds due to (15), and the weak inequality follows from the definition of $\hat{c}_r$ and the fact that $I(\mathbb{P}'_0, \mathbb{P}_2) = r_2 < r$.

In the remainder of the proof, we will argue that the prediction disappointment $\mathbb{P}_2^\infty(c(x, \mathbb{P}_2) > \hat{c}(x, \hat{\mathbb{P}}_T))$ under model $\mathbb{P}_2$ decays at a rate of at most $r_2 < r$ as the sample size T tends to infinity. In analogy to the proof of Theorem 3, we define the set of disappointing estimator realizations as

$$\mathcal{D}(x, \mathbb{P}_2) = \{\mathbb{P}' \in \mathcal{P} : c(x, \mathbb{P}_2) > \hat{c}(x, \mathbb{P}')\}.$$

This set contains $\mathbb{P}'_0$ due to the strict inequality in (16). Moreover, as $\hat{c} \in \mathcal{C}$ is continuous, $\mathcal{D}(x, \mathbb{P}_2)$ is an open subset of $\mathcal{P}$. Thus, we find

$$\begin{aligned} \inf_{\mathbb{P}' \in \text{int}(\mathcal{D}(x, \mathbb{P}_2))} I(\mathbb{P}', \mathbb{P}_2) &= \inf_{\mathbb{P}' \in \mathcal{D}(x, \mathbb{P}_2)} I(\mathbb{P}', \mathbb{P}_2) \\ &\leq I(\mathbb{P}'_0, \mathbb{P}_2) = r_2, \end{aligned}$$

where the inequality holds because $\mathbb{P}'_0 \in \mathcal{D}(x, \mathbb{P}_0)$, and the last equality follows from the definition of r_2 . As the empirical distributions $\{\hat{\mathbb{P}}_T\}_{T \in \mathbb{N}}$ obey the LDP lower bound (7b) under $\mathbb{P}_2 > 0$, we finally conclude that

$$\begin{aligned} -r < -r_2 &\leq -\inf_{\mathbb{P}' \in \text{int}(\mathcal{D}(x, \mathbb{P}_2))} I(\mathbb{P}', \mathbb{P}_2) \\ &\leq \liminf_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}_2^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}(x, \mathbb{P}_2)). \end{aligned}$$

The previous chain of inequalities implies, however, that $\hat{c}$ is infeasible in problem (5). This contradicts our initial assumption, and thus, $\hat{c}_r$ must indeed be a strong solution of (5). $\square$

Theorem 4 asserts that the distributionally robust predictor $\hat{c}_r$ is optimal among all data-driven predictors representable as continuous functions of the empirical distribution $\hat{\mathbb{P}}_T$. That is, any attempt to make it less conservative invariably increases the out-of-sample prediction disappointment. We remark that the class of predictors that depend on the data only through $\hat{\mathbb{P}}_T$ is vast. These predictors constitute arbitrary continuous functions of the data that are independent of the order in which the samples were observed. As the samples are independent and identically distributed (i.i.d.), there are in fact no meaningful data-driven predictors that display a more general dependence on the data.

Note that in the previous discussion all guarantees are fundamentally asymptotic in nature. Using Theorem 2 one can show, however, that $\hat{c}_r$ also satisfies finite sample guarantees.

Theorem 5 (Finite Sample Guarantee). *The out-of-sample disappointment of the distributionally robust predictor $\hat{c}_r$ enjoys the following finite sample guarantee under any model $\mathbb{P} \in \mathcal{P}$ and for any $x \in X$:*

$$\mathbb{P}^\infty(c(x, \mathbb{P}) > \hat{c}_r(x, \hat{\mathbb{P}}_T)) \leq (T + 1)^d e^{-rT} \quad \forall T \in \mathbb{N}. \quad (17)$$

Proof. The proof of this result widely parallels that of Theorem 3 but uses the strong LDP upper bound (9) in lieu of the weak upper bound (7a). Details are omitted for brevity. $\square$

4.2. Distributionally Robust Prescriptors

The distributionally robust predictor $\hat{c}_r$ of Definition 6 induces a corresponding prescriptor.

Definition 7 (Distributionally Robust Prescriptors). Denote by $\hat{c}_r$, $r \geq 0$, the distributionally robust data-driven predictor of Definition 6. We can then define the data-driven prescriptor $\hat{x}_r : \mathcal{P} \rightarrow X$ as a quasi-continuous function satisfying

$$\hat{x}_r(\mathbb{P}') \in \arg \min_{x \in X} \hat{c}_r(x, \mathbb{P}') \quad \forall \mathbb{P}' \in \mathcal{P}. \quad (18)$$

Note that the minimum in (18) is attained because X is compact and $\hat{c}_r$ is continuous due to Proposition 3. Thus, there exists at least one function $\hat{x}_r$ satisfying (18). In the next proposition, we argue that this function can be chosen to be quasi-continuous as desired.

Proposition 4 (Quasi-Continuity of $\hat{x}_r$). *If $r \geq 0$, then there exists a quasi-continuous data-driven predictor $\hat{x}_r$ satisfying (18).*

Proof. Denote by $\Gamma(\mathbb{P}') = \arg \min_{x \in X} \hat{c}_r(x, \mathbb{P}')$ the argmin-mapping of problem (10), and observe that $\Gamma(\mathbb{P}')$ is

compact and nonempty for every $\mathbb{P}' \in \mathcal{P}$ because $\hat{c}_r$ is continuous and X is compact. As X is independent of $\mathbb{P}'$, Berge's maximum theorem (Berge 1963) further implies that Γ is upper semicontinuous. As $\mathcal{P}$ is a Baire space and X is a metric space, Matejdes (1987, corollary 4) finally guarantees that there exists a quasi-continuous function $\hat{x}_r : \mathcal{P} \rightarrow X$ with $\hat{x}_r(\mathbb{P}') \in \Gamma(\mathbb{P}')$ for all $\mathbb{P}' \in \mathcal{P}$. $\square$

Propositions 3 and 4 imply that $(\hat{c}_r, \hat{x}_r)$ belongs to the family $\mathcal{X}$ of all data-driven predictor-prescriptor pairs. Using a similar reasoning as in Theorem 3, we now demonstrate that the out-of-sample disappointment of $\hat{x}_r$ decays at rate at least r as T tends to infinity. Thus, $\hat{x}_r$ provides trustworthy prescriptions.

Theorem 6 (Feasibility of $(\hat{c}_r, \hat{x}_r)$). *If $r \geq 0$, then the predictor-prescriptor-pair $(\hat{c}_r, \hat{x}_r)$ is feasible in (6).*

Proof. Propositions 3 and 4 imply that $(\hat{c}_r, \hat{x}_r) \in \mathcal{X}$. It remains to be shown that the out-of-sample disappointment of $\hat{x}_r$ decays at a rate of at least r . To this end, define $\mathcal{D}(x, \mathbb{P})$ and $\mathcal{J}(\mathbb{P})$ as in the proof of Theorem 3, and recall that $\mathcal{D}(x, \mathbb{P}) \subseteq \mathcal{J}(\mathbb{P})$ for every decision $x \in X$ and model $\mathbb{P} \in \mathcal{P}$. Thus, for every fixed estimator realization $\mathbb{P}' \in \mathcal{P}$ the following implication holds:

$$\begin{aligned} c(\hat{x}_r(\mathbb{P}'), \mathbb{P}) &> \hat{c}_r(\hat{x}_r(\mathbb{P}'), \mathbb{P}') \\ \Rightarrow \exists x \in X \text{ with } c(x, \mathbb{P}) &> \hat{c}_r(x, \mathbb{P}') \\ \Rightarrow \mathbb{P}' \in \bigcup_{x \in X} \mathcal{D}(x, \mathbb{P}) \\ \Rightarrow \mathbb{P}' \in \mathcal{J}(\mathbb{P}), \end{aligned}$$

which in turn implies

$$\begin{aligned} \limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(c(\hat{x}_r(\hat{\mathbb{P}}_T), \mathbb{P}) > \hat{c}_r(\hat{x}_r(\hat{\mathbb{P}}_T), \hat{\mathbb{P}}_T)) \\ \leq \limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{J}(\mathbb{P})) \leq -r, \end{aligned}$$

for every model $\mathbb{P} \in \mathcal{P}$. Note that the second inequality in this expression has already been established in the proof of Theorem 3. Thus, the claim follows. $\square$

Next, we argue that $(\hat{c}_r, \hat{x}_r)$ is a strongly optimal solution for the vector optimization problem (6).

Theorem 7 (Optimality of $(\hat{c}_r, \hat{x}_r)$). *If $r > 0$, then $(\hat{c}_r, \hat{x}_r)$ is strongly optimal in (6).*

Proof. Assume for the sake of argument that $(\hat{c}_r, \hat{x}_r)$ fails to be a strong solution for (6). Thus, there exists a data-driven prescriptor $(\hat{c}, \hat{x}) \in \mathcal{X}$ that is feasible in (6) but not dominated by $(\hat{c}_r, \hat{x}_r)$, that is, $(\hat{c}_r, \hat{x}_r) \not\preceq_{\mathcal{X}} (\hat{c}, \hat{x})$. This means that there exists $\mathbb{P}'_0 \in \mathcal{P}$ with $\hat{c}_r(\hat{x}_r(\mathbb{P}'_0), \mathbb{P}'_0) > \hat{c}(\hat{x}(\mathbb{P}'_0), \mathbb{P}'_0)$. As X is compact and $\hat{c}$ is continuous, the cost $\hat{c}(\hat{x}(\mathbb{P}'), \mathbb{P}')$ of the prescriptor $\hat{x}$ under the corresponding predictor $\hat{c}$ is continuous in $\mathbb{P}'$ (Berge 1963). Similarly, $\hat{c}_r(\hat{x}_r(\mathbb{P}'), \mathbb{P}')$ is continuous in $\mathbb{P}'$. Recall also that $\hat{x}$ is quasi-continuous and therefore

continuous on a dense subset of $\mathcal{P}$ (Bledsoe 1952). Thus, we may assume without loss of generality that $\hat{x}$ is continuous at $\mathbb{P}'_0$. For later reference, we set $\epsilon = \hat{c}_r(\hat{x}(\mathbb{P}'_0), \mathbb{P}'_0) - \hat{c}(\hat{x}(\mathbb{P}'_0), \mathbb{P}'_0) > 0$.

In the remainder of the proof, we will demonstrate that $(\hat{c}, \hat{x})$ cannot be feasible in (6), which contradicts our initial assumption. To this end, let $\mathbb{P}_0 \in \mathcal{P}$ be an optimal solution of problem (10) at $x = \hat{x}(\mathbb{P}'_0)$ and $\mathbb{P}' = \mathbb{P}'_0$. Thus, we have $I(\mathbb{P}'_0, \mathbb{P}_0) \leq r$ and

$$\hat{c}_r(\hat{x}(\mathbb{P}'_0), \mathbb{P}'_0) = c(\hat{x}(\mathbb{P}'_0), \mathbb{P}_0). \quad (19)$$

Next, we first perturb $\mathbb{P}_0$ to obtain a model $\mathbb{P}_1$ that is strictly $\frac{\epsilon}{2}$ -suboptimal in (10) but satisfies $I(\mathbb{P}'_0, \mathbb{P}_1) = r_1 < r$. Subsequently, we perturb $\mathbb{P}_1$ to obtain a model $\mathbb{P}_2$ that is strictly ϵ -suboptimal in (10) but satisfies $I(\mathbb{P}'_0, \mathbb{P}_2) = r_2 < r$ as well and $\mathbb{P}_2 > 0$. The distributions $\mathbb{P}_1$ and $\mathbb{P}_2$ can be constructed exactly as in the proof of Theorem 4. Details are omitted for brevity. Thus, we find

$$\begin{aligned} \hat{c}(\hat{x}(\mathbb{P}'_0), \mathbb{P}'_0) &= \hat{c}_r(\hat{x}(\mathbb{P}'_0), \mathbb{P}'_0) - \epsilon = c(\hat{x}(\mathbb{P}'_0), \mathbb{P}_0) \\ &- \epsilon < c(\hat{x}(\mathbb{P}'_0), \mathbb{P}_2) \leq \hat{c}_r(\hat{x}(\mathbb{P}'_0), \mathbb{P}'_0), \end{aligned} \quad (20)$$

where the first equality follows from the definition of ϵ , and the second equality exploits (19). Moreover, the strict inequality holds because $\mathbb{P}_2$ is strictly ϵ -suboptimal in (10), whereas the weak inequality follows from the definition of $\hat{c}_r$ and the fact that $I(\mathbb{P}'_0, \mathbb{P}_2) = r_2 < r$.

It remains to be shown that the prediction disappointment $\mathbb{P}_2^\infty(c(\hat{x}(\hat{\mathbb{P}}_T), \mathbb{P}_2) > \hat{c}(\hat{x}(\hat{\mathbb{P}}_T), \hat{\mathbb{P}}_T))$ under model $\mathbb{P}_2$ decays at a rate of at most $r_2 < r$ as the sample size T tends to infinity. To this end, we define the set of disappointing estimator realizations as

$$\mathcal{D}(\mathbb{P}_2) = \{\mathbb{P}' \in \mathcal{P} : c(\hat{x}(\mathbb{P}'), \mathbb{P}_2) > \hat{c}(\hat{x}(\mathbb{P}'), \mathbb{P}')\}.$$

This set contains $\mathbb{P}'_0$ due to the strict inequality in (20). Recall now that $\hat{x}$ is continuous at $\mathbb{P}' = \mathbb{P}'_0$ due to our choice of $\mathbb{P}'_0$. As the predictors $\hat{c}$ and $\hat{c}_r$ are both continuous on their entire domain, the compositions $\hat{c}(\hat{x}(\mathbb{P}'), \mathbb{P}')$ and $c(\hat{x}(\mathbb{P}'), \mathbb{P}_2)$ are both continuous at $\mathbb{P}' = \mathbb{P}'_0$. This implies that $\mathbb{P}'_0$ belongs actually to the interior of $\mathcal{D}(\mathbb{P}_2)$. Thus, we find

$$\inf_{\mathbb{P}' \in \text{int}(\mathcal{D}(\mathbb{P}_2))} I(\mathbb{P}', \mathbb{P}_2) \leq I(\mathbb{P}'_0, \mathbb{P}_2) = r_2,$$

where the last equality follows from the definition of r_2 . As the empirical distributions $\{\hat{\mathbb{P}}_T\}_{T \in \mathbb{N}}$ obey the LDP lower bound (7b) under $\mathbb{P}_2 > 0$, we finally conclude that

$$\begin{aligned} -r < -r_2 &\leq - \inf_{\mathbb{P}' \in \text{int}(\mathcal{D}(\mathbb{P}_2))} I(\mathbb{P}', \mathbb{P}_2) \\ &\leq \liminf_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}_2^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}(\mathbb{P}_2)). \end{aligned}$$

The previous chain of inequalities implies, however, that $(\hat{c}, \hat{x})$ is infeasible in problem (6). This contradicts our initial assumption, and thus, $(\hat{c}_r, \hat{x}_r)$ must indeed be a strong solution of (6). $\square$

All guarantees discussed so far are asymptotic in nature. As in the case of the predictor $\hat{c}_r$, however, the prescriptor $\hat{x}_r$ can also be shown to satisfy finite sample guarantees.

Theorem 8 (Finite Sample Guarantee). *The out-of-sample disappointment of the distributionally robust prescriptor $\hat{x}_r$ enjoys the following finite sample guarantee under any model $\mathbb{P} \in \mathcal{P}$:*

$$\mathbb{P}^\infty(c(\hat{x}_r(\hat{\mathbb{P}}_T), \mathbb{P}) > \hat{c}_r(\hat{x}_r(\hat{\mathbb{P}}_T), \hat{\mathbb{P}}_T)) \leq (T+1)^d e^{-rT} \quad \forall T \in \mathbb{N}. \quad (21)$$

Proof. The proof of this result parallels those of Theorems 3 and 6 but uses the strong LDP upper bound (9) in lieu of the weak upper bound (7a). Details are omitted for brevity. $\square$

We stress that the finite sample guarantees of Theorems 5 and 8 as well as the strong optimality properties portrayed in Theorems 4 and 7 are independent of a particular data set. They guarantee that $\hat{c}_r$ and $\hat{x}_r$ provide trustworthy predictions and prescriptions, respectively, before the data are revealed.

Remark 4 (Optimal Hypothesis Testing). Bertsimas et al. (2018b) propose to construct predictors and prescriptors from statistical hypothesis tests. A hypothesis test uses i.i.d. samples $\xi_1, \dots, \xi_T$ drawn from the unknown true distribution $\mathbb{P}^*$ to decide whether the null hypothesis $\mathbb{P}^* = \mathbb{P}$ is false for a fixed model $\mathbb{P} \in \mathcal{P}$. Specifically, the null hypothesis is rejected (it is declared that $\mathbb{P}^* \neq \mathbb{P}$) if the empirical distribution $\hat{\mathbb{P}}_T$ associated with the observed sample path falls outside of a (measurable) acceptance region $A_T(\mathbb{P}) \subseteq \mathcal{P}$, which depends on the conjectured model $\mathbb{P}$ and the sample size T . Otherwise, it is deemed that there is insufficient data to reject the null hypothesis.

Bertsimas et al. (2018b) associate with each hypothesis test a predictor,

$$\hat{c}(x, \mathbb{P}') = \begin{cases} \sup & c(x, \mathbb{P}) \\ \text{s.t.} & \mathbb{P} \in \mathcal{P} \\ & \mathbb{P}' \in A_T(\mathbb{P}), \end{cases} \quad (22)$$

which evaluates the worst-case expected cost across all models $\mathbb{P} \in \mathcal{P}$ that pass the hypothesis test in view of the realization $\mathbb{P}' \in \mathcal{P}_T = \mathcal{P} \cap \{0, 1/T, \dots, (T-1)/T, 1\}^d$ of the empirical distribution $\hat{\mathbb{P}}_T$.

The quality of a hypothesis test is usually measured by its type I error $\mathbb{P}^\infty(\hat{\mathbb{P}}_T \notin A_T(\mathbb{P}))$, that is, the probability of falsely rejecting the null hypothesis, as well as

its type II error $\mathbb{Q}^\infty(\hat{\mathbb{P}}_T \in A_T(\mathbb{P}))$, that is, the probability of falsely accepting the null hypothesis if the data follows a distribution $\mathbb{Q} \neq \mathbb{P}$. A particularly popular test is the likelihood ratio test, which uses the acceptance region,

$$A_T^*(\mathbb{P}) = \left\{ \mathbb{P}' \in \mathcal{P}_T : \mathbb{P}^\infty(\hat{\mathbb{P}}_T = \mathbb{P}') / \sup_{\mathbb{Q} \neq \mathbb{P}} \mathbb{Q}^\infty(\hat{\mathbb{P}}_T = \mathbb{P}') \geq e^{-rT} \right\}.$$

Zeitouni et al. (1992) prove that the likelihood ratio test is optimal in the following sense. Among all hypothesis tests whose type I error decays at a rate of at least r , $\limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(\hat{\mathbb{P}}_T \notin A_T(\mathbb{P})) \leq -r$, the likelihood ratio test minimizes the negative decay rate of the type II error $\limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{Q}^\infty(\hat{\mathbb{P}}_T \in A_T(\mathbb{P}))$ simultaneously for all models $\mathbb{Q} \in \mathcal{P}$ with $\mathbb{Q} \neq \mathbb{P}$. Cover and Thomas (2006, theorem 11.1.2) further establish that the likelihood ratio of an estimator realization $\mathbb{P}' \in \mathcal{P}_T$ under two alternative distributions $\mathbb{Q}$ and $\mathbb{P}$ satisfies $\log(\mathbb{P}^\infty(\hat{\mathbb{P}}_T = \mathbb{P}') / \mathbb{Q}^\infty(\hat{\mathbb{P}}_T = \mathbb{P}')) = -T(I(\mathbb{P}', \mathbb{P}) - I(\mathbb{P}', \mathbb{Q}))$. The acceptance region of the likelihood ratio test thus simplifies to

$$\begin{aligned} A_T^*(\mathbb{P}) &= \left\{ \mathbb{P}' \in \mathcal{P}_T : I(\mathbb{P}', \mathbb{P}) \leq r + \inf_{\mathbb{Q} \neq \mathbb{P}} I(\mathbb{P}', \mathbb{Q}) \right\} \\ &= \{ \mathbb{P}' \in \mathcal{P}_T : I(\mathbb{P}', \mathbb{P}) \leq r \}, \end{aligned}$$

where the equality holds because $\inf_{\mathbb{Q} \neq \mathbb{P}} I(\mathbb{P}', \mathbb{Q}) = 0$. Hence, the distributionally robust predictor $\hat{c}_r$ that is strongly optimal in the meta-optimization problem (5) coincides with the hypothesis test-based predictor (22) corresponding to the likelihood ratio test.

5. Extension to Continuous State Spaces

Assume now that the realizations of the random parameter ξ may range over an arbitrary compact set $\Xi \subseteq \mathbb{R}^d$ that is not necessarily finite. In analogy to the discrete case, we denote by $\mathcal{P}$ the family of all Borel probability distributions supported on Ξ . Note that $\mathcal{P}$ is now a convex subset of an infinite-dimensional space, which significantly complicates the problem of finding optimal predictors and prescriptors. We equip $\mathcal{P}$ with the standard topology of weak convergence of distributions, recalling that the weak topology is metrized by the Prokhorov metric (Prokhorov 1956). Consequently, we equip $X \times \mathcal{P}$ with the product of the standard Euclidean topology on X and the weak topology on $\mathcal{P}$. In the remainder of this section, we analyze to what extent—and under what additional conditions—the results for finite state spaces carry over to the more general continuous case. As this analysis requires more subtle mathematical techniques, we relegate all proofs to Online Appendix A.

We first note that the definitions of model-based predictors and prescriptors require no changes. To evaluate the expectation in the definition of the model-based predictor $c(x, \mathbb{P}) = \int_{\Xi} \gamma(x, \xi) d\mathbb{P}(\xi)$, however, we now need to evaluate an integral with respect to $\mathbb{P}$ instead of a finite sum. Throughout this section, we assume that the cost function $\gamma(x, \xi)$ is jointly continuous in x and ξ . This implies via the compactness of X and Ξ that $c(x, \mathbb{P})$ is continuous in x and $\mathbb{P}$, which in turn guarantees that a model-based prescriptor $x^*(\mathbb{P}) \in \arg \min_{x \in X} c(x, \mathbb{P})$ exists for every $\mathbb{P} \in \mathcal{P}$.

Lemma 1 (Continuity of Model-Based Predictors). *If $\gamma(x, \xi)$ is continuous on the compact set $X \times \Xi$, then $c(x, \mathbb{P})$ is continuous on $X \times \mathcal{P}$.*

As in the case of a discrete state space, we study data-driven predictors and prescriptors that depend on the training data $\{\xi_t\}_{t=1}^T$ only through the empirical distribution. Because Ξ may now have infinite cardinality, we redefine the empirical distribution as $\hat{\mathbb{P}}_T = \frac{1}{T} \sum_{t=1}^T \delta_{\xi_t}$, where δ_{ξ_t} denotes the Dirac point mass at ξ_t . Using this new definition of $\hat{\mathbb{P}}_T$, we then define data-driven predictors and prescriptors exactly as in Section 2.1. As Ξ is compact, one can show that $\mathcal{P}$ is compact in the weak topology (Prokhorov 1956). Moreover, as the weak topology is metrized by the Prokhorov metric, $\mathcal{P}$ constitutes a (locally) compact metric space. The Baire category theorem thus implies that $\mathcal{P}$ is a Baire space (Baire 1899). Corollary 4 in Matejdes (1987), which applies because $\mathcal{P}$ is a Baire space and X is a metric space, further ensures that for any valid (continuous) predictor $\hat{c}$, the set-valued mapping $\arg \min_{x \in X} \hat{c}(x, \mathbb{P}')$ admits a quasi-continuous selector $\hat{x}$, which serves as a valid data-driven prescriptor. Using the exact same reasoning as in Section 2.1, one can show that the points of continuity of any quasi-continuous prescriptor are dense in $\mathcal{P}$.

The best predictors and predictor-prescriptor pairs can again be found by solving the meta-optimization problems (5) and (6), respectively. To construct near-optimal solutions for these meta-optimization problems, we recall the definition of the relative entropy between arbitrary distributions $\mathbb{P}'$ and $\mathbb{P}$ on a compact set $\Xi \subseteq \mathbb{R}^d$.

Definition 8 (Generalized Relative Entropy). The relative entropy of $\mathbb{P}' \in \mathcal{P}$ with respect to $\mathbb{P} \in \mathcal{P}$ is defined as

$$I(\mathbb{P}', \mathbb{P}) = \begin{cases} \int_{\Xi} \log(d\mathbb{P}'/d\mathbb{P}(\xi)) d\mathbb{P}'(\xi) & \text{if } \mathbb{P}' \ll \mathbb{P}, \\ +\infty & \text{otherwise,} \end{cases}$$

where $\mathbb{P}' \ll \mathbb{P}$ means that $\mathbb{P}'$ is absolutely continuous with respect to $\mathbb{P}$, whereas $d\mathbb{P}'/d\mathbb{P}(\xi)$ denotes the Radon-Nikodym derivative of $\mathbb{P}'$ with respect to $\mathbb{P}$, which exists if $\mathbb{P}' \ll \mathbb{P}$ (Nikodym 1930).

The properties of the relative entropy portrayed in Proposition 1 hold verbatim in the more general setting considered here (Van Erven and Harremoës 2014). Using the generalized definition of the relative entropy, the distributionally robust predictor $\hat{c}_r$ and the corresponding prescriptor $\hat{x}_r$ can be constructed as in Definitions 6 and 7, respectively. In the following, we will show that the predictor $\hat{c}_r$ is continuous, which ensures that the prescriptor $\hat{x}_r$ can always be chosen to be quasi-continuous. To this end, we first derive a dual representation for $\hat{c}_r$.

Proposition 5 (Dual Representation Revisited). *If $r > 0$ and $\bar{\gamma}(x) = \max_{\xi \in \Xi} \gamma(x, \xi)$ is the worst-case cost function, then the distributionally robust predictor admits the dual representation,*

$$\hat{c}_r(x, \mathbb{P}') = \min_{\alpha \geq \bar{\gamma}(x)} \alpha - e^{-r} \cdot \exp\left(\int_{\Xi} \log(\alpha - \gamma(x, \xi)) d\mathbb{P}'(\xi)\right). \quad (23)$$

Problem (23) has a minimizer $\alpha^* \leq \frac{\bar{\gamma}(x) - e^{-r} c(x, \mathbb{P}')}{1 - e^{-r}}$.

Proposition 5 extends Proposition 2 to compact continuous state spaces and suggests that $\hat{c}_r(x, \mathbb{P}')$ can be computed via bisection or other line search methods. Thus, the computational tractability of problem (23) largely hinges on our ability to efficiently evaluate the geometric mean $\exp(\int_{\Xi} \log(\alpha - \gamma(x, \xi)) d\mathbb{P}'(\xi))$ for any fixed α . For example, if $\mathbb{P}'$ coincides with (a realization of) the empirical distribution $\hat{\mathbb{P}}_T$, we recover the geometric mean of $\alpha - \gamma(x, \xi)$ along a sample path, which can be reformulated as the optimal value of a tractable second-order cone program involving $\mathcal{O}(T)$ constraints and auxiliary variables (Nesterov and Nemirovskii 1994, section 6.2.3.5):

$$\exp\left(\int_{\Xi} \log(\alpha - \gamma(x, \xi)) d\mathbb{P}'(\xi)\right) = \left(\prod_{t=1}^T (\alpha - \gamma(x, \xi_t))\right)^{1/T}.$$

To our best knowledge, the dual representation (23) is new. The closest result we are aware of is the dual representation of the negative log-entropic risk measure derived in Ahmadi-Javid (2012, theorem 5.1). Indeed, the negative log-entropic risk of $\gamma(x, \xi)$ coincides with the restricted distributionally robust predictor $\bar{c}_r(x, \mathbb{P}')$. Recall from (13) that $\bar{c}_r(x, \mathbb{P}')$ differs from $c_r(x, \mathbb{P}')$ only in that it imposes the additional constraint $\mathbb{P} \ll \mathbb{P}'$ when evaluating the worst-case expected cost. Using theorem 5.1 by Ahmadi-Javid (2012), one can thus show that the dual representation of $\bar{c}_r(x, \mathbb{P}')$ differs from (23) only in that it replaces $\bar{\gamma}(x)$ with $\inf\{\bar{\gamma} : \mathbb{P}[\gamma(x, \xi) \leq \bar{\gamma}] = 1\} \leq \bar{\gamma}(x)$. Maybe surprisingly, however, the derivation of (23) provided here is substantially more challenging.

Proposition 6 (Continuity of $\hat{c}_r$, Revisited). *If $r \geq 0$, then the distributionally robust predictor $\hat{c}_r$ is continuous on $X \times \mathcal{P}$.*

Proposition 6 ensures that $\hat{c}_r \in \mathcal{C}$. As any continuous predictor induces a quasi-continuous prescriptor, we may thus conclude that there exists a valid distributionally robust prescriptor $\hat{x}_r$ such that $(\hat{c}_r, \hat{x}_r) \in \mathcal{X}$. It now only remains to establish that these predictors and predictor-prescriptor pairs are the unique strong solutions of the meta-optimization problems (5) and (6), respectively. In Section 4, this was achieved by leveraging the weak LDP portrayed in Theorem 1. Luckily, this LDP carries over to the more general setting considered here—albeit with a subtle difference.

Theorem 9 (Weak LDP Revisited). *If the samples $\{\xi_t\}_{t \in \mathbb{N}}$ are drawn independently from some $\mathbb{P} \in \mathcal{P}$, then for every set $\mathcal{D} \subseteq \mathcal{P}$ the sequence of empirical distributions $\{\hat{\mathbb{P}}_T\}_{T \in \mathbb{N}}$ satisfies*

$$\limsup_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}) \leq - \inf_{\mathbb{P}' \in \text{cl}(\mathcal{D})} I(\mathbb{P}', \mathbb{P}), \quad (24a)$$

$$\liminf_{T \rightarrow \infty} \frac{1}{T} \log \mathbb{P}^\infty(\hat{\mathbb{P}}_T \in \mathcal{D}) \geq - \inf_{\mathbb{P}' \in \text{int}(\mathcal{D})} I(\mathbb{P}', \mathbb{P}). \quad (24b)$$

Proof. See Csiszár (2006, section 2). $\square$

Formally, Theorem 9 is almost identical to Theorem 1. However, the weak LDP upper bound (24a) differs from (7a) in that the minimization over all estimator realizations $\mathbb{P}'$ on the right-hand side runs over the closure of $\mathcal{D}$. This subtle difference invalidates the proof of Theorem 3, and thus we need a new approach to show that $\hat{c}_r$ is feasible in (5). Moreover, the weak LDP lower bound (24b) does not rely on any structural assumptions about $\mathbb{P}$. Note that the condition $\mathbb{P} > 0$ in Theorem 1 was only imposed for convenience to simplify the proof of (7b) in Online Appendix A.

As in the case of finite state spaces, one can now show that the distributionally robust predictor $\hat{c}_r$ is the unique strong solution of the meta-optimization problem (5).

Theorem 10 (Feasibility and Optimality of $\hat{c}_r$, Revisited). *If $r \geq 0$, then the predictor $\hat{c}_r$ is feasible in (5). Moreover, if $r > 0$, then $\hat{c}_r$ is strongly optimal in (5).*

Although we did not manage to prove that $(\hat{c}_r, \hat{x}_r)$ is feasible in the meta-optimization problem (6), we still could show that it is essentially feasible and strongly optimal in a precise sense.

Theorem 11 (Feasibility and Optimality of $(\hat{c}_r, \hat{x}_r)$, Revisited). *If $r \geq 0$, then the shifted predictor-prescriptor pair $(\hat{c}_r + \epsilon, \hat{x}_r)$ is feasible in (6) for every $\epsilon > 0$. Moreover, if $r > 0$, then $(\hat{c}_r, \hat{x}_r)$ is preferred to every feasible solution of (6)—even though it may be infeasible.*

Theorem 11 asserts that $(\hat{c}_r, \hat{x}_r)$ is less conservative than any predictor-prescriptor pair feasible in the

meta-optimization problem (6) and that $(\hat{c}_r, \hat{x}_r)$ can be made feasible in (6) by shifting the distributionally robust predictor $\hat{c}_r$ up by just a tiny amount. For practical purposes, this means that $(\hat{c}_r, \hat{x}_r)$ is indeed essentially optimal. Whether $(\hat{c}_r, \hat{x}_r)$ itself is feasible in (6) remains open.

We also emphasize that the strong LDP portrayed in Theorem 2 has no continuous counterpart, which implies that the finite sample guarantees of Theorems 5 and 8 cannot be generalized.

Endnote

¹ Here, the interior of $\mathcal{D}$ is taken with respect to the subspace topology on $\mathcal{P}$. Recall that a set $\mathcal{D} \subseteq \mathcal{P}$ is open in the subspace topology on $\mathcal{P}$ if $\mathcal{D} = \mathcal{P} \cap \mathcal{O}$ for some set $\mathcal{O} \subseteq \mathbb{R}^d$ that is open in the Euclidean topology on $\mathbb{R}^d$.

References

- Ahmadi-Javid A (2012) Entropic value-at-risk: A new coherent risk measure. *J. Optim. Theory Appl.* 155(3):1105–1123.
- Baire R (1899) Sur les fonctions de variables réelles. *Ann. Mat. Pura Appl.* 3(3):1–123.
- Bayraksan G, Love DK (2015) Data-driven stochastic programming using phi-divergences. Aleman DM, Thiele AC, eds. *The Operations Research Revolution* (INFORMS, Catonsville, MD), 1–19.
- Ben-Tal A, Den Hertog D, De Waegenaere A, Melenberg B, Rennen G (2013) Robust solutions of optimization problems affected by uncertain probabilities. *Management Sci.* 59(2):341–357.
- Berge C (1963) *Topological Spaces: Including a Treatment of Multi-Valued Functions, Vector Spaces, and Convexity* (Oliver & Boyd, London).
- Bertsimas D, Gupta V, Kallus N (2018a) Robust sample average approximation. *Math. Program.* 171(1):217–282.
- Bertsimas D, Gupta V, Kallus N (2018b) Data-driven robust optimization. *Math. Programming* 167(2):235–292.
- Bledsoe W (1952) Neighborly functions. Hedlund GA, Hochschild GP, Schaeffer AC, eds. *Proceedings of the American Mathematical Society*, vol. 3 (American Mathematical Society, Menasha, WI), 114–115.
- Calafiore GC (2007) Ambiguous risk measures and optimal robust portfolios. *SIAM J. Optim.* 18(3):853–877.
- Cover TM, Thomas JA (2006) *Elements of Information Theory* (John Wiley & Sons, Hoboken, NJ).
- Csiszár I (2006) A simple proof of Sanov's theorem. *Bull. Brazilian Math. Soc.* 37(4):453–459.
- Delage E, Ye Y (2010) Distributionally robust optimization under moment uncertainty with application to data-driven problems. *Oper. Res.* 58(3):595–612.
- Duchi J, Glynn P, Namkoong H (2016) Statistics of robust optimization: A generalized empirical likelihood approach. Preprint, submitted October 11, <https://arxiv.org/abs/1610.03425>.
- Dupačová J, Wets R (1988) Asymptotic behavior of statistical estimators and of optimal solutions of stochastic optimization problems. *Ann. Statist.* 16(4):1517–1549.
- Erdoğan E, Iyengar G (2006) Ambiguous chance constrained problems and robust optimization. *Math. Programming* 107(1–2): 37–61.
- Gupta V (2019) Near-optimal Bayesian ambiguity sets for distributionally robust optimization. *Management Sci.* 65(9):4242–4260.
- Hu Z, Hong LJ (2013) Kullback-Leibler divergence constrained distributionally robust optimization. Accessed November 11, 2020, http://www.optimization-online.org/DB_FILE/2012/11/3677.pdf.
- Jiang R, Guan Y (2018) Risk-averse two-stage stochastic program with distributional ambiguity. *Oper. Res.* 66(5):1390–1405.

- Kullback S, Leibler RA (1951) On information and sufficiency. *Ann. Math. Statist.* 22(1):79–86.
- Lam H (2016) Robust sensitivity analysis for stochastic systems. *Math. Oper. Res.* 41(4):1248–1275.
- Lam H (2019) Recovering best statistical guarantees via the empirical divergence-based distributionally robust optimization. *Oper. Res.* 67(4):1090–1105.
- Le Maître OP, Knio OM (2010) Introduction: Uncertainty quantification and propagation. Le Maître OP, Knio OM, eds. *Spectral Methods for Uncertainty Quantification* (Springer, Dordrecht, Netherlands).
- Matejdes M (1987) Sur les sélecteurs des multifonctions. *Math. Slovaca* 37(1):111–124.
- Michaud RO (1989) The Markowitz optimization enigma: Is ‘optimized’ optimal? *Financial Analysts J.* 45(1):31–42.
- Mohajerin Esfahani P, Kuhn D (2018) Data-driven distributionally robust optimization using the Wasserstein metric: Performance guarantees and tractable reformulations. *Math. Programming* 171(1–2):115–166.
- Nesterov Y, Nemirovskii A (1994) *Interior-Point Polynomial Algorithms in Convex Programming* (SIAM, Philadelphia).
- Nikodym O (1930) Sur une généralisation des intégrales de M.J. Radon. *Fundamenta Math.* 15(1):131–179.
- Owen AB (1988) Empirical likelihood ratio confidence intervals for a single functional. *Biometrika* 75(2):237–249.
- Parpas P, Ustun B, Webster M, Tran Q (2015) Importance sampling in stochastic programming: A Markov chain Monte Carlo approach. *INFORMS J. Comput.* 27(2):358–377.
- Pflug G, Wozabal D (2007) Ambiguity in portfolio selection. *Quant. Finance* 7(4):435–442.
- Prokhorov YV (1956) Convergence of random processes and limit theorems in probability theory. *Theory Probab. Appl.* 1(2):157–214.
- Rockafellar RT, Wets RJ-B (1998) *Variational Analysis* (Springer, Dordrecht, Netherlands).
- Shapiro A, Dentcheva D, Ruszczyński A (2014) *Lectures on Stochastic Programming: Modeling and Theory* (SIAM, Philadelphia).
- Smith JE, Winkler RL (2006) The optimizer’s curse: Skepticism and postdecision surprise in decision analysis. *Management Sci.* 52(3):311–322.
- Sun H, Xu H (2016) Convergence analysis for distributionally robust optimization and equilibrium problems. *Math. Oper. Res.* 41(2):377–401.
- Van Erven T, Harremoës P (2014) Rényi divergence and Kullback-Leibler divergence. *IEEE Trans. Inform. Theory* 60(7):3797–3820.
- Wang Z, Glynn PW, Ye Y (2016) Likelihood robust optimization for data-driven newsvendor problems. *Comput. Management Sci.* 12(2):241–261.
- Zeitouni O, Ziv J, Merhav N (1992) When is the generalized likelihood ratio test optimal? *IEEE Trans. Inform. Theory* 38(5):1597–1602.
- Zhao C, Guan Y (2018) Data-driven risk-averse stochastic optimization with Wasserstein metric. *Oper. Res. Lett.* 46(2):262–267.