

A BAYESIAN APPROACH TO ONLINE SIMULATION OPTIMIZATION WITH STREAMING INPUT DATA

Tianyi Liu

Yifan Lin

Enlu Zhou

School of Industrial and Systems Engineering
Georgia Institute of Technology
755 Ferst Drive NW
Atlanta, GA 30332, USA

ABSTRACT

We consider simulation optimization under input uncertainty, where the unknown input parameter is estimated from streaming data arriving in batches over time. Moreover, data may depend on the decision of the time when they are generated. We take an online approach to jointly estimate the input parameter via Bayesian posterior distribution and update the decision by applying stochastic gradient descent (SGD) on the Bayesian average of the objective function. We show the convergence of our approach. In particular, our consistency result of Bayesian posterior distribution with decision-dependent data might be of independent interest to Bayesian estimation. We demonstrate the empirical performance of our approach on a simple numerical example.

1 INTRODUCTION

The input model (or input distribution) of stochastic simulation is often estimated from a finite amount of input data, resulting in uncertainty about the input model. Ignoring input uncertainty in simulation optimization could lead to risky decisions (see, e.g., Zhou and Xie (2015), Zhou and Wu (2017)). Recent research has made strides towards simulation optimization under input uncertainty. For example, for simulation optimization with a discrete and finite solution space, i.e., the ranking and selection (R&S) problem, Gao et al. (2017), Xiao and Gao (2018), Xiao et al. (2020), Fan et al. (2020) take a distributionally robust approach to select the solution with the best performance in the worst case of input uncertainty; Wu and Zhou (2017) extend the optimal computing budget allocation (OCBA) approach to handle input uncertainty with the goal to select the true best solution; Corlu and Biller (2013), Corlu and Biller (2015) develop procedures to return a subset of superior solutions with desired confidence in the presence of input uncertainty; Song and Nelson (2019) extend the multiple comparisons with the best (MCB) framework from Chang and Hsu (1992) to incorporate input uncertainty. For simulation optimization on a general solution space, Zhou and Xie (2015) propose a Bayesian risk optimization (BRO) framework to reformulate the simulation optimization problem by estimating the input uncertainty with a Bayesian posterior distribution and imposing a risk function with respect to the posterior distribution on the objective function; later Wu et al. (2018) show that this reformulation is consistent and robust against input uncertainty.

The aforementioned works all assume that the input data are given as one fixed batch. However, in modern applications data are often collected over time, and the decision maker often needs to make decisions in an online fashion given all the available data. Such streaming data have only been considered in the simulation community until very recent. Wu and Zhou (2019) consider ranking and selection with streaming input data, and propose a moving average estimator to aggregate simulation outputs across

different time stages together with sequential elimination procedures to select the best solution. Song and Shanbhag (2019) consider the simulation optimization problem with streaming input data and study the convergence properties of stochastic approximation with the maximum likelihood estimation (MLE) of the input parameter. On a related note, Zhou and Liu (2018) develop an online procedure for quantifying input uncertainty with streaming input data; Wu and Zhou (2017) develop a two-stage procedure to allocate budget between collecting input data and doing ranking and selection to balance the input uncertainty and optimization error; Xu et al. (2020) study the problem of joint resource allocation between collecting input data and running simulation replications.

In this paper, we consider the general simulation optimization problem where batches of data arrive over time. We assume a parameterized input model, and thus the distribution family is known but the true input parameter is unknown. At each time stage, our procedure consists of two steps: 1) we use the current batch of data to update the Bayesian posterior distribution of the input parameter, and 2) we take the Bayesian average of the objective function and apply SGD on this reformulated objective function. We consider both cases of exogenous (decision-independent) and endogenous (decision-dependent) input data. In the former case, data follow a fixed distribution that only involves the input parameter. In the latter case, the data follow a time-varying distribution depending not only on the input parameter but also the decision at the current time. This is motivated by many real-world problems where decisions can influence the data process. For example, in the newsvendor problem, the demand rates of customers can depend on the inventory levels decided by the retailer, who only observes data of customer demand and needs to decide the order amount to minimize the expected cost (see Lee et al. (2012)). Due to the correlation and non-stationarity of the data across time stages, the Bayesian estimation with decision-dependent input data is drastically different from the classical Bayesian updating with independent and identically distributed (i.i.d.) data, which poses a great challenge to showing the consistency of the Bayesian posterior distribution.

Our proposed approach can be viewed as an online extension of the BRO framework in Wu et al. (2018), where the risk function is taken as the expectation and the reformulated problem is updated over time stages. We consider the same problem as Song and Shanbhag (2019), but differ in two key aspects: first, we take a Bayesian approach to estimate the input parameter and solve the Bayesian average problem, whereas they estimate the input parameter by MLE and solve the problem with the plug-in MLE; second, we also consider the case of decision-dependent input data, whereas they only consider decision-independent input data. We note that decision-dependent input data have been studied in machine learning (e.g. Brückner et al. (2012), Juan et al. (2020)) and stochastic optimization (e.g. Lappas and Gounaris (2018), Luo and Mehrotra (2020) and Dmitriy and Lin (2020)). However, their problem settings are in general different from the setting in this paper.

Our contributions are summarized as follows. First, we propose a new approach to online simulation optimization with streaming data. Second, we show the convergence of our approach in the decision-dependent case. Third, we show the consistency of the Bayesian posterior distribution with endogenous data, which might be of independent interest to Bayesian estimation outside the scope of this paper. Our theoretical results are also verified in our numerical experiments.

2 ONLINE SIMULATION OPTIMIZATION WITH DECISION-INDEPENDENT INPUT DATA

In this section, we consider simulation optimization with decision-independent input distribution, i.e.,

$$\min_{x \in \mathcal{X}} H(x) := \mathbb{E}_{\xi \sim f^c} [h(x, \xi)], \quad (1)$$

where $\mathcal{X} \subset \mathbb{R}^d$ is the decision space assumed to be hyper-rectangular, $\xi \in \Xi \subset \mathbb{R}^m$ is a random vector representing the randomness in the simulation, $h: \mathbb{R}^d \times \mathbb{R}^m \rightarrow \mathbb{R}$ is a function evaluated through simulation. The expectation is taken with respect to (w.r.t.) the distribution of ξ , which we denote by f^c . The true distribution f^c is usually unknown, and hence the true problem (1) is not available and calls for reformulation.

We assume that f^c belongs to a parameterized family of distributions $\{f(\cdot; \theta), \theta \in \Theta\}$, where $\Theta \subset \mathbb{R}^l$ is the parameter space and θ^c denotes the true but unknown input parameter. At each time stage t , we observe

a batch of data $\mathbf{y}_t = \{y_{t,j}, j = 1, \dots, D\}$ of size D , where $\{y_{t,j}\}$ are independent and identically distributed (i.i.d.) with $f(\cdot; \theta^c)$. To estimate θ^c , we take a Bayesian approach by viewing the unknown parameter as a random variable θ and assume a prior distribution π_0 on θ . Then the posterior distribution is updated according to the Bayes rule: $\pi_t(\theta) = \frac{\pi_{t-1}(\theta)f(\mathbf{y}_t; \theta)}{\int \pi_{t-1}(\theta)f(\mathbf{y}_t; \theta)d\theta}$. With the posterior distribution summarizing all the information contained in the data so far, we consider a Bayesian average of the original objective function:

$$\min_{x \in \mathcal{X}} \mathbb{E}_{\theta \sim \pi_t} [H(x, \theta)], \quad (2)$$

where $H(x, \theta) = \mathbb{E}_{\xi \sim f(\cdot; \theta)} [h(x, \xi)]$. To solve (2), we use the SGD algorithm for K steps within each time stage. Note that (2) is the same as the BRO formulation in Wu et al. (2018) with the risk function taken as the expectation. The algorithm is presented below.

Algorithm 1: Online Simulation Optimization with Decision-Independent Input Data

- At time 0, choose an initial decision x_0 , prior distribution $\pi_0(\theta)$, and step size sequence $\{a_{t,j}, t = 1, 2, \dots; j = 0, \dots, K-1\}$.
- At time $t \geq 1$, a batch of data $y_{t,1}, \dots, y_{t,D} \stackrel{\text{i.i.d.}}{\sim} f(\cdot; \theta^c)$ arrives.
 1. posterior update: $\pi_t(\theta) \propto \pi_{t-1}(\theta) \cdot \prod_{j=1}^D f(y_{t,j}; \theta)$.
 2. SGD: take K SGD steps in x , i.e.,
 - 2.1. set $x_{t,0} := x_{t-1}$.
 - 2.2. for $j = 0, \dots, K-1$, draw $\theta_{t,j} \sim \pi_t(\theta)$ and $\xi_{t,j} \sim f(\cdot; \theta_{t,j})$, and update the solution

$$x_{t,j+1} := \text{Proj}_{\mathcal{X}} \{x_{t,j} - a_{t,j} \nabla_x h(x_{t,j}, \xi_{t,j})\}.$$

2.3 set $x_t := x_{t,K}$ and $t := t+1$; go to Step 2.

- Output x_T as the final decision for some terminal time T or when some stopping criteria is satisfied.

In the SGD step of Algorithm 1, we use infinitesimal perturbation analysis (IPA, refer to Fu (2008)) to estimate the gradient of the objective function in (2). We assume we can interchange the expectation and differentiation, i.e., $\nabla_x \mathbb{E}_{\theta \sim \pi_t} [\mathbb{E}_{\xi \sim \theta_t} [h(x, \xi_t)]] = \mathbb{E}_{\theta \sim \pi_t} [\mathbb{E}_{\xi \sim \theta_t} [\nabla_x h(x, \xi_t)]]$. Then $\nabla_x h(x, \xi)$ with $\xi \sim f(\cdot; \theta)$, where $\theta \sim \pi_t(\theta)$, is an unbiased gradient estimator of the objective function in (2). This assumption is satisfied with bounded gradient, where the interchange is justified by the dominated convergence theorem.

3 ONLINE SIMULATION OPTIMIZATION WITH DECISION-DEPENDENT INPUT DATA

In this section, we consider simulation optimization with decision-dependent input distribution, i.e.,

$$\min_{x \in \mathcal{X}} H(x, \theta^c) = \mathbb{E}_{\xi \sim f(\cdot; x, \theta^c)} [h(x, \xi)], \quad (3)$$

where the random variable ξ follows the decision-dependent input distribution $f(\cdot; x, \theta^c)$, and other notations are defined as in (1). In the online setting, the decision is updated at each time t and denoted by x_t , while the data batch of that stage depends on x_t . Specifically, at time stage t we observe data $\mathbf{y}_t = \{y_{t,j}, j = 1, \dots, D\}$ of size D , where $\{y_{t,j}\}_j$ are i.i.d. from $f(\cdot; x_t, \theta^c)$. Note that data batches $\{\mathbf{y}_t\}_t$ are correlated and differently distributed across stages since the decision x_t is updated based on previous decision(s). Regardless of the correlation and non-stationarity, we use the data batches to update the posterior distribution according to:

$$\pi_t(\theta) = \frac{\pi_{t-1}(\theta)f(\mathbf{y}_t; x_t, \theta)}{\int \pi_{t-1}(\theta)f(\mathbf{y}_t; x_t, \theta)d\theta}.$$

Similarly as the previous section, at time stage t we consider optimizing the Bayesian average of the objective function:

$$\mathbb{E}_{\theta \sim \pi_t} [H(x, \theta)] = \mathbb{E}_{\theta \sim \pi_t} \mathbb{E}_{\xi \sim f(\cdot; x_t, \theta)} [h(x_t, \xi)]. \quad (4)$$

To apply SGD on (4), we need to estimate the gradient of the objective function. Assuming we can interchange integration and differentiation, we have

$$\begin{aligned}\nabla_x \mathbb{E}_\theta[H(x, \theta)] &= \mathbb{E}_\theta \mathbb{E}_\xi[\nabla_x h(x, \xi) f(\xi; x, \theta)] + \mathbb{E}_\theta \left[\int_{\Xi} h(x, \xi) \nabla_x f(\xi; x, \theta) d\xi \right] \\ &= \mathbb{E}_\theta \mathbb{E}_\xi[\nabla_x h(x, \xi) f(\xi; x, \theta)] + \mathbb{E}_\theta \left[\int_{\Xi} h(x, \xi) \frac{\nabla_x f(\xi; x, \theta)}{f(\xi; x, \theta)} f(\xi; x, \theta) d\xi \right] \\ &= \mathbb{E}_\theta \mathbb{E}_\xi \left[\nabla_x h(x, \xi) + h(x, \xi) \frac{\nabla_x f(\xi; x, \theta)}{f(\xi; x, \theta)} \right].\end{aligned}$$

For simplicity, we denote $\hat{f}_t(\cdot; x) = E_{\theta \sim \pi_t}[f(\cdot; x, \theta)]$. From the equation above, it is easy to see an unbiased gradient estimator of (4) is as follows:

$$\nabla_x h(x_t, \xi_t) + h(x_t, \xi_t) \frac{\nabla_x \hat{f}_t(\xi_t; x_t)}{\hat{f}_t(\xi_t; x_t)}.$$

We now present our algorithm below.

Algorithm 2: Online Simulation Optimization with Decision-Dependent Input Data

- At time 0, choose an initial decision x_0 , prior distribution $\pi_0(\theta)$, and step size sequence $\{a_{t,j}, t = 1, 2, \dots; j = 0, \dots, K-1\}$.
- At time $t \geq 1$, a batch of data $y_{t,1}, \dots, y_{t,D} \stackrel{\text{i.i.d.}}{\sim} f(\cdot; x_t, \theta^c)$ arrives.
 1. posterior update: $\pi_t(\theta) \propto \pi_{t-1}(\theta) \cdot \prod_{j=1}^D f(y_{t,j}; \theta)$.
 2. SGD: take K SGD steps in x , i.e.,
 - 2.1. set $x_{t,0} := x_{t-1}$.
 - 2.2. for $j = 0, \dots, K-1$, draw $\theta_{t,j} \sim \pi_t(\theta)$ and $\xi_{t,j} \sim f(\cdot; x_{t,j}, \theta_{t,j})$, and update the solution

$$x_{t,j+1} := \text{Proj}_{\mathcal{X}} \left\{ x_{t,j} - a_{t,j} \left(\nabla_x h(x_{t,j}, \xi_{t,j}) + h(x_{t,j}, \xi_{t,j}) \frac{\nabla_x \hat{f}_t(\xi_{t,j}; x_{t,j})}{\hat{f}_t(\xi_{t,j}; x_{t,j})} \right) \right\}. \quad (5)$$

2.3 set $x_t := x_{t,K}$ and $t := t+1$; go to Step 2.

- Output x_T as the final decision for some terminal time T or when some stopping criteria is satisfied.

4 CONVERGENCE ANALYSIS

In this section, we focus on the decision-dependent case and theoretically study the convergence behavior of Algorithm 2. Before proceeding to detailed analysis, we first introduce the following notations and assumptions. Define the Bayesian prior π_0 on $(\Theta, \mathcal{B}(\Theta))$, where $\mathcal{B}(\Theta)$ is the Borel σ -algebra on Θ . Let $\mathcal{Y}$ denote the data (observation) space. Define a probability space $(\Omega, \mathcal{F}, \mathbb{P})$, such that $\mathbb{P}(\theta \in A) = \pi_0(A), \forall A \in \mathcal{B}(\Theta)$ and $\mathbb{P}(y_t \in C | x_t, \theta) = \int_C f(y; x_t, \theta) dy, \forall C \in \mathcal{B}(\mathcal{Y})$. Furthermore, define the σ -filtration $\mathcal{F}_t = \sigma\{x_\tau, y_\tau, \tau \leq t\}$. Without loss of generality, we assume at each time stage the data batch size $D = 1$ and the number of SGD steps $K = 1$.

Assumption 1

- The parameter space Θ is finite, i.e., $\Theta = \{\theta_1, \dots, \theta_k\}$. Moreover, $\theta^c \in \Theta$.
- The prior distribution $\pi_0(\theta^c) > 0$.
- The decision space $\mathcal{X} = \{x | l_i \leq x^{(i)} \leq u_i, \forall i = 1, \dots, d\}$, where $-\infty < l_i < u_i < \infty$ and $x^{(i)}$ is the i^{th} component of x .
- The observation space $\mathcal{Y}$ is bounded.

- $h(x, y)$ is C -Lipschitz continuous in x and uniformly in y , i.e., $\|\nabla_x h(x, y)\|_2 \leq C, \forall x \in \mathcal{X}, y \in \mathcal{Y}$, for some $C > 0$.
- $f(y; x, \theta)$ is continuously differentiable in x and has bounded first order derivative, i.e.,

$$\|\nabla_x f(y; x, \theta)\|_2 \leq C', \forall x, y, \theta, \text{ for some } C' > 0.$$

- The step size a_t satisfies $\sum_{t=1}^{\infty} a_t = \infty, \lim_{t \rightarrow \infty} a_t = 0, a_t > 0, \forall t > 0$.

The assumptions above are regularity conditions and easy to be verified in practice. The correlated and differently distributed data $\{y_t\}$ pose a great challenge to analyzing the consistency of the Bayesian posterior distribution π_t . To prove the consistency, we first show the following intermediate result.

Lemma 2 Suppose Assumption 1 holds. Recall $\hat{f}_t(\cdot; x) = \sum_{\theta} \pi_t(\theta) f(\cdot; x, \theta)$, and denote $f^*(\cdot; x) := f(\cdot; x, \theta^c)$, for any $x \in \mathcal{X}$. At decision x_{t+1} , the Kullback-Leibler (K-L) divergence between f^* and $\hat{f}_t$ is defined as d_t , i.e., $d_t := KL(f^*(\cdot; x_{t+1}) \parallel \hat{f}_t(\cdot; x_{t+1}))$. Then we have

$$\lim_{t \rightarrow \infty} d_t = 0 \quad \text{and} \quad \sum_{t=1}^{\infty} d_t < \infty, \text{ a.s..}$$

Please refer to Appendix A.1 for detailed proof of Lemma 2. Intuitively, Lemma 2 implies that with more observation data even at different decisions, we know more about the true parameter θ^c and are able to provide a more precise estimation of the density f^* at the next decision. Moreover, if we know that each θ is identifiable as rigorously defined in the following assumption, we can further prove the consistency of π_t regardless of the correlation and non-stationarity of the observation data.

Assumption 3 (Linear Independence) For any x in $\mathcal{X}$, $\{f(\cdot; x, \theta_i)\}_{i=1}^k$ are linearly independent in $\mathcal{Y}$, i.e.,

$$c_1 f(y; x, \theta_1) + c_2 f(y; x, \theta_2) + \dots + c_k f(y; x, \theta_k) = 0, \quad \forall y \in \mathcal{Y} \Rightarrow c_1 = c_2 = \dots = c_k = 0.$$

Assumption 3 intuitively requires that for any decision x , the observation distributions generated from different θ 's are distinguishable. This is usually satisfied in practice and for many distributions. For example, let $f(y; x, \theta_i) = \theta_i x \exp(\theta_i x)$ be the density function of the exponential distribution with rate $\theta_i x$. One can verify that the Wronskian Determinant is nonzero when θ_i 's are distinct, which directly implies the linear independence of $\{f(y; x, \theta_i)\}_i$. We then have the following proposition on the strong consistency of the posterior distribution $\{\pi_t\}$.

Proposition 4 Under Assumptions 1 and 3, $\pi_t(\theta) \rightarrow \delta_{\theta^c}, \text{ a.s. as } t \rightarrow \infty$.

Please refer to Appendix A.2 for detailed proof. Proposition 4 guarantees that although our observation depends on our current decision, it can provide enough information to make sure the posterior distribution will finally concentrate on the true parameter. The consistency of π_t ensures that our gradient estimator is accurate enough, and Algorithm 2 can converge.

We study the asymptotic behavior of Algorithm 2 by the ordinal differential equation (ODE) method (refer to Kushner and Yin (2003)). The main idea is that SGD can be viewed as a noisy discretization of an ODE. Under certain conditions, the noise in SGD averages out asymptotically, such that the SGD iterates converge to the solution trajectory of an ODE. Before proceeding to our main convergence result, we introduce some more notations. Let $N(t)$ be the unique n such that $\sum_{i=0}^{n-1} a_i \leq t < \sum_{i=0}^n a_i$ for $t > 0$, and 0 for $t < 0$. Define the interpolated continuous process X as $X(0) = x_0$ and $X(t) = x_{N(t)}$ for any $t > 0$, and the shifted process as $X^n(s) = X(s + \sum_{i=0}^{n-1} a_i)$ for any $s \geq 0$. We then have the following theorem.

Theorem 5 Let $D^d[0, \infty)$ be the space of $\mathbb{R}^d$ -valued operators which are right continuous and have left-hand limits for each dimension. Under Assumptions 1 and 3, for each subsequence of $\{X^n(\cdot)\}_n$, there exists a further subsequence $\{X^{n_k}(\cdot)\}_{n_k}$ and a process $X^*(\cdot)$ such that $X^{n_k}(\cdot) \Rightarrow X^*(\cdot)$ in the weak sense as $t \rightarrow \infty$.

in the space $D^d[0, \infty)$, where $X^*(\cdot)$ satisfies the following ODE:

$$\dot{X} = -\nabla H(X, \theta^c) + z, \quad z \in -C(X), \quad X(0) = x_0, \quad (6)$$

where z is the minimum force needed to keep the trajectory of the ODE in $\mathcal{X}$. Let $L_{\mathcal{X}}$ be the set of limit points of (6) in $\mathcal{X}$. Then there exist $\mu_n \rightarrow 0$ and $T_n \rightarrow \infty$ such that

$$\lim_n P \left\{ \sup_{t \leq T_n} \text{Dist}(X^n(t), L_{\mathcal{X}}) \geq \mu_n \right\} = 0,$$

where $\text{Dist}(x, A) = \inf_{y \in A} \|x - y\|_2$ for any set A and point $x \in \mathcal{X}$.

Remark 6 The set $C(X)$ in (6) is defined as follows. For $X \in \mathcal{X}^0$, the interior of $\mathcal{X}$, $C(X)$ contains only the zero element; for $X \in \alpha\mathcal{X}$, the boundary of $\mathcal{X}$, let $C(X)$ be the infinite convex cone generated by the outer normals at X of the faces on which X lies.

Remark 7 Theorem 5 shows the weak convergence of Algorithm 2. The SGD iterates specified in (5) approaches the solution trajectory of the ODE (6) and eventually converges to a limit point of the ODE, which is a point x^* satisfying $\nabla H(x^*, \theta^c) = 0$ if the point is in the interior of $\mathcal{X}$. Hence, such a point is a stationary point of problem (3) and can be a local optimal solution if it is stable. The weak convergence result implies that once the trajectory enters the domain of attraction of a local optimal solution, the chance of escaping from it goes to 0 in the limit.

Now we prove Theorem 5 below.

Proof. Recall that at time $t + 1$, Algorithm 2 takes the following update

$$x_{t+1} = x_t - a_t \left(\nabla_x h(x_t, \xi_t) + h(x_t, \xi_t) \frac{\nabla_x \hat{f}_t(\xi_t; x_t)}{\hat{f}_t(\xi_t; x_t)} \right) + a_t z_t,$$

where $\xi_t \sim f(\cdot; x_t, \theta_t)$, $\theta_t \sim \pi_t$ and $a_t z_t$ is the projection term. Note that

$$\begin{aligned} \mathbb{E}[\nabla_x h(x_t, \xi_t) | \mathcal{F}_t] &= \mathbb{E}[\nabla_x h(x_t, \xi_t) | x_t, \pi_t] = \mathbb{E}[\mathbb{E}[\nabla_x h(x_t, \xi_t) | \theta_t] | x_t, \pi_t] \\ &= \mathbb{E} \left[\int_{\mathcal{Y}} \nabla_x h(x_t, y) f(y; x_t, \theta_t) dy | x_t, \pi_t \right] = \sum_{\theta} \pi_t(\theta) \int_{\mathcal{Y}} \nabla_x h(x_t, y) f(y; x_t, \theta) dy \\ &= \int_{\mathcal{Y}} \sum_{\theta} \pi_t(\theta) \nabla_x h(x_t, y) f(y; x_t, \theta) dy = \mathbb{E}_{y \sim \hat{f}_t(\cdot; x_t)} \nabla_x h(x_t, y) \\ &= \mathbb{E}_{y \sim f^*(\cdot; x_t)} \nabla_x h(x_t, y) + \left(\mathbb{E}_{y \sim \hat{f}_t(\cdot; x_t)} \nabla_x h(x_t, y) - \mathbb{E}_{y \sim f^*(\cdot; x_t)} \nabla_x h(x_t, y) \right) \\ &= \mathbb{E}_{y \sim f^*(\cdot; x_t)} \nabla_x h(x_t, y) + \beta_{t,1}, \end{aligned}$$

where $f^*(\cdot; x) = f(\cdot; x, \theta^c)$ for $x \in \mathcal{X}$, and $\beta_{t,1} = \mathbb{E}_{y \sim \hat{f}_t(\cdot; x_t)} \nabla_x h(x_t, y) - \mathbb{E}_{y \sim f^*(\cdot; x_t)} \nabla_x h(x_t, y)$ is the bias. Further, we have

$$\begin{aligned} \mathbb{E} \left[h(x_t, \xi_t) \frac{\nabla_x \hat{f}_t(\xi_t; x_t)}{\hat{f}_t(\xi_t; x_t)} | \mathcal{F}_t \right] &= \mathbb{E} \left[h(x_t, \xi_t) \frac{\nabla_x \hat{f}_t(\xi_t; x_t)}{\hat{f}_t(\xi_t; x_t)} | x_t, \pi_t \right] = \int_{\mathcal{Y}} h(x_t, y) \nabla_x \hat{f}_t(y; x_t) dy \\ &= \int_{\mathcal{Y}} h(x_t, y) \nabla_x f^*(y; x_t) dy + \left(\int_{\mathcal{Y}} h(x_t, y) \nabla_x \hat{f}_t(y; x_t) dy - \int_{\mathcal{Y}} h(x_t, y) \nabla_x f^*(y; x_t) dy \right) \\ &= \int_{\mathcal{Y}} h(x_t, y) \nabla_x f^*(y; x_t) dy + \beta_{t,2}, \end{aligned}$$

where $\beta_{t,2} = \int_{\mathcal{Y}} h(x_t, y) \nabla_x \hat{f}_t(y; x_t) dy - \int_{\mathcal{Y}} h(x_t, y) \nabla_x f^*(y; x_t) dy$. Note that

$$\int_{\mathcal{Y}} h(x_t, y) \nabla_x f^*(y; x_t) dy + \mathbb{E}_{y \sim f^*(\cdot; x_t)} \nabla_x h(x_t, y) = \nabla_x H(x, \theta^c),$$

and we can rewrite the update as

$$x_{t+1} = x_t - a_t \nabla_x H(x_t, \theta^c) - a_t \beta_{t,1} - a_t \beta_{t,2} - a_t \delta M_t + a_t z_t,$$

where $\delta M_t = h(x_t, \xi_t) - \mathbb{E}[\nabla_x h(x_t, \xi_t) | \mathcal{F}_t]$ is a martingale difference sequence. Suppose we can show $\lim_{t \rightarrow \infty} \beta_{t,1} = 0$ a.s. and $\lim_{t \rightarrow \infty} \beta_{t,2} = 0$ a.s., then the rest of the update is exactly the discretization of ODE (6). Utilizing the consistency of π_t , we show the two bias terms $\beta_{t,1}$ and $\beta_{t,2}$ vanish in the limit, i.e., $\lim_{t \rightarrow \infty} \beta_{t,1} = 0$ a.s. (see Lemma 8 in the Appendix A.3) and $\lim_{t \rightarrow \infty} \beta_{t,2} = 0$ a.s. (see Lemma 9 in the Appendix A.4). Then Theorem 5 is proved by a straightforward application of Theorem 2.1 in Chapter 7.2 in Kushner and Yin (2003). $\square$

5 NUMERICAL EXPERIMENTS

We carry out numerical experiments on a simple quadratic problem: $h(x, \xi) = (x - 5)^2 + 0.5\xi x$. In the decision-independent case we assume $\xi \sim \mathcal{N}(\theta^c, \sigma^2)$, and in the decision-dependent case $\xi \sim \mathcal{N}(x + \theta^c, \sigma^2)$. In both cases, $\sigma = 4$, $\theta^c = 9$, and the parameter space $\Theta = \{1, 2, \dots, 10\}$.

- **Experiment 1.** In this experiment, we demonstrate the performance of Algorithm 1 in the decision-independent case. It is easy to check $H(x, \theta) = x^2 - 5.5x + 25$, and the true optimal decision is taken at $x^* = 2.75$. At each time t , the gradient estimator in Algorithm 1 is $\nabla_x h(x_t, \xi_t) = 2x_t - 10 + 0.5\xi_t$. We use a uniform prior distribution on the parameter. We set the initial solution $x_0 = 20$, batch data size $D = 1$, number of SGD steps $K = 1$, step size $a_t = \frac{2}{t+5}$, and the total number of time stages $T = 1000$. We run Algorithm 1 for 100 times on the problem. The mean and standard deviation of the solution error $|x_t - x^*|$ over time is shown in Figure 1. Figure 1 shows that with decreasing step size, the solution sequence in Algorithm 1 converges to the true optimal solution.

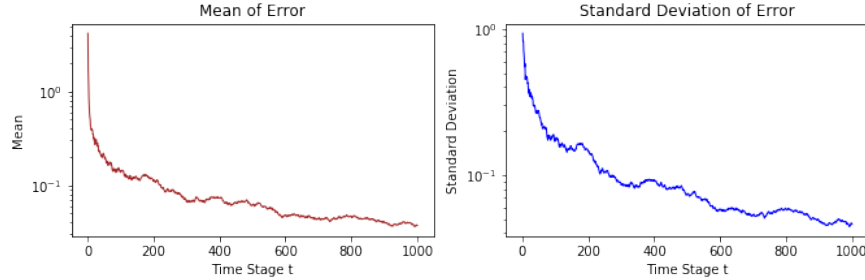


Figure 1: Mean and standard deviation of $|x_t - x^*|$ of 100 runs of Algorithm 1 with $a_t = \frac{2}{t+5}$, $D = 1$, $K = 1$.

In the decision-dependent case, it is easy to verify that $H(x, \theta) = (x - 5)^2 + 0.5(x + \theta^c)x = 1.5x^2 - 5.5x + 25$ and the true optimal decision is $x^* = \frac{11}{6}$. The gradient estimator in Algorithm 2 at each time t is $\nabla_x h(x_t, \xi_t) + h(x_t, \xi_t) \frac{\nabla_x \hat{f}_t(\xi_t; x_t)}{\hat{f}_t(\xi_t; x_t)}$, which can be computed as $(2x_t - 10 + 0.5\xi_t) + ((x_t - 5)^2 + 0.5\xi_t x_t) \frac{\sum_{\theta} \pi_t(\theta) \cdot \nabla_x f(\xi_t; x_t, \theta_t)}{\sum_{\theta} \pi_t(\theta) \cdot f(\xi_t; x_t, \theta_t)}$, where $f(\xi_t; x_t, \theta_t) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(\xi_t - (x_t + \theta_t))^2}{2\sigma^2}\right)$, $\nabla_x f(\xi_t; x_t, \theta_t) = \frac{\xi_t - (x_t + \theta_t)}{\sqrt{2\pi}\sigma^3} \exp\left(-\frac{(\xi_t - (x_t + \theta_t))^2}{2\sigma^2}\right)$. We use a uniform prior distribution on the parameter. We set the step size $a_t = \frac{2}{t+5}$ and initial solution $x_0 = 20$. We run Algorithm 2 for 100 times to solve the problem (2).

- **Experiment 2.** In this experiment, we first demonstrate the convergence of the posterior distribution π_t . We set the batch data size $D = 1$, the number of SGD steps at each time stage as $K = 1$, and the

total number of time stages $T = 500$. Figure 2 shows that mean of $\pi_t(\theta^c)$, the posterior probability mass on the true parameter θ^c , increases and converges to 1, and the 95% confidence interval of $\pi_t(\theta^c)$ goes to 0. This empirically verifies our theoretical result Lemma 4. Next, Figure 3 shows that the mean and standard deviation of the solution error $|x_t - x^*|$ over $T = 2000$ time stages, which demonstrates the convergence of the solution sequence to the true optimal solution.

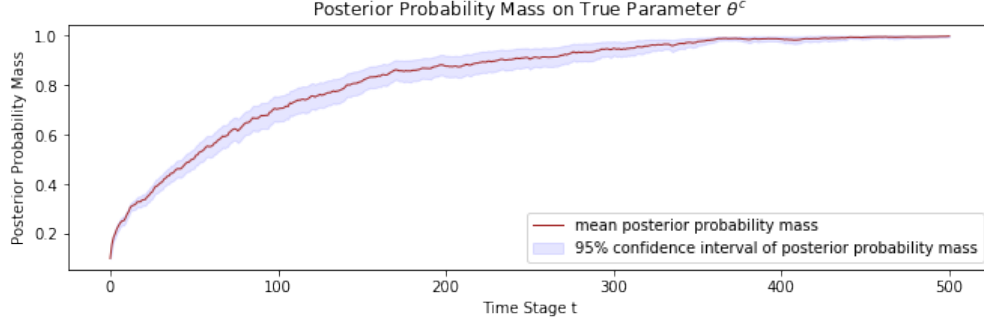


Figure 2: Mean and 95% confidence interval of $\pi_t(\theta^c)$ of 100 runs of Algorithm 2. $a_t = \frac{2}{t+5}$, $D = 1$, $K = 1$.

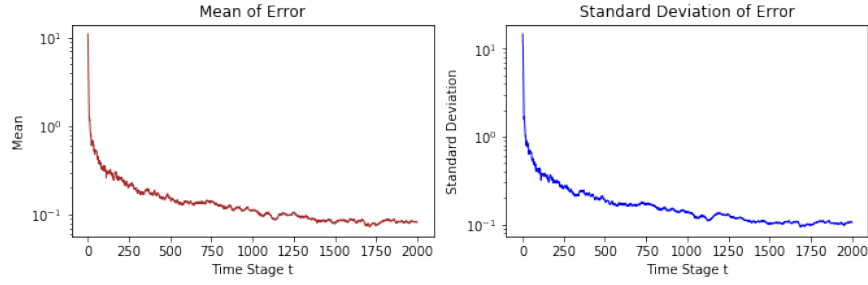


Figure 3: Mean and standard deviation of $|x_t - x^*|$ of 100 runs of Algorithm 2. $a_t = \frac{2}{t+5}$, $D = 1$, $K = 1$.

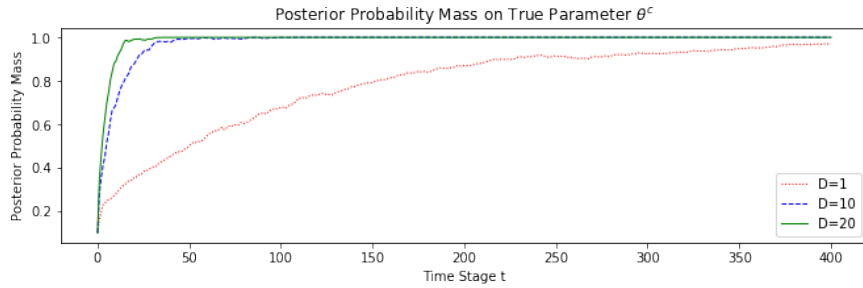


Figure 4: Mean of $\pi_t(\theta^c)$ of 100 runs of Algorithm 2 under different batch size D . $a_t = \frac{2}{t+5}$, $K = 1$.

- **Experiment 3.** In this experiment, we empirically study the impact of the batch size D and number of SGD steps K on the convergence rate. Figure 4 plots $\pi_t(\theta^c)$ over $T = 400$ time stages under different data batch size D . It shows that as we observe more data at each time stage, the posterior distribution converges faster to a delta function concentrated on the true parameter θ^c . Figure 5 plots the mean of the solution error $|x_t - x^*|$ of 100 runs of Algorithm 2 under different batch size D . As a benchmark, we replace the gradient estimator in Algorithm 2 with the gradient estimator that uses the true parameter θ^c (denoted by $D = \infty$), i.e., $\nabla_x h(x, \xi) + h(x, \xi) \frac{\nabla_x f(\xi; x, \theta^c)}{f(\xi; x, \theta^c)}$, where $\xi \sim f(\cdot; x, \theta^c)$.

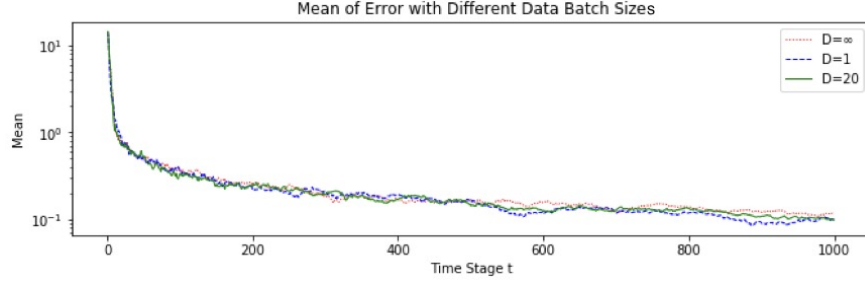


Figure 5: Mean of $|x_t - x^*|$ of 100 runs of Algorithm 2 under different batch size D . $a_t = \frac{2}{t+5}$, $K = 1$.

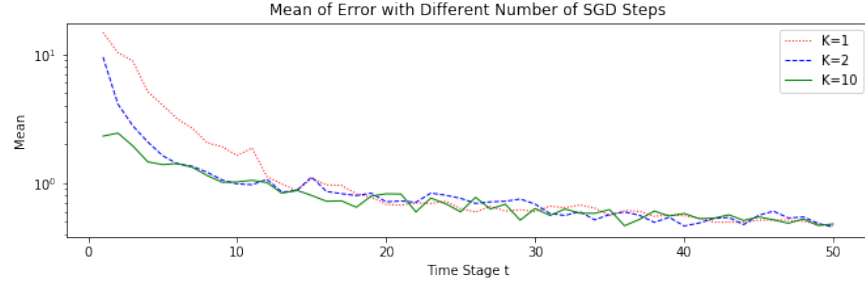


Figure 6: Mean of $|x_t - x^*|$ of 100 runs of Algorithm 2 under different number of SGD steps K . $D = 1$, $a_t = \frac{2}{t+5}$.

Figure 4 and Figure 5 together show that there is no significant difference in the convergence rate under different data batch sizes, even though the posterior distribution converges faster with larger data batch size. It implies that the Bayesian average of the objective function (4) in this example is a good estimate of the true objective function regardless of the inaccuracy of the posterior distribution at the beginning time stages. Figure 6 plots the solution error of Algorithm 2 under different number of SGD steps K in each time stage. It shows that, when we take more SGD steps at each time stage, the algorithm converges faster in the beginning but eventually behaves similarly.

ACKNOWLEDGMENTS

The authors gratefully acknowledge the support by the Air Force Office of Scientific Research under Grant FA9550-19-1-0283.

A APPENDICES

A.1 Proof of Lemma 2

Proof. Define $w_t = -\log \pi_t(\theta^c)$. One can easily verify that $w_t \geq 0$. Then we have

$$\begin{aligned} \mathbb{E}[w_{t+1}] &= \mathbb{E}[\mathbb{E}[w_{t+1} | \mathcal{F}_t, x_{t+1}]] \\ &= \mathbb{E} \left[\mathbb{E} \left[-\log \frac{\pi_t(\theta^c) f(y_{t+1}; x_{t+1}, \theta^c)}{\sum_{\theta} \pi_t(\theta) f(y_{t+1}; x_{t+1}, \theta)} \middle| \mathcal{F}_t, x_{t+1} \right] \right] \\ &= \mathbb{E} \left[-\log \pi_t(\theta^c) - \mathbb{E} \left[\log \frac{f(y_{t+1}; x_{t+1}, \theta^c)}{\sum_{\theta} \pi_t(\theta) f(y_{t+1}; x_{t+1}, \theta)} \middle| \mathcal{F}_t, x_{t+1} \right] \right] \\ &= \mathbb{E}[w_t] - \mathbb{E}[KL(f^*(\cdot; x_{t+1}) || \hat{f}_t(\cdot; x_{t+1}))]. \end{aligned}$$

This implies that $\mathbb{E}[d_t] = \mathbb{E}[w_t] - \mathbb{E}[w_{t+1}]$. For any $T > 0$, we have

$$\sum_{t=0}^T \mathbb{E}[d_t] = \sum_{t=0}^T \mathbb{E}[w_t] - \mathbb{E}[w_{t+1}] = w_0 - \mathbb{E}[w_{T+1}] \leq w_0 < \infty.$$

Then we have $\sum_{t=0}^{\infty} \mathbb{E}[d_t] \leq w_0$. $\forall \varepsilon > 0$, we have

$$\sum_{t=0}^{\infty} \mathbb{P}(d_t \geq \varepsilon) \leq \frac{1}{\varepsilon} \sum_{t=0}^{\infty} \mathbb{E}[d_t] < \infty.$$

By Borel-Cantelli Lemma, we know that $\mathbb{P}(d_t \geq \varepsilon, i.o.) = 0$, which further implies $\lim_{t \rightarrow \infty} d_t = 0, a.s..$ Moreover, since $d_t \geq 0$, by Tonelli's Theorem, we have

$$\mathbb{E} \left[\sum_{t=0}^{\infty} d_t \right] = \sum_{t=0}^{\infty} \mathbb{E}[d_t] \leq w_0.$$

Since $\sum_{t=0}^{\infty} d_t$ has bounded expectation, it must be finite almost surely. $\square$

A.2 Proof of Proposition 4

Proof. Without loss of generality, we assume $\theta^c = \theta^1$. Recall that $f^*(y; x_{t+1}) = f^*(y; x_{t+1}, \theta^1)$ and $\hat{f}_t(y; x_{t+1}) = \sum_i \pi_t(\theta^i) f(y; x_{t+1}, \theta^i)$. Then

$$f^*(y; x_{t+1}) - \hat{f}_t(y; x_{t+1}) = (1 - \pi_t(\theta^1)) f(y; x_{t+1}, \theta^1) + \sum_{i>1} \pi_t(\theta^i) f(y; x_{t+1}, \theta^i).$$

Note that $(\pi_t(\theta^1), \dots, \pi_t(\theta^k))$ is bounded. Take any convergence subsequence $\{t_1, t_2, \dots\}$ with limit $(p_1^*, p_2^*, \dots, p_k^*)$. From this subsequence, take a further subsequence $\{\tau_1, \tau_2, \dots\}$ such that x_{τ_i+1} converges to x' . We can find this subsequence since $\mathcal{X}$ is bounded. Then

$$f^*(y; x_{\tau_i+1}) - \hat{f}_{\tau_i}(y; x_{\tau_i+1}) \rightarrow (1 - (p_1^*)) f(y; x', \theta^1) + \sum_{i>1} p_i^* f(y; x', \theta^i).$$

Moreover, since K-L divergence dominates total variation distance between two distributions, we have

$$\int_{\mathcal{Y}} |f^*(y; x_{t+1}) - \hat{f}_t(y; x_{t+1})| dy \leq d_t. \quad (7)$$

From (7) and Lemma 2, we know that $\int_{\mathcal{Y}} |f^*(y; x_{t+1}) - \hat{f}_t(y; x_{t+1})| dy \rightarrow 0, a.s..$ By dominant convergence theorem, we have $\int_{\mathcal{Y}} |(1 - (p_1^*)) f(y; x^*, \theta^1) + \sum_{i>1} p_i^* f(y; x^*, \theta^i)| dy = 0$, which implies

$$(1 - (p_1^*)) f(y; x^*, \theta^1) + \sum_{i>1} p_i^* f(y; x^*, \theta^i) = 0, \forall y.$$

By linear independence, we know $p_1^* = 1, p_2^* = \dots = p_k^* = 0$. Since every convergent subsequence of $\{(\pi_t(\theta^1), \dots, \pi_t(\theta^k))\}_t$ have the same limit, we know $\pi_t(\theta) \rightarrow \delta_{\theta^c}$. $\square$

A.3 Lemma 8 and Proof

Lemma 8 Under Assumptions 1 and 3, $\lim_{t \rightarrow \infty} \beta_{t,1} = 0$ a.s..

Proof. We bound $|\beta_{t,1}|$ as follows.

$$\begin{aligned} |\beta_{t,1}| &= |\mathbb{E}_{y \sim \hat{f}_t(\cdot; x_t)} \nabla_x h(x_t, y) - \mathbb{E}_{y \sim f^*(\cdot; x_t)} \nabla_x h(x_t, y)| \\ &\leq \max |\nabla_x h(x, y)| \int_{\mathcal{Y}} |f^*(y; x_t) - \hat{f}_t(y; x_t)| dy \\ &\leq C \int_{\mathcal{Y}} |f^*(y; x_t) - f^*(y; x_{t+1})| + |f^*(y; x_{t+1}) - \hat{f}_t(y; x_{t+1})| + |\hat{f}_t(y; x_t) - \hat{f}_t(y; x_{t+1})| dy. \end{aligned}$$

By Lipschitz condition, $|f^*(y; x_t) - f^*(y; x_{t+1})| \leq L|x_t - x_{t+1}| \leq LD a_t$ and $|\hat{f}_t(y; x_t) - \hat{f}_t(y; x_{t+1})| \leq LD a_t$ for some $L, D > 0$. Thus, given the boundedness of $\mathcal{Y}$, there exists some constant $C'' > 0$, such that

$$\int_{\mathcal{Y}} |f^*(y; x_t) - f^*(y; x_{t+1})| + |\hat{f}_t(y; x_t) - \hat{f}_t(y; x_{t+1})| dy \leq C'' a_t. \quad (8)$$

Moreover, since K-L divergence dominates total variation distance between two distributions, we have

$$\int_{\mathcal{Y}} |f^*(y; x_{t+1}) - \hat{f}_t(y; x_{t+1})| dy \leq d_t. \quad (9)$$

Then combine (8) and (9), we know that $\lim_{t \rightarrow \infty} |\beta_{t,1}| = 0$ a.s.. $\square$

A.4 Lemma 9 and Proof

Lemma 9 Under Assumptions 1 and 3, $\lim_{t \rightarrow \infty} \beta_{t,2} = 0$ a.s..

Proof. We bound $\beta_{t,2}$ as follows.

$$\begin{aligned} |\beta_{t,2}| &= \left| \int_{\mathcal{Y}} h(x_t, y) \nabla_x \hat{f}_t(y; x_t) dy - \int_{\mathcal{Y}} h(x_t, y) \nabla_x f^*(y; x_t) dy \right| \\ &= \left| \int_{\mathcal{Y}} h(x_t, y) \left(\nabla_x \hat{f}_t(y; x_t) - \nabla_x f^*(y; x_t) \right) dy \right| \\ &= \left| \int_{\mathcal{Y}} h(x_t, y) \left(\sum_{\theta \in \Theta} (\pi_t(\theta) - \delta_{\theta^c}(\theta)) \nabla_x f(y; x_t, \theta) \right) dy \right| \\ &\leq C * C' \int_{\mathcal{Y}} 1 dy \left| \sum_{\theta \in \Theta} \pi_t(\theta) - \delta_{\theta^c}(\theta) \right| \rightarrow 0, \end{aligned}$$

almost surely as $t \rightarrow \infty$, since $\mathcal{Y}$ is bounded and we have the consistency of $\pi_t(\theta)$ from Proposition 4. $\square$

REFERENCES

- Brückner, M., C. Kanzow, and T. Scheffer. 2012. “Static Prediction Games for Adversarial Learning Problems”. *The Journal of Machine Learning Research* 13(1):2617–2654.
- Chang, J. Y., and J. C. Hsu. 1992. “Optimal designs for multiple comparisons with the best”. *Journal of Statistical Planning and Inference* 30(1):45–62.
- Corlu, C. G., and B. Biller. 2013. “A subset selection procedure under input parameter uncertainty”. In *Proceedings of the 2013 Winter Simulation Conference*, edited by R. Pasupathy, S.-H. Kim, A. Tolk, R. Hill, and M. E. Kuhl, 463–473. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Corlu, C. G., and B. Biller. 2015. “Subset selection for simulations accounting for input uncertainty”. In *Proceedings of the 2015 Winter Simulation Conference*, edited by L. Yilmaz, W. K. V. Chan, I. Moon, T. M. K. Roeder, C. Macal, and M. D. Rossetti, 437–446. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- D. Dmitriy and X. Lin 2020. “Stochastic optimization with decision-dependent distributions”. arXiv preprint arXiv:2011.11173, accessed 6th March.

- Fan, W., L. J. Hong, and X. Zhang. 2020. "Distributionally Robust Selection of the Best". *Management Science* 66(1):190–208.
- Fu, M. C. 2008. "What you should know about simulation and derivatives". *Naval Research Logistics* 55(8):723–736.
- Gao, S., H. Xiao, E. Zhou, and W. Chen. 2017. "Robust ranking and selection with optimal computing budget allocation". *Automatica* 81:30–36.
- Juan, P., Z. Tijana, M.-D. Celestine, and H. Moritz. 2020. "Performative Prediction". In *Proceedings of the 37th International Conference on Machine Learning*, edited by H. D. III and A. Singh, Volume 119, 7599–7609. New York, NY: Association for Computing Machinery.
- Kushner, H., and G. Yin. 2003. *Stochastic Approximation and Recursive Algorithms and Applications*. New York, NY: Springer.
- Lappas, N. H., and C. E. Gounaris. 2018. "Robust optimization for decision-making under endogenous uncertainty". *Computers & Chemical Engineering* 111:252–266.
- Lee, S., T. Homem-de Mello, and A. J. Kleywegt. 2012. "Newsvendor-type models with decision-dependent uncertainty". *Mathematical Methods of Operations Research* 76(2):189–221.
- Luo, F., and S. Mehrotra. 2020. "Distributionally robust optimization with decision dependent ambiguity sets". *Optimization Letters* 14:2565–2594.
- Song, E., and B. L. Nelson. 2019. "Input–Output Uncertainty Comparisons for Discrete Optimization via Simulation". *Operations Research* 67(2):562–576.
- Song, E., and U. V. Shanbhag. 2019. "Stochastic Approximation for simulation Optimization under Input Uncertainty with Streaming Data". In *Proceedings of the 2019 Winter Simulation Conference*, edited by N. Mustafee, K.-H. Bae, S. Lazarova-Molnar, M. Rabe, C. Szabo, P. Haas, and Y.-J. Son, 3597–3608. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Wu, D., and E. Zhou. 2017. "Ranking and Selection under Input Uncertainty: a Budget Allocation Formulation". In *Proceedings of the 2017 Winter Simulation Conference*, edited by W. K. V. Chan, A. D'Ambrogio, G. Zacharewicz, N. Mustafee, G. Wainer, and E. Page, 2245–2256. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Wu, D., and E. Zhou. 2019. "Fixed Confidence Ranking and Selection under Input Uncertainty". In *Proceedings of the 2019 Winter Simulation Conference*, edited by N. Mustafee, K.-H. Bae, S. Lazarova-Molnar, M. Rabe, C. Szabo, P. Haas, and Y.-J. Son, 3717–3727. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Wu, D., H. Zhu, and E. Zhou. 2018. "A Bayesian Risk Approach to Data-driven Stochastic Optimization: Formulations and Asymptotics". *SIAM Journal on Optimization* 28(2):1588–1612.
- Xiao, H., F. Gao, and L. H. Lee. 2020. "Optimal computing budget allocation for complete ranking with input uncertainty". *IISE Transactions* 52(5):489–499.
- Xiao, H., and S. Gao. 2018. "Simulation Budget Allocation for Selecting the Top-m Designs With Input Uncertainty". *IEEE Transactions on Automatic Control* 63(9):3127–3134.
- Xu, J., P. Glynn, and Z. Zheng. 2020. "Joint Resource Allocation for Input Data Collection and Simulation". In *Proceedings of the 2020 Winter Simulation Conference*, edited by K.-H. Bae, B. Feng, S. Kim, S. Lazarova-Molnar, Z. Zheng, T. Roeder, and R. Thiesing, 2126–2137. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Zhou, E., and T. Liu. 2018. "Online Quantification Of Input Uncertainty For Parametric Models". In *Proceedings of the 2018 Winter Simulation Conference*, edited by M. Rabe, A. A. Juan, N. Mustafee, A. Skoogh, S. Jain, and B. Johansson, 1587–1598. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Zhou, E., and D. Wu. 2017. "Simulation Optimization Under Input Model Uncertainty". In *Advances in Modeling and Simulation*, edited by A. Tolk, J. Fowler, G. Shao, and E. Yücesan. Springer Cham.
- Zhou, E., and W. Xie. 2015. "Simulation optimization when facing input uncertainty". In *Proceedings of the 2015 Winter Simulation Conference*, edited by L. Yilmaz, W. K. V. Chan, I. Moon, T. M. K. Roeder, C. Macal, and M. D. Rossetti, 3714–3724. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.

AUTHOR BIOGRAPHIES

TIANYI LIU is a Ph.D. candidate in the School of Industrial and Systems Engineering at Georgia Institute of Technology. His research interests include machine learning, stochastic optimization, and simulation. His email address is liu341@gatech.edu.

YIFAN LIN is a Ph.D. student in the School of Industrial and Systems Engineering at Georgia Institute of Technology. His research interests include reinforcement learning and simulation optimization. His email address is ylin429@gatech.edu.

ENLU ZHOU is an Associate Professor in the School of Industrial and Systems Engineering at Georgia Institute of Technology. Her research interests include control, simulation, and stochastic optimization. Her email address is enlu.zhou@isye.gatech.edu.