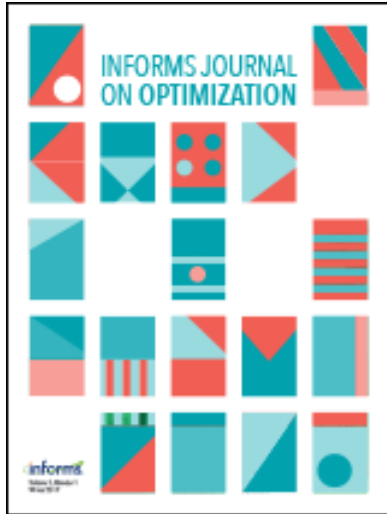




INFORMS Journal on Optimization

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

Computational Aspects of Bayesian Solution Estimators in Stochastic Optimization

Danial Davarnia, Burak Kocuk, Gérard Cornuéjols

To cite this article:

Danial Davarnia, Burak Kocuk, Gérard Cornuéjols (2020) Computational Aspects of Bayesian Solution Estimators in Stochastic Optimization. INFORMS Journal on Optimization 2(4):256-272. <https://doi.org/10.1287/ijoo.2019.0035>

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>

Computational Aspects of Bayesian Solution Estimators in Stochastic Optimization

Danial Davarnia,^a Burak Kocuk,^b Gérard Cornuéjols^c

^a Department of Industrial and Manufacturing Systems Engineering, Iowa State University, Ames, Iowa 50011; ^b Industrial Engineering Program, Sabancı University, Istanbul 34956, Turkey; ^c Tepper School of Business, Carnegie Mellon University, Pittsburgh, Pennsylvania 15213

Contact: davarnia@iastate.edu,  <https://orcid.org/0000-0002-8602-4454> (DD); burakkocuk@sabanciuniv.edu,  <https://orcid.org/0000-0002-4218-1116> (BK); gc0v@andrew.cmu.edu (GC)

Received: December 23, 2018

Revised: September 29, 2019

Accepted: April 23, 2020

Published Online in Articles in Advance:
 October 5, 2020

<https://doi.org/10.1287/ijoo.2019.0035>

Copyright: © 2020 INFORMS

Abstract. We study a class of stochastic programs in which some of the elements in the objective function are random and their probability distribution has unknown parameters. The goal is to find a good estimate for the optimal solution of the stochastic program using data sampled from the distribution of the random elements. We investigate two common optimization criteria for evaluating the quality of a solution estimator—one based on the difference in objective values and the other based on the Euclidean distance between solutions. We use *risk* as the expected value of such criteria over the sample space. Under a Bayesian framework, where a prior distribution is assumed for the unknown parameters, two natural estimation-optimization strategies arise. A *separate* scheme first finds an estimator for the unknown parameters and then uses this estimator in the optimization problem. A *joint* scheme combines the estimation and optimization steps by directly adjusting the distribution in the stochastic program. We analyze the risk difference between the solutions obtained from these two schemes for several classes of stochastic programs and provide insight on the computational effort to solve these problems.

Funding: This work was supported by the National Science Foundation, Division of Civil, Mechanical and Manufacturing Innovation [Grant CMMI-1560828] and the Office of Naval Research [Grant 000141812129].

Supplemental Material: The e-companion is available at <https://doi.org/10.1287/ijoo.2019.0035>.

Keywords: stochastic optimization • Bayesian inference • statistical estimation • solution estimators

1. Introduction

The interaction between data and optimization can be complex and sometimes counterintuitive. As an example, consider portfolio optimization: Construct a portfolio of assets that maximizes expected return over some future period, subject to a constraint on risk (Markowitz 1959). A major issue in practice is the estimation of the expected returns of the individual assets. The theory of efficient markets tells us that the latest data contain the best available information. Surprisingly, Jorion (1986) recommends not to use the data directly to estimate the expected asset returns but instead to shrink the data on individual assets toward a “grand average” involving all the other assets as well. This seems counterintuitive, but it works well in practice. This paradox can be traced back to the work of Stein (1956). Such results are best understood under a Bayesian framework. In this paper, we focus on the computational implications of parameter estimation in the context of stochastic optimization.

Consider the following stochastic program:

$$\max_{x \in \mathcal{X}} \mathbb{E}_{\xi|\theta} [f(\xi, x)]. \quad (1)$$

In (1), the vector $\xi \in \mathbb{R}^n$ represents random variables with joint probability density function (pdf) $g(\xi|\theta)$, where θ denotes a vector of parameters. Vector $x \in \mathbb{R}^m$ represents decision variables that belong to a closed set $\mathcal{X} \subseteq \mathbb{R}^m$. Function $f(\xi, x) : \mathbb{R}^{n+m} \rightarrow \mathbb{R}$ contains both random variables and decision variables, and its expectation with respect to the distribution $g(\xi|\theta)$ forms the objective function of the stochastic program. We assume that this expectation is finite for all $x \in \mathcal{X}$. In our introductory portfolio example, the asset returns ξ could have a multivariate normal distribution g with mean θ and known covariance matrix, the portfolio weights form the decision vector x , and the portfolio return is $f(\xi, x) = \xi^\top x$.

If the parameters θ are known, then problem (1) reduces to a classical stochastic program with an optimal solution $x^*(\theta)$ and optimal value $\mathbb{E}_{\xi|\theta}[f(\xi, x^*(\theta))]$. In *parametric* stochastic programs, it is assumed that θ is not fully known. In such situations, $T \geq 1$ independent observations of the random variables ξ drawn from the distribution $g(\xi|\theta)$ are used to estimate the unknown parameters. The question is how to use these data to find a *good* estimator for the true optimal solution $x^*(\theta)$ of (1).

Parametric stochastic programs encompass numerous applications; see Birge and Louveaux (2011) for applications in optimization and Donti et al. (2017) for an application in machine learning. Various solution techniques have been proposed in the literature. Examples are Monte Carlo, bootstrapping, and scenario-reduction techniques such as sample-average approximation (SAA); see Shapiro (2003). Jiang and Shanbhag (2016) consider a setup where the parameter may be obtained through the solution of a learning problem. See also Wang et al. (2017).

One of the main tools in studying parametric stochastic problems is point estimation, which originates in statistical inference and decision theory. We refer the interested reader to Ferguson (1967) or Lehmann and Casella (1998) and many references therein. Estimation plays a central role in constructing optimization models that involve uncertainty. For instance, the theory of minimaxity in estimation is closely related to robust optimization. We refer the reader to Ben-Tal et al. (2009) and Bertsimas et al. (2011), for an introduction to the theory and applications of robust optimization, and to Lim et al. (2006) and references therein for a detailed account of the connection between robust optimization and minimaxity.

In this paper, we focus on stochastic optimization and its connection to statistical estimation theory under a Bayesian setting, where a prior distribution is assumed for the unknown parameters; see Kadane (2011) for an overview of basics in the Bayesian framework. This is a common framework when modeling parametric stochastic programs; see Lim et al. (2006). For the portfolio optimization problem introduced earlier, this setting implies that the returns have a certain distribution, whose mean is random and follows a prior distribution. There are two natural schemes to find a solution estimator in this setting. In the first scheme, the parameter estimation is performed separately, and its estimator is used as an input for the optimization problem. In the second scheme, the estimation process is incorporated within the optimization step to obtain a solution.

Although several bodies of work provide comparisons between the “true” stochastic optimization problem and its approximate “deterministic” counterpart for classical stochastic programs (see Kleywegt and Shapiro 2004), sometimes framed under the “separate” versus “joint” estimation-optimization framework (see, e.g., Thomas et al. 2006, Levine et al. 2016, Donti et al. 2017), a detailed analysis of this nature is lacking in parametric stochastic optimization under the Bayesian framework. We provide such an analysis for the above schemes from two angles: the quality of the solution estimators obtained by each scheme and the computational efficiency of performing them. In particular, we show how different problem structures and probability distributions affect the quality and efficiency of the separate and joint schemes, by providing theoretical and computational results.

Our study is organized as follows. In Section 2, we formally define the key problem settings: risk criteria and the separate and joint estimation-optimization schemes. In Section 3, we illustrate the trade-off between accuracy and computational efficiency by presenting experiments on two applications: a portfolio optimization problem and a stochastic geometric program. We observe that the separate and joint schemes can lead to problems with different computational complexity and solution stability. In Section 4, we identify conditions under which the two schemes yield the same solutions and provide examples with an arbitrarily large risk difference when these conditions do not hold. In Section 5, we study a class of stochastic piecewise-linear programs under our two schemes. This class of stochastic programs generalizes the newsvendor and median problems and models applications in education, psychology, and artificial intelligence. We derive explicit bounds on the risk difference between the solution estimators of the separate and joint schemes for various probability distributions. Through a higher dimensional extension, this class also describes the “rectifier” used as an activation function for deep neural networks in machine learning. We conclude the section by performing computational experiments on a core structure of such models. Section 6 contains concluding remarks.

2. Loss Function and Risk

Consider the stochastic program (1). The form of the distribution $g(\xi|\theta)$ is known, but the parameters θ are unknown. However, we have $T \geq 1$ observations sampled from this distribution. To simplify notation, we will assume a single vector ξ of observations; that is, $T = 1$. This is done without loss of generality, since in our

Bayesian context the samples are used to derive posterior distributions, which are readily obtainable for any size T .

Since the true optimal solution $x^*(\theta)$ is unknown (as θ is unknown), we estimate it with a *solution estimator* $\hat{x}(\bar{\xi}) \in \mathcal{X}$ as a function of the observation $\bar{\xi}$. There are many choices for a solution estimator; hence, we need a suitable criterion to evaluate its performance. We use *risk*, a popular criterion in statistical inference. To this end, we define two natural *loss* functions, one based on the difference in objective values and the other based on the distance between solutions.

2.1. Loss as the Difference in Objective Values

The *loss function* $\mathcal{L}^L$ is defined as the difference between the objective function values of the solution estimator $\hat{x}(\bar{\xi})$ and the true optimal solution $x^*(\theta)$,

$$\mathcal{L}^L(x^*(\theta), \hat{x}(\bar{\xi})) = \mathbb{E}_{\xi|\theta}[f(\xi, x^*(\theta))] - \mathbb{E}_{\xi|\theta}[f(\xi, \hat{x}(\bar{\xi}))], \quad (2)$$

where the expectations are taken with respect to the distribution $g(\xi|\theta)$ of ξ given the true θ . The superscript represents that the loss is *linear* in objective function values. Note that $\mathcal{L}^L(x^*(\theta), x^*(\theta)) = 0$ and that $\mathcal{L}^L(x^*(\theta), \hat{x}(\bar{\xi})) \geq 0$ for any $\hat{x}(\bar{\xi}) \in \mathcal{X}$ as it is feasible to (1), which has $x^*(\theta)$ as an optimal solution.

We evaluate the loss under a popular framework in the Bayesian school, where the unknown parameters θ are assumed to be random with a prior distribution $\pi(\theta)$. Since $\mathcal{L}^L(x^*(\theta), \hat{x}(\bar{\xi}))$ is a function of both the unknown parameters θ and the observation $\bar{\xi}$, it is a random quantity. To obtain a measure of overall performance, a *risk* is defined as

$$\mathcal{R}^L(x^*(\theta), \hat{x}(\bar{\xi})) = \mathbb{E}_{\theta} \mathbb{E}_{\xi|\theta}[\mathcal{L}^L(x^*(\theta), \hat{x}(\bar{\xi}))], \quad (3)$$

where the outer expectation is taken with respect to the prior distribution $\pi(\theta)$ of θ , and the inner expectation is taken with respect to the likelihood distribution $g(\xi|\theta)$.

A *best* solution estimator is defined as one that achieves minimum risk. It is called a *Bayes solution estimator*. The proofs of the following proposition and some others in this section follow from standard techniques in Bayesian analysis. We include these proofs in the electronic companion for the sake of completeness.

Proposition 1. Consider stochastic program (1). Assume that $\xi \sim g(\xi|\theta)$ and $\theta \sim \pi(\theta)$. Then, any solution estimator $\hat{x}^{J,L}(\bar{\xi})$ that solves

$$\max_{x \in \mathcal{X}} \mathbb{E}_{\theta|\bar{\xi}} \mathbb{E}_{\xi|\theta}[f(\xi, x)] \quad (4)$$

is a Bayes solution estimator under the loss function (2), where the inner expectation is taken with respect to the likelihood distribution $g(\xi|\theta)$, and the outer expectation is taken with respect to the posterior distribution $\Pi(\theta|\bar{\xi})$ of the parameters θ given the observation $\bar{\xi}$.

The Bayes solution estimator $\hat{x}^{J,L}(\bar{\xi})$ found in Proposition 1 minimizes risk for the loss function (2) in a joint estimation and optimization (Joint-EO) scheme. The superscripts in $\hat{x}^{J,L}(\bar{\xi})$ represent the *Joint* scheme with respect to the *Linear* loss (2).

We next give another representation of the Joint-EO problem (4) based on the expectation with respect to a specific marginal distribution called *posterior predictive*.

Definition 1. The *posterior predictive* distribution of ξ given $\bar{\xi}$ is defined as $h(\xi|\bar{\xi}) = \int_{\theta} g(\xi|\theta) \Pi(\theta|\bar{\xi}) d\theta$, which is obtained by marginalizing the distribution of ξ given θ over the posterior distribution of θ given observation $\bar{\xi}$.

Proposition 2. Assume that $\xi \sim g(\xi|\theta)$, $\theta \sim \pi(\theta)$, and that $\mathbb{E}_{\xi|\bar{\xi}}[|f(\xi, x)|]$ is finite for any $x \in \mathcal{X}$. Then, any solution estimator $\hat{x}^{J,L}(\bar{\xi})$ that solves

$$\max_{x \in \mathcal{X}} \mathbb{E}_{\xi|\bar{\xi}}[f(\xi, x)] \quad (5)$$

is a Bayes solution estimator under the loss function (2), where the expectation is taken with respect to the posterior predictive distribution $h(\xi|\bar{\xi})$.

Computing $\hat{x}^{J,L}(\bar{\xi})$ can be extremely challenging in practice. An alternative, which can be suboptimal in risk, is to apply a separate estimation and optimization (Separate-EO) scheme. First, a Bayes estimator $\hat{\theta}^B(\bar{\xi})$ for θ is computed under the squared error loss, then this estimator is used in place of θ in (1) to obtain an optimal

solution $\hat{x}^S(\bar{\xi})$ that serves as a solution estimator for the true optimal solution $x^*(\theta)$. The superscript in $\hat{x}^S(\bar{\xi})$ represents the *Separate* scheme.

It can be shown that the Bayes estimator of θ under the squared error loss $\mathcal{L}^Q(\theta, \hat{\theta}(\bar{\xi})) = \|\theta - \hat{\theta}(\bar{\xi})\|^2$ is $\hat{\theta}^B(\bar{\xi}) = \mathbb{E}_{\theta|\bar{\xi}}[\theta]$, where the expectation is taken with respect to the posterior distribution $\Pi(\theta|\bar{\xi})$. When the likelihood and prior belong to the *conjugate* family, the posterior mean is representable as a convex combination between the prior mean and the observation (maximum likelihood estimator for multiple observations); see Diaconis and Ylvisaker (1979). Such estimators are sometimes referred to as *shrinkage* estimators, as they shrink the maximum likelihood estimator toward another vector. Jorion (1986)'s shrinkage mentioned in the introduction can be viewed in this light. We refer the interested reader to Davarnia and Cornuéjols (2017) and Basu et al. (2019) for a non-Bayesian perspective.

Proposition 2 allows for a direct comparison between the Separate-EO and the Joint-EO problems. For the Separate-EO, we solve (1), where the expectation is taken with respect to the distribution $g(\xi|\hat{\theta}^B(\bar{\xi}))$. For the Joint-EO under the linear loss (2), we solve (5), where the expectation is taken with respect to the posterior predictive distribution $h(\xi|\bar{\xi})$. Therefore, the difference between these two estimation schemes stems from the difference between the distributions with respect to which the expectation of $f(\xi, x)$ is computed.

2.2. Loss as the Distance Between Solutions

Since an optimization problem can have multiple optimal solutions, we let $\mathcal{D}(W, v) = \min_{w \in W} \|w - v\|$, where $W \subseteq \mathbb{R}^n$ and $v \in \mathbb{R}^n$. We define the loss function

$$\mathcal{L}^Q(\mathcal{X}^*(\theta), \hat{x}(\bar{\xi})) = \mathcal{D}^2(\mathcal{X}^*(\theta), \hat{x}(\bar{\xi})), \quad (6)$$

where $\mathcal{X}^*(\theta)$ denotes the set of optimal solutions of (1) when θ is known. The superscript indicates that the loss is *quadratic* in solutions. We define risk similarly to (3) as

$$\mathcal{R}^Q(x^*(\theta), \hat{x}(\bar{\xi})) = \mathbb{E}_{\theta} \mathbb{E}_{\xi|\theta} [\mathcal{L}^Q(x^*(\theta), \hat{x}(\bar{\xi}))]. \quad (7)$$

Note that the Separate-EO method does not incorporate information about the optimization problem in the estimation step. Hence, the change of the loss function does not affect this method, and we obtain the same solution estimator $\hat{x}^S(\bar{\xi})$ as given in Section 2.1. For the Joint-EO method, however, the underlying optimization problem is different and therefore a different definition for solution estimators is necessary.

Proposition 3. Assume that $\xi \sim g(\xi|\theta)$ and $\theta \sim \pi(\theta)$. Then, any solution estimator $\hat{x}^{J,Q}(\bar{\xi})$ that solves

$$\min_{x \in \mathcal{X}} \mathbb{E}_{\theta|\bar{\xi}} [\mathcal{D}^2(x^*(\theta), x)] \quad (8)$$

is a Bayes solution estimator under the loss function (6), where the expectation is taken with respect to the posterior distribution $\Pi(\theta|\bar{\xi})$.

We note that the problem (8) is indeed a minimization problem as opposed to the problems (4) and (5).

The superscript in $\hat{x}^{J,Q}(\bar{\xi})$ indicates the *Joint* method under the *Quadratic* loss function (6). In the sequel, the expectation operator $\mathbb{E}[\cdot]$ applied on a vector implies component-wise expectation. The following result shows that $\hat{x}^{J,Q}(\bar{\xi})$ can be obtained as the projection of $\mathbb{E}_{\theta|\bar{\xi}}[x^*(\theta)]$ onto $\mathcal{X}$, when (1) has a unique optimal solution.

Corollary 1. Assume that (1) has a unique optimal solution $x^*(\theta)$ for any given θ . Then, the Bayes solution estimator of (8) under the loss function (6) is obtained as the optimal solution of

$$\min_{x \in \mathcal{X}} \|\mathbb{E}_{\theta|\bar{\xi}}[x^*(\theta)] - x\|^2. \quad (9)$$

Further, when $\mathcal{X}$ is convex, the Bayes solution estimator is $\hat{x}^{J,Q}(\bar{\xi}) = \mathbb{E}_{\theta|\bar{\xi}}[x^*(\theta)]$.

We note that computing $\hat{x}^{J,Q}(\bar{\xi})$ can be more challenging than $\hat{x}^{J,L}(\bar{\xi})$, as the former requires an explicit derivation of the optimal solution set $x^*(\theta)$, which can be very hard for general problem structures.

3. Trade-Off Between Accuracy and Computational Efficiency

As established in Section 2, the Joint-EO method yields the best solution estimator. On the other hand, the Separate-EO method may sometimes be computationally more efficient. In this section, we present two applications that illustrate contrasting points on this trade-off spectrum.

The first application is a portfolio construction model, and we show that the Joint-EO method dominates the Separate-EO method in both computation time and solution quality. For the second application, we investigate a stochastic geometric program. Here, the Joint-EO method is computationally expensive and does not yield stable solutions, whereas the Separate-EO method is tractable and exhibits a robust performance.

3.1. Application I: Portfolio Construction

Consider an investor with unit initial wealth. She can invest in a combination of n stocks whose random return is denoted by the vector ξ . We assume that random returns have a joint distribution $g(\xi|\mu)$ with unknown mean μ . We further assume that a prior distribution for the unknown parameters, denoted by $\pi(\mu)$, is available and a set of observations $\xi^1, \dots, \xi^T$ are drawn from the likelihood distribution.

We assume that the utility function of the investor is given as $u : [0, \infty) \rightarrow \mathbb{R}$, and the aim is to maximize the expected utility. Denoting the portfolio weights by x , we formulate this problem as

$$\max_{x \in \Delta^n} \mathbb{E}_{\xi|\mu} [u(1 + \xi^\top x)], \quad (10)$$

where $\Delta^n := \{x \in \mathbb{R}_+^n : \sum_{j=1}^n x_j = 1\}$. Recall that the Separate-EO estimator can be obtained as an optimal solution of

$$\max_{x \in \Delta^n} \mathbb{E}_{\xi|\mu^B(\bar{\xi})} [u(1 + \xi^\top x)], \quad (11)$$

where $\mu^B(\bar{\xi})$ is the Bayes estimator of μ , whereas the Joint-EO estimator under the linear loss can be computed as an optimal solution of

$$\max_{x \in \Delta^n} \mathbb{E}_{\xi|\bar{\xi}} [u(1 + \xi^\top x)]. \quad (12)$$

Now, consider a situation where the likelihood and prior distributions are multivariate Normal with known covariance matrices and prior mean; that is, $\xi|\mu \sim \mathcal{N}(\mu, \Sigma)$ and $\mu \sim \mathcal{N}(\mu_0, \Sigma_0)$. Then, we obtain the posterior and predictive distributions, respectively, as $\mu|\xi \sim \mathcal{N}(\mu'_0, \Sigma'_0)$ and $\xi|\xi \sim \mathcal{N}(\mu'_0, \Sigma'_0 + \Sigma)$, where ξ is the sample average, $\Sigma'_0 := (\Sigma_0^{-1} + m\Sigma^{-1})^{-1}$ and $\mu'_0 := \Sigma'_0(\Sigma_0^{-1}\mu_0 + m\Sigma^{-1}\xi)$. Here, m is the sample size and μ'_0 is the posterior mean (and, hence, the Bayes estimator for μ).

Let the utility function be of the popular exponential form $u(\omega) := 1 - e^{-\lambda\omega}$ for some given $\lambda > 0$. It follows that $1 + \xi^\top x|\mu'_0 \sim \mathcal{N}(1 + x^\top \mu'_0, x^\top \Sigma'_0 x)$ and $1 + \xi^\top x|\xi \sim \mathcal{N}(1 + x^\top \mu'_0, x^\top (\Sigma'_0 + \Sigma)x)$. Using the moment-generating function of the Normal distribution, we can obtain the Separate-EO estimator as an optimal solution of

$$\max_{x \in \Delta^n} \left\{ 1 - e^{-\lambda(1+x^\top \mu'_0) + \frac{1}{2}\lambda^2(x^\top \Sigma'_0 x)} \right\},$$

and the Joint-EO estimator under the linear loss can be computed as an optimal solution of

$$\max_{x \in \Delta^n} \left\{ 1 - e^{-\lambda(1+x^\top \mu'_0) + \frac{1}{2}\lambda^2(x^\top (\Sigma'_0 + \Sigma)x)} \right\}.$$

In fact, since the exponential function is monotone, the above problems can be converted to the quadratic programs

$$\max_{x \in \Delta^n} \left\{ (x^\top \mu'_0) - \frac{1}{2} \lambda (x^\top \Sigma'_0 x) \right\} \quad (13)$$

and

$$\max_{x \in \Delta^n} \left\{ (x^\top \mu'_0) - \frac{1}{2} \lambda (x^\top (\Sigma'_0 + \Sigma)x) \right\}, \quad (14)$$

respectively.

To compare the empirical performance of the Separate-EO and Joint-EO solution estimators, we use a real data set from Kocuk and Cornuéjols (2018) based on 11 sectors in the Standard & Poor's 500 index spanning 360 months. In particular, let $\xi^1, \dots, \xi^{360}$ be the sector returns, and let $w^1, \dots, w^{360}$ be the weights of the market portfolio constructed based on the market capitalization of each sector. We use the Black–Litterman approach to construct a prior distribution; see Black and Litterman (1991) and (1992) for the original derivation and Bertsimas et al. (2012) for a modern interpretation.

Let m be a fixed integer denoting the number of data points used in the estimation. For each $v = 181, \dots, 360$, we conduct the following experiment:

- We compute the sample average and sample covariance $\hat{\mu}^{v-1}$ and $\hat{\Sigma}^{v-1}$ of the return values $\xi^{v-m}, \dots, \xi^{v-1}$.
- Using the Black–Litterman approach, we set the prior mean as $\lambda \hat{\Sigma}^{v-1} w^{v-1}$ and prior covariance as $\tau \hat{\Sigma}^{v-1}$, for some $\tau > 0$. Here, small (large) τ indicates strong (weak) prior, and small (large) λ represents low (high) risk aversion.
- We set the observed data mean as $\hat{\mu}^{v-1}$ and likelihood covariance as $\hat{\Sigma}^{v-1}$.
- We solve problems (13) and (14) to obtain the Separate-EO estimator $\hat{x}^S$ and the Joint-EO estimator $\hat{x}^{J,L}$.
- We compute the realized portfolio returns as $r_v^S := \xi^{v\top} \hat{x}^S$ and $r_v^{J,L} := \xi^{v\top} \hat{x}^{J,L}$.
- Finally, we compute some performance measures such as average, standard deviation, and the Sharpe ratio, given the realized returns $r_{181}^S, \dots, r_{360}^S$ and $r_{181}^{J,L}, \dots, r_{360}^{J,L}$.

The results with varying values of risk aversion λ and prior confidence τ are given in Figure 1. We observe that the performance of the Joint-EO estimator is better than that of the Separate-EO estimator in terms of all the measures considered, especially when the risk-aversion coefficient λ is small. Due to our choice of the conjugate families, the computational effort of the separate and joint schemes is similar, as the likelihood and posterior predictive are both Normal. In addition, the accuracy of the joint method is uniformly higher than the separate method. Next, we present another application that exhibits an opposite performance.

3.2. Application II: Geometric Programs

Geometric programs (GPs) have applications in a wide variety of areas, including circuit design, optimal control, nonlinear network design, and chemical equilibrium problems; see Boyd et al. (2007) for a tutorial on geometric programming and for a comprehensive list of applications.

The objective and constraints of GPs are described by *posynomial* functions of the form $f(z) = \sum_k c_k \prod_i z_i^{\alpha_{ik}}$, where $c_k > 0$ is a constant, $z_i > 0$ is a decision variable, and $\alpha_{ik} \in \mathbb{R}$ is an exponent. The trick to solve these nonconvex programs is to use the transformation $z_i = e^{x_i}$ for all i and replace $f(z)$ with $\log f(e^x)$ in the model. Such problems can be formulated as

$$\min_{x \in \mathcal{X}} \left\{ \sum_k c_k e^{\alpha_k^\top x} \mid \log f_j(e^x) \leq 0, \forall j = 1, \dots, m \right\}, \quad (15)$$

where $f_j(z)$ is a posynomial function, and $\mathcal{X}$ is a polyhedron.

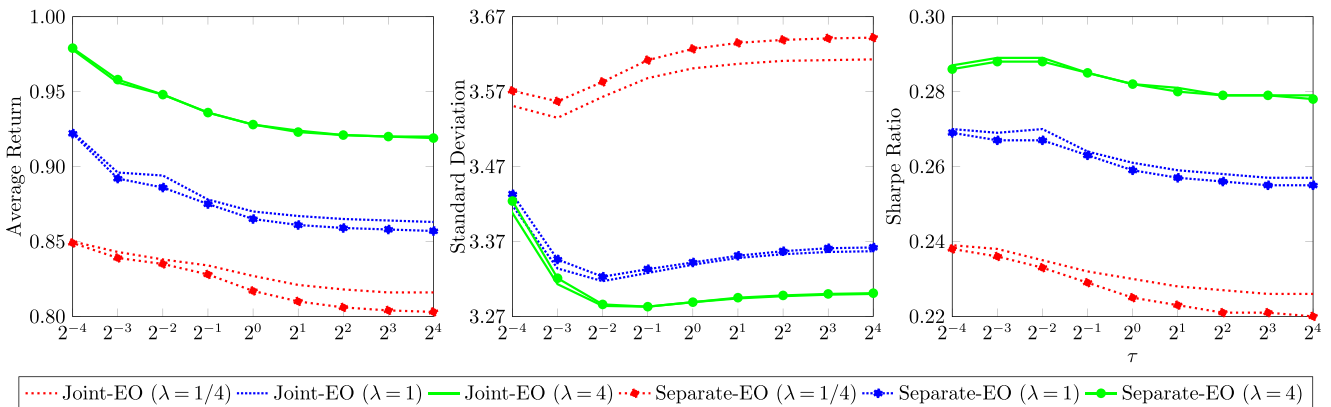
In a stochastic variant of (15), the exponents α_{ik} are random with distribution $g_{ik}(\xi_{ik}|\theta_{ik})$. Including all constraints in $\mathcal{X}$, we write the general form of the stochastic geometric program as

$$\min_{x \in \mathcal{X}} \mathbb{E}_{\xi|\theta} \left[\sum_k c_k e^{\xi_k^\top x} \right], \quad (16)$$

where the expectation is taken with respect to the joint distribution of the random variables.

Under the assumption that the random exponents ξ_{ik} are independent, the objective function of (16) decomposes into separable terms as $\mathbb{E}_{\xi|\theta} [\sum_k c_k e^{\xi_k^\top x}] = \sum_k c_k \prod_i \mathbb{E}_{\xi_{ik}|\theta_{ik}} [e^{\xi_{ik} x_i}]$. The unique property of such decomposition

Figure 1. (Color online) Empirical Performance Comparison of Separate-EO and Joint-EO Solution Estimators ($m = 12$)



is that the term $\mathbb{E}_{\xi_{ik}|\theta_{ik}}[e^{\xi_{ik}x_i}]$ is equal to the *moment generating function* $M_{\xi_{ik}|\theta_{ik}}(x_i)$ of the distribution $g_{ik}(\xi_{ik}|\theta_{ik})$. As a result, the stochastic geometric problem (16) reduces to

$$\min_{x \in \mathcal{X}} \sum_k c_k \prod_i M_{\xi_{ik}|\theta_{ik}}(x_i). \quad (17)$$

Assume now that the parameter θ_{ik} is random with prior $\pi(\theta_{ik})$. We next use the results of Section 2.1 to evaluate the Separate-EO and Joint-EO solution estimators for (17). We present the risk comparison under the linear loss (2).

Corollary 2. Consider the stochastic geometric program (17). Assume that ξ_{ik} has distribution $g_{ik}(\xi_{ik}|\theta_{ik})$ and its parameter θ_{ik} has prior $\pi(\theta_{ik})$. Further, assume that an observation $\bar{\xi}_{ik}$ is available. Then we have the following:

- (i) $\hat{x}^S(\bar{\xi}) \in \operatorname{argmin}_{x \in \mathcal{X}} \sum_k c_k \prod_i M_{\xi_{ik}|\hat{\theta}_{ik}^B}(x_i)$, where $\hat{\theta}_{ik}^B$ is the Bayes estimator of θ_{ik} under the squared error loss.
- (ii) $\hat{x}^{J,L}(\bar{\xi}) \in \operatorname{argmin}_{x \in \mathcal{X}} \sum_k c_k \prod_i M_{\xi_{ik}|\bar{\xi}_{ik}}(x_i)$, where $M_{\xi_{ik}|\bar{\xi}_{ik}}$ is the moment-generating function of the posterior predictive distribution $h(\xi_{ik}|\bar{\xi}_{ik})$.

For a numerical illustration, consider the stochastic geometric program (17) for the Exponential-Gamma conjugate pair with prior hyperparameters α_k and β_k . We assume that a set of observations $\bar{\xi}_k$, $k \in [K]$ for some $K \in \mathbb{Z}$, is available. According to Corollary 2, obtaining the Separate-EO solution estimator amounts to solving the convex program

$$z^S := \min_{x \in \mathcal{X}} \left\{ \sum_{k=1}^K c_k \frac{\lambda_k(\bar{\xi}_k)}{\lambda_k(\bar{\xi}_k) - x_k} \right\}, \quad (18)$$

where $\lambda_k(\bar{\xi}_k) = \frac{\alpha_k + 1}{\beta_k + \bar{\xi}_k}$. Here, we implicitly use the fact that the moment-generating function of the Exponential distribution is available in closed form. However, this approach is not applicable for the Joint-EO method, as the moment-generating function of the posterior predictive distribution, Lomax, is known to be mathematically intractable. Therefore, we require to make use of sampling-based methods (such as SAA) to obtain the Joint-EO solution estimator. In particular, let $\tilde{\xi}_{rk}$ be a Lomax random variable¹ with parameters $\beta_k + \bar{\xi}_k$ and $\alpha_k + 1$, for $k \in [K]$ and $r \in [R]$, where R denotes the sample size. Then applying the SAA method to (15) for the Lomax distribution, we can approximate the Joint-EO solution estimator with the solution of the following convex program:

$$z^{J,L}(R) := \min_{x \in \mathcal{X}} \left\{ \sum_{k=1}^K c_k \frac{1}{R} \sum_{r=1}^R e^{\tilde{\xi}_{rk} x_k} \right\}. \quad (19)$$

We now discuss an experimental setting. We define the feasible region $\mathcal{X}$ as the intersection of the standard simplex $\Delta_K := \{x \in \mathbb{R}_+^K : \sum_{k=1}^K x_k = 1\}$ with a polyhedron $\mathcal{P} := \{x \in \mathbb{R}^K : Ax = b\}$, where $A \in \mathbb{R}^{M \times K}$ and $b \in \mathbb{R}^M$. Such linear constraints are common in geometric programs when the original posynomial function $f_j(z)$ contains a single summation, that is, $f_j(z) = c \prod_i z_i^{\alpha_i}$. In this case, the aforementioned geometric transformation yields $f_j(e^x) = ce^{\alpha^\top x}$, and hence the log constraint as presented in (15) becomes linear in x ; that is, $\alpha^\top x \leq -\log c$. We randomly generate the parameters according to the following rules:

$$\begin{aligned} c_k &= 1 & \alpha_k &\sim \text{Unif}(5, 10) & \beta_k &\sim \text{Unif}(5, 10) \\ A_{mk} &\sim \text{Unif}(1, 10) & b_m &\sim \text{Unif}(10, 50) & \bar{\xi}_k &= \frac{\alpha_k}{\beta_k}. \end{aligned}$$

In Table 1, we compare the Separate-EO and the Joint-EO solution estimators with $R = 100, 1,000$, and $10,000$ for five randomly generated instances in dimension $K = 1,000$ with $M = 10$ linear constraints. To compare the quality of the solutions obtained from the separate and approximate joint approach, we evaluate the solutions in the objective function of (19) for a new set of $R' = 10,000$ independent samples generated from the Lomax distribution. We report these values in column SEO for the Separate-EO and in columns JEO(R) for the Joint-EO methods. The results give several interesting observations from different perspectives.

First, we observe that obtaining the JEO(R) solution estimators becomes increasingly demanding with larger R values in terms of the computational effort. For example, for moderate-sized instances considered here, it typically takes about two minutes to obtain a solution estimator when a sample size of $10,000$ is used. For these problem instances, we ran into memory issues with $R = 100,000$ sample points. The SEO estimators, on

Table 1. Computational Results for Randomly Generated Stochastic Geometric Programs with Exponential-Gamma Pair ($K = 1,000$, $M = 10$)

Ins.	SEO		JEO (100)		JEO (1,000)		JEO (10,000)	
	Obj. val.	Time (s)	Obj. val.	Time (s)	Obj. val.	Time (s)	Obj. val.	Time (s)
1	1,061.09	2.55	1,009.50	3.57	1,009.34	13.22	1,009.16	131.27
2	5,944.84	2.69	1,016.91	3.79	1,014.87	12.68	1,014.39	100.33
3	76,641.14	2.93	1,313.29	3.65	1,166.62	12.75	1,064.67	95.55
4	2,020.64	2.69	4,891.68	3.73	2,936.59	12.37	1,547.22	89.58
5	1,721.14	2.83	25,154.60	3.58	4,540.34	12.86	2,121.23	110.28

Note. Five problem instances (Ins.) are solved, and the corresponding objective values (Obj. val.) and computational times in seconds (s) are recorded.

the other hand, do not suffer from these issues since they are obtained as a solution to a deterministic convex program, which scales reasonably well with the instance size.

Second, because of the above issue, the objective values $z^{J,L}(R)$ for some instances of the Joint-EO method (instances 3, 4, and 5) exhibit a slow convergence. This causes a serious concern for determining a reliable stopping criterion for the SAA algorithms, as not only the objective values but also the optimal solutions are unstable for different sampling sizes. These obstacles for the computation of the Joint-EO solution estimators stem from the slow convergence result of the SAA method for certain problem structures. In particular, Birge (2016) argues that, despite the asymptotic convergence of the SAA method with the growth of the sample size, the optimization problem is still prone to severe estimation errors when the number of uncertain elements (often linked to the problem scale) is large. In contrast, the Separate-EO solution estimators, despite being suboptimal, enjoy a fast and stable solution process.

Third, unlike the pattern observed in the portfolio application, the approximate solution of the Joint-EO does not always give a better solution than that of the Separate-EO method. In particular, in instance 5, the SEO solution consistently has a better objective value than the JEO solution for different sizes. In instance 4, interestingly, the SEO solution has a better objective value than the JEO for sample sizes 100 and 1,000, whereas, for sample size 10,000, the JEO solution dominates the SEO at the price of a larger computation time. This observation clearly shows that the computational limitations dictated by the problem structure and prior-likelihood distributions can change the dynamic between the performance of solution estimators obtained from the separate and joint approaches.

The applications presented in Sections 3.1 and 3.2 show a remarkable difference between the separate and joint methods when it comes to the accuracy versus computational efficiency trade-off. This suggests a closer study of these two methods when applied to various problem structures and probability distributions. We note here that if the separate and joint problems are solved approximately using sampling methods such as SAA, and both likelihood and posterior distributions are similarly convenient to sample from, then the two models will have a similar computational complexity. In such a situation, the joint problem is clearly the preferred method. However, in this paper, we investigate the exact solutions of these two problems, which potentially leads to a difference in the computational effort, as illustrated in this section. In the remainder of this paper, we develop a theoretical foundation for such a comparison.

4. Risk Comparison Between the Separate and Joint Estimators

Because of the trade-off observed in Section 3, two basic questions arise: (i) Under what conditions do the Separate-EO and Joint-EO methods lead to the same solution estimator? (ii) When the two estimators are different, is there a bound on the risk difference?

4.1. Sufficient Conditions for the Estimators to be the Same

First, we present sufficient conditions under which the Separate-EO and the Joint-EO yield the same solution estimators for the linear loss (2).

Proposition 4. Assume that the objective function $\mathbb{E}_{\xi|\theta}[f(\xi, \mathbf{x})]$ of (1) is multilinear in θ and that for any pair $(i, j) \in [n] \times [n]$, $i \neq j$ for which the product $\theta_i \theta_j$ appears in $\mathbb{E}_{\xi|\theta}[f(\xi, \mathbf{x})]$, we have that ξ_i and ξ_j , as well as θ_i and θ_j , are independent. Then, the set of Bayes solution estimators obtained from the Joint-EO scheme under the loss (2) is equal to the one obtained from the Separate-EO scheme; that is, $\hat{\mathbf{x}}^{J,L}(\bar{\xi}) = \hat{\mathbf{x}}^S(\bar{\xi})$.

Next, we present conditions for the quadratic loss (6).

Proposition 5. Assume that (1) has a unique optimal solution $x^*(\theta)$ for any given θ . Assume also that each component of $x^*(\theta)$ is multilinear in θ and that for any pair $(i, j) \in [n] \times [n]$, $i \neq j$ for which the product $\theta_i \theta_j$ appears in a component, we have that ξ_i and ξ_j , as well as θ_i and θ_j , are independent. Then, the set of Bayes solution estimators obtained from the Joint-EO scheme under the loss (6) is equal to the one obtained from the Separate-EO scheme; that is, $\hat{x}^{J,Q}(\bar{\xi}) = \hat{x}^S(\bar{\xi})$.

The proofs of these two propositions are given in the electronic companion.

We give three examples of the portfolio selection model to illustrate how problem modeling and parameter distributions affect the solution estimators, leading to equal and/or different values.

In the first case, all three solution estimators are the same.

Example 1. For the classical Markowitz (mean-variance) formulation (Markowitz 1959)

$$\max_{x \in \mathcal{X}} \mu^\top x - \frac{\tau}{2} x^\top \Sigma x, \quad (20)$$

assume that $\mathcal{X} = \mathbb{R}^n$. This occurs when both long and short positions are allowed on all assets and a riskless asset is available.

Assume now that μ is unknown and has a prior distribution $\pi(\mu)$, and assume that Σ is known and positive definite. The unconstrained problem (20) is convex and has a unique optimal solution $x^*(\mu) = \frac{1}{\tau} \Sigma^{-1} \mu$. The conditions of Propositions 4 and 5 are satisfied. Therefore, all three solution estimators are equal to $\hat{x}^{J,L}(\bar{\xi}) = \hat{x}^{J,Q}(\bar{\xi}) = \hat{x}^S(\bar{\xi}) = \frac{1}{\tau} \Sigma^{-1} \hat{\mu}^B(\bar{\xi})$, where $\hat{\mu}^B(\bar{\xi})$ is the Bayes estimator (posterior mean) of θ , which is a shrinkage vector under Exponential distributions.

Next, we give an example where the solution estimators of the Separate-EO and the Joint-EO under the loss (2) are equal but different from that of the Joint-EO under the loss (6).

Example 2. Consider the variation of portfolio selection where we want to maximize expected return subject to bounding the risk of the portfolio. This problem can be modeled as (1), where $f(\xi, x) = \xi^\top x$ and $\mathcal{X} = \{x \in \mathbb{R}^n | x^\top \Sigma x \leq r^2\}$. We obtain that

$$\max_{x \in \mathcal{X}} \mu^\top x. \quad (21)$$

Assume now that μ is unknown and has a prior distribution $\pi(\mu)$. Since the objective function of (21) is linear in μ , it follows from Propositions 4 that the solution estimators of the Separate-EO and the Joint-EO under the loss (2) are equal; that is, $\hat{x}^{J,L}(\bar{\xi}) = \hat{x}^S(\bar{\xi}) = \frac{\Sigma^{-1} \hat{\mu}^B(\bar{\xi})}{\|\Sigma^{-1/2} \hat{\mu}^B(\bar{\xi})\|} r$. On the other hand, the unique optimal solution of (21) is obtained as $x^*(\mu) = \frac{\Sigma^{-1} \mu}{\|\Sigma^{-1/2} \mu\|} r$, which is not a linear function of μ . As a result, the conditions of Proposition 5 are not satisfied. This yields $\hat{x}^{J,Q}(\bar{\xi}) = \mathbb{E}_{\mu|\bar{\xi}}[\frac{\Sigma^{-1} \mu}{\|\Sigma^{-1/2} \mu\|} r]$, a quantity that can be very difficult to compute depending on the distributions.

Finally, we illustrate the converse of Example 2, where the solution estimators of the Separate-EO and the Joint-EO under the loss (6) are equal but different from that of the Joint-EO under the loss (2).

Example 3. Consider the setting of Example 1, where the mean and covariance matrix of the asset returns are known, but the risk aversion factor τ is random with Exponential distribution $\tau \sim \text{Exp}(\lambda)$. This problem can be modeled as (1), where $f(\tau, x) = \mu^\top x - \frac{\tau}{2} x^\top \Sigma x$ and $\mathcal{X} = \mathbb{R}^n$. We obtain that

$$\max_{x \in \mathcal{X}} \mu^\top x - \frac{1}{2\lambda} x^\top \Sigma x. \quad (22)$$

Assume now that parameter λ has a prior $\pi(\lambda)$. The unique optimal solution of the problem is $x^*(\lambda) = \lambda \Sigma^{-1} \mu$. Since $x^*(\lambda)$ is linear in λ , it follows from Proposition 5 that the solution estimators of the Separate-EO and the Joint-EO under the loss (6) are equal; that is, $\hat{x}^{J,Q}(\bar{\xi}) = \hat{x}^S(\bar{\xi}) = \hat{\lambda}^B(\bar{\xi}) \Sigma^{-1} \mu$, where $\hat{\lambda}^B(\bar{\xi})$ is the Bayes estimator (posterior mean) of λ . On the other hand, since the objective function of (22) is not linear in λ , it does not satisfy the conditions of Proposition 4. In particular, $\hat{x}^{J,L}(\bar{\xi})$ is the maximizer of $\max_{x \in \mathcal{X}} \mu^\top x - \mathbb{E}_{\lambda|\bar{\xi}}[\frac{1}{2\lambda}] x^\top \Sigma x$. This yields $\hat{x}^{J,L}(\bar{\xi}) = \frac{\Sigma^{-1} \mu}{\mathbb{E}_{\lambda|\bar{\xi}}[\frac{1}{\lambda}]}$, which may be computable only numerically depending on the distributions.

4.2. The Separate EO Solution Estimator Can Be Arbitrarily Poor

The next example shows that the Separate-EO solution can yield arbitrarily poor solutions under the linear loss (2).

Example 4. Assume that random variables ξ_i for $i \in [n]$ are independent and follow a Normal distribution $\mathcal{N}(\mu_i, \sigma_i^2)$, where μ_i is unknown and σ_i^2 is known. Assume also that the parameters μ_i for $i \in [n]$ follow a Normal distribution $\mathcal{N}(\lambda_i, \delta_i^2)$, where λ_i and δ_i^2 are known. In this case, the posterior $\mu_i | \bar{\xi}_i$ has a Normal distribution $\mathcal{N}(\eta_i, \zeta_i^2)$, where $\eta_i = \frac{\sigma_i^2}{\sigma_i^2 + \delta_i^2} \lambda_i + \frac{\delta_i^2}{\sigma_i^2 + \delta_i^2} \bar{\xi}_i$ and $\zeta_i^2 = \frac{\sigma_i^2 \delta_i^2}{\sigma_i^2 + \delta_i^2}$. Consider an instance of the stochastic program (1), where the objective function is $\mathbb{E}_{\xi|\mu}[f(\xi, x)] = \sum_{i=1}^n \mu_i^2 x_i$ and $\mathcal{X}$ is an n -simplex; that is,

$$\max \left\{ \sum_{i=1}^n \mu_i^2 x_i \mid \sum_{i=1}^n x_i = 1, x_i \geq 0, \forall i \in [n] \right\}. \quad (23)$$

To obtain the Bayes solution estimator for the Joint-EO method under the linear loss (2), we solve (23) by computing the posterior expectation of its objective function; see (4). We obtain the objective function $\sum_{i=1}^n (\eta_i^2 + \zeta_i^2) x_i$ since $\mathbb{E}_{\mu|\bar{\xi}}[\mu_i^2] = \eta_i^2 + \zeta_i^2$. The optimal solution (Bayes solution estimator) is attained at $\hat{x}^{J,L}(\bar{\xi}) = e^j$, where e^j is the j th unit vector and j is the index of the maximum value among $\{\eta_i^2 + \zeta_i^2\}_{i \in [n]}$. In contrast, the minimizer (worst solution) of the above problem is attained at $x = e^k$, where k is the index of the minimum value among $\{\eta_i^2 + \zeta_i^2\}_{i \in [n]}$. Now we calculate the estimator obtained from the Separate-EO method. The Bayes estimator of the unknown parameter μ under the squared error loss is the posterior mean η and therefore the resulting objective function in (23) is $\sum_{i=1}^n \eta_i^2 x_i$. Using similar arguments as above, the optimal solution of this problem is $\hat{x}^S(\bar{\xi}) = e^l$, where l is the index of the maximum value among $\{\eta_i^2\}_{i \in [n]}$. In order to compare the quality of the estimators obtained from the Joint-EO and Separate-EO methods, we need to evaluate the objective value of e^l in the Joint-EO problem given above. For any values of the parameters σ , δ , and data $\bar{\xi}$ that satisfy $l = k$, the optimal solution of the Separate-EO method matches the worst solution in the Joint-EO problem. This shows that the Separate-EO estimator can be arbitrarily weak compared with the Joint-EO estimator.

For a numerical illustration, assume that $n = 2$, $\bar{\xi}_1 = 2$, $\bar{\xi}_2 = 1$, $\lambda_1 = \lambda_2 = 0$, $\sigma_1 = \sigma_2 = 1$, $\delta_1 = 1$, and $\delta_2 = 3$. We compute $\eta_1 = 1$, $\eta_2 = \frac{9}{10}$, $\zeta_1^2 = \frac{1}{2}$, and $\zeta_2^2 = \frac{9}{10}$. It follows that $\eta_1^2 > \eta_2^2$ and $\eta_1^2 + \zeta_1^2 < \eta_2^2 + \zeta_2^2$. This shows that the Bayes solution estimator for the Joint-EO is $\hat{x}^{J,L}(\bar{\xi}) = (0, 1)$ and the solution estimator for the Separate-EO is $\hat{x}^S(\bar{\xi}) = (1, 0)$.

A similar example is presented for the quadratic loss (6) in the electronic companion. Even though these examples show that the Separate-EO method can yield arbitrarily poor solution estimators compared with the Joint-EO method, there are problem classes for which explicit bounds on the difference between these two estimators can be derived. We present a popular such class in the next section.

5. Piecewise-Linear Functions

Piecewise-linear functions encompass a variety of applications in optimization. They are used to model approximate nonlinear functions, constraints involving partial differential equations, discount policy in marketing, fixed-charge problems, and so forth; see Vielma et al. (2010). The stochastic variants of piecewise-linear functions can be categorized into two major classes: one in which the stochasticity appears in the location of breakpoints (intercept), and one in which the stochasticity appears on the variable coefficients (slope). The former class includes stochastic newsvendor and median problems, and the latter includes rectifier activation functions in deep neural networks. In this section, we investigate examples of both of these classes using our Separate-EO and Joint-EO methods.

5.1. Piecewise-Linear Functions with One Break Point

Define $f(\xi, x) : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ as

$$f(\xi, x) = \begin{cases} \bar{f}(\xi, x), & \text{if } \xi \leq x \\ \tilde{f}(\xi, x), & \text{if } \xi \geq x' \end{cases} \quad (24)$$

where $\bar{f}(\xi, x) = \bar{a}\xi + \bar{b}x + \bar{c}$ and $\tilde{f}(\xi, x) = \tilde{a}\xi + \tilde{b}x + \tilde{c}$. We assume that $f(\xi, x)$ is continuous over its domain, including the breakpoint, which yields $\bar{f}(x, x) = \tilde{f}(x, x)$. In the above relation, x is a decision variable and ξ is a random variable with distribution function $g(\xi|\theta)$, and θ is a random parameter with the prior distribution $\pi(\theta)$.

Such functions are commonly used to model inventory control and pricing problems. For instance, consider a single-stage newsvendor problem, where the purchase cost per unit of a product is c and the selling price per unit is s . The random demand ξ follows a distribution $g(\xi|\theta)$ with a random parameter θ that has a prior $\pi(\theta)$. Let x be the decision variable representing the order quantity. Then, the profit is $f(\xi, x) = s \min\{x, \xi\} - cx$,

which is of the form (24), where $\bar{f}(\xi, x) = s\xi - cx$ and $\tilde{f}(\xi, x) = (s - c)x$. Despite being simple, the newsvendor problem encompasses a rich literature in optimization and is considered as one of the most fundamental structures in stochastic programming; see, e.g., Chu et al. (2008) and Liyanage and Shanthikumar (2005). Other areas of application for piecewise-linear functions with uncertain breakpoints emerge in modeling two-phase linear-linear processes such as latent growth behaviors; see Kohli and Harring (2013). These models are widely used to describe developmental processes in education and psychology.

Assume that the likelihood cumulative distribution function (cdf) G satisfies $\lim_{t \rightarrow -\infty} tG(t|\theta) = 0$, a property that holds for most of the standard probability distributions (such as Normal, Exponential, and Geometric distributions). Then, using an integration by part, the expectation of the function (24) is computed as

$$\begin{aligned}\mathbb{E}_{\xi|\theta}[f(\xi, x)] &= \left[\int_{z \leq x} \bar{f}(z, x)g(z|\theta)dz + \int_{z \geq x} \tilde{f}(z, x)g(z|\theta)dz \right] \\ &= \bar{a}\mathbb{E}_{\xi|\theta}[\xi] + \tilde{b}x + \tilde{c} + (\bar{a} - \bar{a}) \int_{z \leq x} G(z|\theta)dz.\end{aligned}\quad (25)$$

It follows from (25) that the function $\mathbb{E}_{\xi|\theta}[f(\xi, x)]$ (i) is infinitely differentiable in x ; (ii) is strictly convex if $\bar{a} - \bar{a} > 0$, strictly concave if $\bar{a} - \bar{a} < 0$, and linear if $\bar{a} - \bar{a} = 0$; and (iii) has a unique stationary point $x^*(\theta) = G^{-1}(\frac{\tilde{b}}{\bar{a} - \bar{a}}|\theta)$ when it is not linear; that is, $\bar{a} - \bar{a} \neq 0$.

Next, we consider the univariate unconstrained stochastic program

$$\max_{x \in \mathbb{R}} \mathbb{E}_{\xi|\theta}[f(\xi, x)], \quad (26)$$

where $f(\xi, x)$ is a piecewise-linear objective function as defined in (24), and where $\bar{a} - \bar{a} > 0$ so that (26) is convex. Assume that the random variable ξ follows distribution $g(\xi|\theta)$ and the parameter θ is random with a prior $\pi(\theta)$. Given an observation $\bar{\xi}$ from the distribution $g(\bar{\xi}|\theta)$, Corollary 3 presents the Separate-EO estimator $\hat{x}^S(\bar{\xi})$ and the Joint-EO estimators $\hat{x}^{J,L}(\bar{\xi})$ and $\hat{x}^{J,Q}(\bar{\xi})$ under the linear and quadratic loss. This result is obtained as a direct consequence of Proposition 2 and Corollary 1.

Corollary 3. Consider the piecewise-linear problem (26), and assume that an observation $\bar{\xi}$ from the distribution $g(\bar{\xi}|\theta)$ is available. Then we have the following:

(i) $\hat{x}^S(\bar{\xi}) = G^{-1}(\frac{\tilde{b}}{\bar{a} - \bar{a}}|\hat{\theta}^B(\bar{\xi}))$, where $\hat{\theta}^B(\bar{\xi})$ is the Bayes estimator of the unknown parameter θ under the squared error loss.

(ii) $\hat{x}^{J,L}(\bar{\xi}) = H^{-1}(\frac{\tilde{b}}{\bar{a} - \bar{a}}|\bar{\xi})$, where $H(\bar{\xi}|\bar{\xi})$ is the cdf of the posterior predictive distribution of ξ given the observation $\bar{\xi}$.

(iii) $\hat{x}^{J,Q}(\bar{\xi}) = \mathbb{E}_{\theta|\bar{\xi}}[G^{-1}(\frac{\tilde{b}}{\bar{a} - \bar{a}}|\theta)]$, where the expectation is taken with respect to the posterior distribution $\Pi(\theta|\bar{\xi})$.

It is clear from Corollary 3 that the solution estimators depend on the distribution of ξ and its corresponding prior. As an example, we analyze the case with a Normal-Normal conjugate pair in detail.

Proposition 6. Consider the stochastic problem (26). Assume that the likelihood distribution is Normal with $\xi \sim \mathcal{N}(\mu, \sigma^2)$ and the prior distribution is Normal with $\mu \sim \mathcal{N}(\mu_0, \delta^2)$. Assume further that the parameters σ^2 , μ_0 and δ^2 are known, and an observation $\bar{\xi}$ is drawn from the likelihood distribution. Then, we have the following:

(i) $\hat{x}^S(\bar{\xi}) = \hat{x}^{J,Q}(\bar{\xi}) = [\rho\mu_0 + (1 - \rho)\bar{\xi}] + \Phi^{-1}(\frac{\tilde{b}}{\bar{a} - \bar{a}})\sigma$.

(ii) $\hat{x}^{J,L}(\bar{\xi}) = [\rho\mu_0 + (1 - \rho)\bar{\xi}] + \Phi^{-1}(\frac{\tilde{b}}{\bar{a} - \bar{a}})\sigma\sqrt{2 - \rho}$.

Here, $\rho := \frac{\sigma^2}{\sigma^2 + \delta^2}$ and $\Phi^{-1}(\cdot)$ is the inverse cdf of a standard Normal random variable.

Proof. (i) Due to Corollary 3(i), we have

$$\hat{x}^S(\bar{\xi}) = G^{-1}\left(\frac{\tilde{b}}{\bar{a} - \bar{a}} \middle| \hat{\mu}^B(\bar{\xi})\right) = \hat{\mu}^B(\bar{\xi}) + \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \bar{a}}\right)\sigma,$$

where $G(\bar{\xi})$ is the cdf of a Normal random variable with mean $\hat{\mu}^B(\bar{\xi}) := \rho\mu_0 + (1 - \rho)\bar{\xi}$ and variance σ^2 . Hence, we obtain $\hat{x}^S(\bar{\xi})$ as stated. We also observe that $\hat{x}^*(\mu)$ is linear in μ . Therefore, by Proposition 5, we conclude that $\hat{x}^S(\bar{\xi}) = \hat{x}^{J,Q}(\bar{\xi})$.

(ii) Due to Corollary 3(ii), we have

$$\hat{x}^{J,L}(\bar{\xi}) = H^{-1}\left(\frac{\tilde{b}}{\bar{a} - \bar{a}} \middle| \bar{\xi}\right) = \hat{\mu}^B(\bar{\xi}) + \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \bar{a}}\right)\sigma\sqrt{2 - \rho},$$

where H is the cdf of a Normal random variable with mean $\hat{\mu}^B(\bar{\xi}) = \rho\mu_0 + (1 - \rho)\bar{\xi}$ and variance $\sigma^2(2 - \rho)$. Hence, the result follows. $\square$

In the previous proposition, the value of ρ is based on $T = 1$ observation. The reader can verify that, in general, for T observations, the proposition holds with $\rho := \frac{\sigma^2}{\sigma^2 + T\delta^2}$. Similar adjustments can be made throughout this section. The computational complexity of obtaining the Separate-EO and Joint-EO estimators is similar, as they both require an inverse Normal cdf computation. We next provide bounds on the difference in their risk values. We will use the following result.

Remark 1. Consider $d(\kappa) = \kappa(\sqrt{1 + 1/(1 + \kappa^2)} - 1)$. Then,

$$(i) \ d(0) = 0, \ (ii) \ \lim_{\kappa \rightarrow \infty} d(\kappa) = 0, \ (iii) \ d^* := \max_{\kappa \geq 0} d(\kappa) \leq 0.22575.$$

Proposition 7. Let $2\tilde{b} \geq \bar{a} - \tilde{a}$. Under the assumptions of Proposition 6, we have

$$\mathcal{R}^L(x^*(\mu), \hat{x}^S(\bar{\xi})) - \mathcal{R}^L(x^*(\mu), \hat{x}^{J,L}(\bar{\xi})) \leq \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \tilde{a}}\right) \tilde{b} \sigma (\sqrt{2 - \rho} - 1) \leq \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \tilde{a}}\right) \tilde{b} \delta d^*.$$

Proof. We first compute an upper bound on the difference in the loss values of the estimators. Define $F(x) := \mathbb{E}_{\xi|\mu}[f(\xi, x)]$ as the objective function of (26). Then, we have

$$\begin{aligned} \mathcal{L}^L(x^*(\theta), \hat{x}^S(\bar{\xi})) - \mathcal{L}^L(x^*(\theta), \hat{x}^{J,L}(\bar{\xi})) &= F(\hat{x}^{J,L}(\bar{\xi})) - F(\hat{x}^S(\bar{\xi})) \\ &\leq F'(\hat{x}^S(\bar{\xi}))(\hat{x}^{J,L}(\bar{\xi}) - \hat{x}^S(\bar{\xi})) \leq \tilde{b} \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \tilde{a}}\right) \sigma (\sqrt{2 - \rho} - 1) \\ &= \tilde{b} \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \tilde{a}}\right) \delta d(\kappa) \leq \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \tilde{a}}\right) \tilde{b} \delta d^*, \end{aligned}$$

where the first equality follows from the definition of the linear loss (2); the first inequality is obtained from the first-order Taylor expansion of the concave function $F(x)$ at point $\hat{x}^{J,L}(\bar{\xi})$ about $\hat{x}^S(\bar{\xi})$; the second inequality follows from (i) $\hat{x}^{J,L}(\bar{\xi}) - \hat{x}^S(\bar{\xi}) = \Phi^{-1}(\frac{\tilde{b}}{\bar{a} - \tilde{a}})\sigma(\sqrt{2 - \rho} - 1)$ because of Proposition 6, (ii) $\Phi^{-1}(\frac{\tilde{b}}{\bar{a} - \tilde{a}}) \geq 0$ because of the assumption $2\tilde{b} \geq \bar{a} - \tilde{a}$, and (iii) $F'(x) \leq \tilde{b}$ for all $x \in \mathbb{R}$, which is deduced from (25); the second equality holds due to the definition of $d(\kappa)$ with $\kappa := \sigma/\delta$ given in Remark 1 and $\rho = \frac{\sigma^2}{\sigma^2 + \delta^2}$; and the last inequality follows from Remark 1(iii). Since the difference in the loss of the two estimators in the above expression is independent of both μ and $\bar{\xi}$, the risk difference obtained by taking the expectations $\mathbb{E}_\mu \mathbb{E}_{\bar{\xi}|\mu}$ as given in (3) remains unchanged, yielding the result. $\square$

The following result is obtained similarly to that of Proposition 7.

Proposition 8. Let $2\tilde{b} \leq \bar{a} - \tilde{a}$. Under the assumptions of Proposition 6, we have

$$\begin{aligned} \mathcal{R}^L(x^*(\mu), \hat{x}^S(\bar{\xi})) - \mathcal{R}^L(x^*(\mu), \hat{x}^{J,L}(\bar{\xi})) &\leq \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \tilde{a}}\right) (\tilde{b} + \bar{a} - \tilde{a}) \sigma (\sqrt{2 - \rho} - 1) \\ &\leq \Phi^{-1}\left(\frac{\tilde{b}}{\bar{a} - \tilde{a}}\right) (\tilde{b} + \bar{a} - \tilde{a}) \delta d^*. \end{aligned}$$

We also analyze the Exponential-Gamma and Geometric-Beta conjugate pairs for which the detailed derivations are provided in the electronic companion. We present the algorithmic summary of these results in Table 2. This table clearly shows how different conjugate pairs affect computational complexity of the stochastic problem (26).

Table 2. Summary of the Computational Effort Required to Obtain the Separate-EO and Joint-EO Solution Estimators for the Univariate Piecewise-Linear Stochastic Problem with Different Likelihood-Prior Pairs

Likelihood	Prior	Separate-EO	Joint-EO (linear loss)	Joint-EO (quadratic loss)
Normal	Normal	Inverse Normal cdf	Inverse Normal cdf	Inverse Normal cdf
Exponential	Gamma	Closed form	Closed form	Closed form
Geometric	Beta	Closed form	Need algorithm	Need numerical integration

5.2. Sum of Piecewise-Linear Functions with One Break Point

In this section, we study the following extension of the piecewise-linear model of Section 5.1,

$$\max_{x \in \mathbb{R}^n} \mathbb{E}_{\xi|\theta} \left[\sum_{i=1}^n f_i(\xi_i, x) \right], \quad (27)$$

where $f_i(\xi_i, x)$ is a piecewise-linear objective function of the form (24); that is, $\tilde{f}_i(\xi_i, x) = \bar{a}_i \xi_i + \bar{b}_i x + \bar{c}_i$ for $\xi_i \leq x$, and $\tilde{f}_i(\xi_i, x) = \tilde{a}_i \xi_i + \tilde{b}_i x + \tilde{c}_i$ for $\xi_i \geq x$. We assume that $\tilde{a}_i - \bar{a}_i > 0$ for each $i \in [n]$, so that (27) is convex. In this problem, x is the single decision variable and ξ_i is a random variable with probability distribution $g_i(\xi_i|\theta_i)$ and the random parameter θ_i has prior distribution $\pi_i(\theta_i)$. In this setting, variables ξ_i are not required to be independent. We further assume that an observation $\bar{\xi}_i$ drawn from $g_i(\bar{\xi}_i|\theta_i)$ for $i \in [n]$ is available.

Stochastic problems of the form (27) have applications in median and location-allocation problems in facility planning. For instance, consider the one-dimensional stochastic median problem, where the position ξ_i of n points on the line is random and each comes from a distribution $g_i(\xi_i|\theta_i)$ for $i \in [n]$. The goal is to find the location x that minimizes the total distance from the random points, that is, $\sum_{i=1}^n |\xi_i - x|$. This problem can be formulated as (27), where $f_i(\xi_i, x) = |\xi_i - x|$, $\tilde{f}_i(\xi_i, x) = x - \xi_i$, and $\tilde{f}_i(\xi_i, x) = \xi_i - x$.

Using (25), the objective function of (27) can be computed as

$$\mathbb{E}_{\xi|\theta} \left[\sum_{i=1}^n f_i(\xi_i, x) \right] = \sum_{i=1}^n \left[\tilde{a}_i \mathbb{E}_{\xi_i|\theta_i}[\xi_i] + \tilde{b}_i x + \tilde{c}_i + (\tilde{a}_i - \bar{a}_i) \int_{z_i \leq x} G_i(z_i|\theta_i) dz_i \right]. \quad (28)$$

Since this is a concave function under the assumption $\tilde{a}_i - \bar{a}_i > 0$ for $i \in [n]$, its maximizer is obtained as the root of the equation

$$\sum_{i=1}^n (\tilde{a}_i - \bar{a}_i) G_i(x|\theta_i) + \sum_{i=1}^n \tilde{b}_i = 0. \quad (29)$$

We next give the Separate-EO estimator $\hat{x}^S(\bar{\xi})$ and the Joint-EO estimator $\hat{x}^{J,L}(\bar{\xi})$ via univariate root finding using (29). We do not present the Joint-EO estimator $\hat{x}^{J,Q}(\bar{\xi})$ since the explicit solution for (29) is not computable in general.

Proposition 9. Consider problem (27). Assume that, for each $i \in [n]$, the likelihood distribution has the pdf $g_i(\xi_i|\theta_i)$ and the prior distribution has the pdf $\pi_i(\theta_i)$. Assume further that an observation $\bar{\xi}_i$ is available for $i \in [n]$. Then, we have the following:

(i) $\hat{x}^S(\bar{\xi})$ is the solution of

$$\sum_{i=1}^n (\tilde{a}_i - \bar{a}_i) G_i(x|\hat{\theta}_i^B(\bar{\xi}_i)) + \sum_{i=1}^n \tilde{b}_i = 0,$$

where $\hat{\theta}_i^B(\bar{\xi}_i)$ is the Bayes estimator of θ_i under the squared error loss, given observation $\bar{\xi}_i$.

(ii) $\hat{x}^{J,L}(\bar{\xi})$ is the solution of

$$\sum_{i=1}^n (\tilde{a}_i - \bar{a}_i) H_i(x|\bar{\xi}_i) + \sum_{i=1}^n \tilde{b}_i = 0,$$

where $H_i(\xi|\bar{\xi}_i)$ is the cdf of the posterior predictive distribution of ξ_i given observation $\bar{\xi}_i$.

Next, we give a generic method to compute an upper bound on the risk difference between the Separate-EO estimator $\hat{x}^S(\bar{\xi})$ and the Joint-EO estimator $\hat{x}^{J,L}(\bar{\xi})$.

Proposition 10. Consider the setting in Proposition 9. Assume that the likelihood $g_i(\xi_i|\theta_i)$ and the posterior predictive $h_i(\xi_i|\bar{\xi}_i)$ are positive at all points in the domain. Then,

$$\begin{aligned} & \mathcal{R}^L(x^*(\theta), \hat{x}^S(\bar{\xi})) - \mathcal{R}^L(x^*(\theta), \hat{x}^{J,L}(\bar{\xi})) \\ & \leq K \mathbb{E}_{\bar{\xi}} \left[\max \left\{ \max_i \left\{ \hat{x}_i^{J,L}(\bar{\xi}_i) \right\} - \min_i \left\{ \hat{x}_i^S(\bar{\xi}_i) \right\}, \max_i \left\{ \hat{x}_i^S(\bar{\xi}_i) \right\} - \min_i \left\{ \hat{x}_i^{J,L}(\bar{\xi}_i) \right\} \right\} \right], \end{aligned}$$

where the expectation $\mathbb{E}_{\bar{\xi}}$ is taken with respect to the marginal distribution of $\bar{\xi}$, $K = \max\{|\sum_{i=1}^n \tilde{b}_i|, |\sum_{i=1}^n \tilde{b}_i + \tilde{a}_i - \bar{a}_i|\}$, and $\hat{x}_i^S(\bar{\xi}_i)$ and $\hat{x}_i^{J,L}(\bar{\xi}_i)$ are the Separate-EO and Joint-EO estimators if there were only a single random variable ξ_i , as computed in Corollary 3.

Proof. We will first obtain intervals that contain the solution estimators $\hat{x}^S(\bar{\xi})$ and $\hat{x}^{J,L}(\bar{\xi})$. Since $g_i(\xi_i|\theta_i)$ and $h_i(\xi_i|\bar{\xi}_i)$ are positive over the domain, their cdf's are strictly increasing. Therefore, we have that

$$\begin{aligned} G_i\left(x \middle| \hat{\theta}_i^B(\bar{\xi}_i)\right) &< \frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} \text{ for } x < G_i^{-1}\left(\frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} \middle| \hat{\theta}_i^B(\bar{\xi}_i)\right), \\ G_i\left(x \middle| \hat{\theta}_i^B(\bar{\xi}_i)\right) &> \frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} \text{ for } x > G_i^{-1}\left(\frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} \middle| \hat{\theta}_i^B(\bar{\xi}_i)\right), \end{aligned}$$

for each $i \in [n]$. Multiplying each of the above relations by $\bar{a}_i - \tilde{a}_i > 0$ and then summing them over i , we deduce that $\hat{x}^S(\bar{\xi}) \in [\min_i\{G_i^{-1}(\frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} | \hat{\theta}_i^B(\bar{\xi}_i))\}, \max_i\{G_i^{-1}(\frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} | \hat{\theta}_i^B(\bar{\xi}_i))\}]$. A similar argument shows that $\hat{x}^{J,L}(\bar{\xi}) \in [\min_i\{H_i^{-1}(\frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} | \bar{\xi}_i)\}, \max_i\{H_i^{-1}(\frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} | \bar{\xi}_i)\}]$.

We now give an upper bound on the loss values of these estimators. Let $F(x) = \mathbb{E}_{\xi|\theta}[\sum_{i=1}^n f_i(\xi_i, x)]$ as the objective function of (27). Since F is concave, we have

$$\begin{aligned} \mathcal{L}^L(x^*(\theta), \hat{x}^S(\bar{\xi})) - \mathcal{L}^L(x^*(\theta), \hat{x}^{J,L}(\bar{\xi})) &= F(\hat{x}^{J,L}(\bar{\xi})) - F(\hat{x}^S(\bar{\xi})) \\ &\leq F'(\hat{x}^S(\bar{\xi}))(\hat{x}^{J,L}(\bar{\xi}) - \hat{x}^S(\bar{\xi})) \leq |F'(\hat{x}^S(\bar{\xi}))| |\hat{x}^{J,L}(\bar{\xi}) - \hat{x}^S(\bar{\xi})| \\ &\leq K \max\left\{\max_i\{\hat{x}_i^{J,L}(\bar{\xi}_i)\} - \min_i\{\hat{x}_i^S(\bar{\xi}_i)\}, \max_i\{\hat{x}_i^S(\bar{\xi}_i)\} - \min_i\{\hat{x}_i^{J,L}(\bar{\xi}_i)\}\right\}, \end{aligned}$$

where the last inequality holds because (i) $|F'(\hat{x}^S)| \leq K$ as $F'(\hat{x}^S) \in [\sum_{i=1}^n \tilde{b}_i + \tilde{a}_i - \bar{a}_i, \sum_{i=1}^n \tilde{b}_i]$ due to (28), and (ii) the previously derived intervals for $\hat{x}^S(\bar{\xi})$ and $\hat{x}^{J,L}(\bar{\xi})$, and the fact that if $w \in [w_1, w_2]$ and $v \in [v_1, v_2]$, then $|w - v| \leq \max\{w_2 - v_1, v_2 - w_1\}$. To compute an upper bound on the risk difference, we take the expectation $\mathbb{E}_{\theta} \mathbb{E}_{\bar{\xi}|\theta}$ or, equivalently, $\mathbb{E}_{\bar{\xi}} \mathbb{E}_{\theta|\bar{\xi}}$ of both sides of the above expression. Since this expression is independent of θ , the desired result is obtained by only taking the expectation $\mathbb{E}_{\bar{\xi}}$. $\square$

The result of Proposition 10 requires an additional expectation with respect to the marginal distribution of $\bar{\xi}$, over the maximum of random variables. As a consequence, the bound given in this proposition can be hard to compute in closed form for general classes of the stochastic program (27). However, it can be computed explicitly for certain objective functions and probability distributions. To illustrate the derivation technique, we next consider the one-dimensional stochastic median problem, where $f_i(\xi_i, x) = |\xi_i - x|$ for $i \in [n]$, and compute the bound on the risk of the Separate-EO and Joint-EO solution estimators for a Normal-Normal conjugate pair. An analysis with Exponential-Gamma pair is provided in the electronic companion.

Proposition 11. Consider the one-dimensional stochastic median problem. Assume that, for each $i \in [n]$, the likelihood distribution is Normal with $\xi_i \sim \mathcal{N}(\mu_i, \sigma_i^2)$ and the prior distribution is Normal with $\mu_i \sim \mathcal{N}(\mu_i^0, \delta_i^2)$. Assume further that the parameters σ_i^2 , μ_i^0 , and δ_i^2 are known, and a realization of locations $\bar{\xi}_i$ is observed. Then, we have

$$\mathcal{R}^L(x^*(\mu), \hat{x}^S(\bar{\xi})) - \mathcal{R}^L(x^*(\mu), \hat{x}^{J,L}(\bar{\xi})) \leq n \left(\max_i\{\mu_i^0\} - \min_i\{\mu_i^0\} + 2\sqrt{\frac{n-1}{n} \sum_{i=1}^n \frac{\delta_i^4}{\sigma_i^2 + \delta_i^2}} \right).$$

Proof. Proposition 6 implies that $v_i := \hat{x}_i^S(\bar{\xi}_i) = \hat{x}_i^{J,L}(\bar{\xi}_i) = \rho_i \mu_i^0 + (1 - \rho_i) \bar{\xi}_i$, where $\rho_i := \frac{\sigma_i^2}{\sigma_i^2 + \delta_i^2}$ for $i \in [n]$. This result follows from the facts that $\frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i} = \frac{1}{2}$ for the median objective function, and therefore $\Phi^{-1}(\frac{\tilde{b}_i}{\bar{a}_i - \tilde{a}_i}) = 0$. Next, we use Proposition 10 to obtain

$$\mathcal{R}^L(x^*(\mu), \hat{x}^S(\bar{\xi})) - \mathcal{R}^L(x^*(\mu), \hat{x}^{J,L}(\bar{\xi})) \leq n \mathbb{E}_{\bar{\xi}} \left[\max_i\{v_i\} - \min_i\{v_i\} \right]. \quad (30)$$

Now, we compute an upper bound on the right-hand-side using a well-known result in probability theory to bound max/min expectations. In particular, for random variables $z_1, \dots, z_n$, we have $\mathbb{E}[\max_i\{z_i\}] \leq \max_i\{\mathbb{E}[z_i]\} + \sqrt{\frac{n-1}{n} \sum_{i=1}^n \text{Var}(z_i)}$ and $\mathbb{E}[\min_i\{z_i\}] \geq \min_i\{\mathbb{E}[z_i]\} - \sqrt{\frac{n-1}{n} \sum_{i=1}^n \text{Var}(z_i)}$; see Aven (1985). Since $\bar{\xi}_i \sim \mathcal{N}(\mu_i^0, \sigma_i^2 + \delta_i^2)$ for

each $i \in [n]$, we have $\mathbb{E}[v_i] = \mu_i^0$ and $\text{Var}(v_i) = \frac{\delta_i^4}{\sigma_i^2 + \delta_i^2}$. Using these relations in the above bounding inequalities, we obtain

$$\mathbb{E}\left[\max_i\{v_i\}\right] \leq \max_i\{\mu_i^0\} + \sqrt{\frac{n-1}{n} \sum_{i=1}^n \frac{\delta_i^4}{\sigma_i^2 + \delta_i^2}} \quad \text{and} \quad \mathbb{E}\left[\min_i\{v_i\}\right] \geq \min_i\{\mu_i^0\} - \sqrt{\frac{n-1}{n} \sum_{i=1}^n \frac{\delta_i^4}{\sigma_i^2 + \delta_i^2}}.$$

Plugging these bounds into (30) gives the desired conclusion. $\square$

5.3. Piecewise-Linear Functions with Random Slope

In this section, we consider a stochastic piecewise-linear model that includes randomness in the coefficients of the decision variables. This class of problems models the *rectified linear unit* (ReLU) activation function often used in neural networks; see Goodfellow et al. (2016). Unlike for the class studied in Sections 5.1 and 5.2, an explicit derivation of Separate-EO and Joint-EO solution estimators seems out of reach in this case. Therefore, we illustrate the difference in the quality of solutions obtained from the separate and joint methods through a computational experiment in a streamlined neural network model with a single layer.

Suppose that there are n features denoted by random variables $\xi_1, \dots, \xi_n$ that define an output O through the following procedure:

$$O = \max\left\{0, \sum_{j=1}^n \xi_j w_j + w_{n+1} + \epsilon\right\}, \quad (31)$$

where $w_1, \dots, w_n, w_{n+1}$ are the true weights of the random features and ϵ is a random error. The goal is to model this problem using a single-layer neural network and to evaluate the accuracy of the Separate-EO and Joint-EO methods in estimating the weights.

To this end, we consider a setting where the likelihood and prior distributions of ξ_i are independent Normal distributions with known variance and prior mean; that is, $\xi_j | \mu_j \sim \mathcal{N}(\mu_j, \sigma_j^2)$ and $\mu_j \sim \mathcal{N}(\mu_{0j}, \delta_{0j}^2)$, $j = 1, \dots, n$. Further, we assume that T observations drawn from the likelihood are available as $\bar{\xi}^1, \dots, \bar{\xi}^T$. For such a conjugate pair, the posterior and predictive distributions are readily available.

The corresponding single-layer neural network model with n input nodes and a single output node is represented as

$$O = \max\left\{0, \sum_{j=1}^n \xi_j x_j + x_{n+1}\right\}, \quad (32)$$

where $x_1, \dots, x_n, x_{n+1}$ are decision variables. To determine these variables, we need to solve the problem

$$\min_{x \in \mathbb{R}^{n+1}} \mathbb{E}_{\xi | \mu} \left| O - \max\left\{0, \sum_{j=1}^n \xi_j x_j + x_{n+1}\right\} \right|. \quad (33)$$

The Separate-EO estimator can be obtained as an optimal solution of

$$\min_{x \in \mathbb{R}^{n+1}} \mathbb{E}_{\xi | \mu^B(\bar{\xi})} \left| O - \max\left\{0, \sum_{j=1}^n \xi_j x_j + x_{n+1}\right\} \right|, \quad (34)$$

where $\mu^B(\bar{\xi})$ is the Bayes estimator of μ under the quadratic loss, whereas the Joint-EO estimator under the linear loss can be computed as an optimal solution of

$$\min_{x \in \mathbb{R}^{n+1}} \mathbb{E}_{\xi | \bar{\xi}} \left| O - \max\left\{0, \sum_{j=1}^n \xi_j x_j + x_{n+1}\right\} \right|. \quad (35)$$

Table 3. Test-Error Results in the Deep Learning Application with Different Problem Sizes

		$\bar{\sigma} = 0.10$		$\bar{\sigma} = 0.15$		$\bar{\sigma} = 0.20$		$\bar{\sigma} = 0.25$	
n	T	SEO	JEO	SEO	JEO	SEO	JEO	SEO	JEO
10	10	0.367	0.366	0.274	0.231	0.236	0.198	0.179	0.174
20	20	0.579	0.493	0.318	0.201	0.277	0.245	0.358	0.210
30	30	0.758	0.533	0.511	0.315	0.255	0.184	0.293	0.203
40	40	0.497	0.387	0.540	0.333	0.448	0.264	0.324	0.212
50	50	0.592	0.378	0.612	0.389	0.494	0.334	0.366	0.252

Note. SEO and JEO represent Separate-EO and Joint-EO, respectively.

Since solving problems (34) and (35) directly does not seem possible, we adopt a sampling approach as common in machine learning. In particular, assuming that R sample points are available from the posterior and predictive distributions, we define an approximate problem

$$\min_{x \in \mathbb{R}^{n+1}} \frac{1}{R} \sum_{r=1}^R \left| O^r - \max \left\{ 0, \sum_{j=1}^n \xi_j^r x_j + x_{n+1} \right\} \right|, \quad (36)$$

which minimizes the average prediction error over the sample space. We denote the optimal solutions of problem (36) as $\hat{x}^S$ and $\hat{x}^{J,L}$ when sample points are obtained from the posterior and predictive distributions, respectively.

To conduct our experiments, we draw sample points of input features from the posterior and predictive distributions and compute the sample output according to (31) using true weights. We first randomly generate the parameters according to the following rules:

$$w_j, w_{n+1} \sim \text{Unif}(-1, 1), \mu_{0j} \sim \text{Unif}(-1, 1), \delta_{0j} \sim \text{Unif}(0, 1), \sigma_j \sim \text{Unif}(0.05, \bar{\sigma}), j = 1, \dots, n.$$

We also set the distribution of the error term as $\epsilon \sim \mathcal{N}(0, 0.1^2)$.

We solve problem (36) with $R = 1,000$ as a mixed-integer linear program after linearizing the objective. This requires assuming an upper bound on the quantity $\sum_{j=1}^n \xi_j x_j + x_{n+1}$, which we set to 10^5 . Once the optimal weights $\hat{x}^S$ and $\hat{x}^{J,L}$ are calculated, we compute the test error on a new sample point. The average test-error results of 100 macro replications are reported in Table 3. We observe that the test errors with the Joint-EO approach are lower than those of the Separate-EO approach. The extension of such computational studies performed on a deeper neural net with more complex structures while employing different simulation techniques is left as a direction of future research.

6. Conclusion

In the presence of uncertainty, data are used to estimate the unknown elements of stochastic programs. Several techniques can be employed. Under a Bayesian framework, two of the most common estimation rules solve the estimation and optimization problems either separately or jointly. In this paper we explore the analytical and computational trade-offs between the quality and complexity of these schemes in parametric stochastic programs. We use risk as an averaging measure for the quality of the solutions based on two optimization criteria: the gap between the objective values, and the distance between the solutions. We show conditions under which the two schemes yield equal solutions, and give examples when the risk difference between the solutions of the two schemes can be very large. We further study several classes of nonlinear stochastic programs with a piecewise-linear structure, while presenting computational experiments on various applications.

Acknowledgments

The authors thank the associate editor and three anonymous referees for suggestions that helped improve the paper.

Endnote

¹ Using the inverse transformation technique, we can generate such random variables as $\beta_k[(1-u)^{-1/(a_k+1)} - 1]$, where u is a uniform $[0, 1]$ random variable.

References

- Aven T (1985) Upper (lower) bounds on the mean of the maximum (minimum) of a number of random variables. *J. Appl. Probab.* 22(3):723–728.
- Basu A, Nguyen T, Sun A (2019) Admissibility of solution estimators for stochastic optimization. Preprint, submitted January 21, <https://arxiv.org/abs/1901.06976>.
- Ben-Tal A, Ghaoui L, Nemirovski A (2009) *Robust Optimization* (Princeton University Press, Princeton, NJ).
- Bertsimas D, Brown DB, Caramanis C (2011) Theory and applications of robust optimization. *SIAM Rev.* 53:464–501.
- Bertsimas D, Gupta V, Paschalidis I (2012) Inverse optimization: A new perspective on the Black-Litterman model. *Oper. Res.* 60(6):1389–1403.
- Birge J (2016) Uses of sub-sample estimates in stochastic optimization models. Working paper, Booth School of Business, University of Chicago.
- Birge J, Louveaux F (2011) *Introduction to Stochastic Programming* (Springer, New York).
- Black F, Litterman R (1991) Asset allocation: Combining investor views with market equilibrium. *J. Fixed Income* 1(2):7–18.
- Black F, Litterman R (1992) Global portfolio optimization. *Financial Analysts J.* 48(5):28–43.
- Boyd S, Kim SJ, Vandenberghe L, Hassibi A (2007) A tutorial on geometric programming. *Optim. Engrg.* 8:67–127.
- Chu LY, Shanthikumar JG, Shen ZJM (2008) Solving operational statistics via a Bayesian analysis. *Oper. Res. Lett.* 36:110–116.
- Davarnia D, Cornuéjols G (2017) From estimation to optimization via shrinkage. *Oper. Res. Lett.* 45:642–646.
- Diaconis P, Ylvisaker D (1979) Conjugate priors for exponential families. *Ann. Statist.* 7:269–281.
- Donti P, Amos B, Kolter Z (2017) Task-based end-to-end model learning. Preprint, submitted March 13, <https://arxiv.org/abs/1703.04529v1>.
- Ferguson T (1967) *Mathematical Statistics—A Decision Theoretic Approach* (Academic Press, New York).
- Goodfellow I, Bengio Y, Courville A (2016) *Deep Learning* (MIT Press, Cambridge, MA).
- Jiang H, Shanbhag UV (2016) On the solution of stochastic optimization and variational problems in imperfect information regimes. *SIAM J. Optim.* 26(4):2394–2429.
- Jorion P (1986) Bayes-Stein estimation for portfolio analysis. *J. Financial Quantization Anal.* 21:279–292.
- Kadane JB (2011) *Principles of Uncertainty* (CRC Press, Boca Raton, FL).
- Kleywegt AJ, Shapiro A (2004) Stochastic optimization. Technical report, Georgia Institute of Technology, Atlanta.
- Kocuk B, Cornuéjols G (2018) Incorporating Black-Litterman views in portfolio construction when stock returns are a mixture of normals. *Omega*, ePub ahead of print March, <https://doi.org/10.1016/j.omega.2018.11.017>.
- Kohli N, Harring JR (2013) Modeling growth in latent variables using a piecewise function. *Multivariate Behav. Res.* 48(3):370–397.
- Lehmann EL, Casella G (1998) *Theory of Point Estimation* (Springer, New York).
- Levine S, Finn C, Darrell R, Abbeel P (2016) End-to-end training of deep visuomotor policies. *J. Machine Learning Res.* 17:1–40.
- Lim AEB, Shanthikumar JG, Shen ZJM (2006) Model uncertainty, robust optimization, and learning. Johnson MP, Norman B, Secomandi N, eds. *Models, Methods, and Applications for Innovative Decision Making*, TutORials in Operations Research (INFORMS, Catonsville, MD), 66–94.
- Liyanage LH, Shanthikumar JG (2005) A practical inventory control policy using operational statistics. *Oper. Res. Lett.* 33:341–348.
- Markowitz HM (1959) *Portfolio Selection: Efficient Diversification of Investments* (Wiley, New York).
- Shapiro A (2003) Monte Carlo sampling methods. Ruszczyński A, Shapiro A, eds. *Stochastic Programming* (Elsevier, Amsterdam), 352–425.
- Stein C (1956) Inadmissibility of the usual estimator for the mean of a multivariate normal distribution. *Proc. 3rd Berkeley Sympos. Math. Statist. Probab.*, vol. 1 (University of California Press, Berkeley, CA), 197–206.
- Thomas RW, Friend DH, Dasilva LA, Mackenzie AB (2006) Cognitive networks: Adaptation and learning to achieve end-to-end performance objectives. *IEEE Comm. Magazine* 44(12):51–57.
- Vielma JP, Ahmed S, Nemhauser G (2010) Mixed-integer models for nonseparable piecewise linear optimization: Unifying framework and extensions. *Oper. Res.* 58:303–315.
- Wang M, Fang EX, Liu H (2017) Stochastic compositional gradient descent: Algorithms for minimizing compositions of expected-value functions. *Math. Programming* 161:419–449.