

SELECTION OF THE MOST PROBABLE BEST UNDER INPUT UNCERTAINTY

Kyoung-Kuk Kim

School of Management Eng.
Korea Advanced Institute of Science and Technology
85 Hoegi-ro
Dongdaemun-gu, Seoul 02455, REPUBLIC OF KOREA

Taeho Kim

Department of Mathematical Sciences
Korea Advanced Institute of Science and Technology
291 Daehak-ro
Yuseong-gu, Daejeon 34141, REPUBLIC OF KOREA

Eunhye Song

Department of Industrial and Manufacturing Eng.
Pennsylvania State University
301 Leonhard Building
University Park, PA 16802, USA

ABSTRACT

We consider a ranking and selection problem whose configuration depends on a common input model estimated from finite real-world observations. To find a solution robust to estimation error in the input model, we introduce a new concept of robust optimality: *the most probable best*. Taking the Bayesian view, the most probable best is defined as the solution whose posterior probability of being the best is the largest given the real-world data. Focusing on the case where the posterior on the input model has finite support, we study the large deviation rate of the probability of incorrectly selecting the most probable best and formulate an optimal computing budget allocation (OCBA) scheme for this problem. We further approximate the OCBA problem to obtain a simple and interpretable budget allocation rule and propose sequential learning algorithms. A numerical study demonstrates good performances of the proposed algorithms.

1 INTRODUCTION

When randomness in a simulation model is driven by input models estimated from finite real-world data, the simulation output is subject to uncertainty caused by estimation error in the input models. This additional uncertainty, distinguished from inherent stochastic simulation error, is often referred to as *input uncertainty*. If the simulation model is applied to find an optimal design or policy for the target real-world system, then input uncertainty must be accounted for the optimization procedure to make a correct statistical inference on the real world performance of a selected solution.

Several optimization via simulation frameworks have been proposed to account for input uncertainty, which we categorize into three groups according to their treatments of input uncertainty. The first is to apply a risk measure with respect to input model estimation error to the simulation output mean and optimize it; Corlu and Biller (2015), Wu and Zhou (2017), Pearce and Branke (2017), Ungredda et al. (2020) use a risk-neutral measure (mean) and Xie and Zhou (2015), Wu et al. (2018), Zhu et al. (2020) explore value at risk or conditional value at risk. The second category focuses on inference to provide a probability guarantee that the best solution chosen under the current input model is in fact optimal; Corlu and Biller (2013), Song et al. (2015), Song and Nelson (2019) take this view, however, point out that when estimation

error of input model is large, the desired probability guarantee may not be attained. The last category takes the distributionally robust optimization approach, which first assumes an ambiguity set on the input model, finds the worst-case input model for each solution within the ambiguity set, and selects the solution with the best worst-case performance; Gao et al. (2017), Fan et al. (2020) fall under this category.

In this work, we consider a ranking and selection (R&S) problem under input uncertainty focusing on the case when all k solutions in comparison share a common input model estimated from data. Specifically, we take the parametric Bayesian approach to input modeling, thus uncertainty about the input model is captured by the posterior distribution on the input model parameters. Because the simulation output mean is a functional of the input model, the configuration of the R&S problem (and its optimum, correspondingly) depends on the input model, which itself is uncertain. To find a solution with robust performance, we define a new concept of robust optimality, *the most probable best*, which is the solution with the largest posterior probability of being the best.

Our work is related to the contextual R&S studied by Gao et al. (2019), Li et al. (2020), and Shen et al. (2021). Personalized decision making is an example of contextual R&S problems in which the problem configuration depends on the covariates of an individual. Thus, their objective is to learn the best policies under all covariate values equitably, which differs from finding the most probable best.

To devise a R&S algorithm to select the most probable best, we study the large deviation property of the probability of incorrectly selecting the most probable best focusing on the case when the posterior distribution of the input parameter is approximated with an empirical distribution. We formulate the optimal computing budget allocation (OCBA) problem to minimize the asymptotic probability of incorrect selection based on the large deviation analysis. We further approximate the OCBA problem to obtain a simple and interpretable budget allocation rule. Since this rule is based on unknown quantities of the problem, we propose sequential learning algorithms that aim to achieve an optimal allocation in the limit.

The remainder of the paper is organized as follows. In Section 2, we present some background on the Bayesian statistical models. Section 3 formally defines our problem of interest. Section 4 presents the OCBA formulation for our problem, which is then approximated to provide easy-to-compute balance conditions in Section 5. Two sequential learning algorithms are introduced and discussed in Section 6 followed by numerical demonstrations in Section 7. Proofs of all theorems and propositions are omitted from the paper due to space limit.

2 PRELIMINARIES

Consider the following R&S problem:

$$i_0 := \arg \min_{1 \leq i \leq k} \mathbb{E}[Y_i(\theta^c)], \quad (1)$$

where $Y_i(\theta^c)$ is the simulation output of the i th solution whose inputs are generated from a parametric input model f_{θ^c} . We assume i_0 is unique.

Typically in real-world applications, θ^c is unknown and must be estimated from observations. We assume that size- m i.i.d. observations $\mathcal{Z}_m := \{\mathbf{Z}_1, \mathbf{Z}_2, \dots, \mathbf{Z}_m\}$ are collected from f_{θ^c} . Taking the Bayesian approach, uncertainty about θ^c can be modeled with θ that has prior distribution π_0 . Given $\mathcal{Z}_m$, its posterior is derived as $\pi_m(\theta) = \frac{\pi_0(\theta)L_m(\theta)}{\int \pi_0(\theta_1)L_m(\theta_1)d\theta_1}$, where $L_m(\theta)$ is the likelihood function of $\mathcal{Z}_m$.

Let $y_i(\theta) := \mathbb{E}[Y_i(\theta)|\theta]$ be the mean response at solution i conditional on realized θ . Note that $y_i(\theta)$ can only be estimated via simulation given θ . We introduce stochastic process $\eta_i(\theta)$ to model uncertainty about $y_i(\theta)$. Instead of taking a purely Bayesian stance, we view $y_i(\theta)$ to be fixed and adopt $\eta_i(\theta)$ to simply assist the learning process of $y_i(\theta)$. The following normal-normal model for $\eta_i(\theta)$ and $Y_i(\theta)$ is adapted as in Ryzhov (2016).

$$Y_i(\theta) \sim N(\eta_i(\theta), \lambda_i^2(\theta)), \quad \eta_i(\theta) \sim N(\mu_{i,0}(\theta), \sigma_{i,0}^2(\theta)), \quad (2)$$

where $\lambda_i(\theta)$ is the simulation error variance, and $\mu_{i,0}(\theta)$ and $\sigma_{i,0}^2(\theta)$ are mean and variance of the prior distribution of $\eta_i(\theta)$, respectively. Model (2) assumes that the simulation error at each (i, θ) pair is normally

distributed. For non-normal simulation error, this assumption can be justified by batching (Kim and Nelson 2006). Furthermore, we assume $\lambda_i(\theta)$ to be known; for an unknown $\lambda_i(\theta)$, a normal-gamma model can be applied; see for instance, Section 5 in Ryzhov (2016).

Suppose n replications have been made under some sampling policy. We denote the number of replications allocated to (i, θ) by $N_i^n(\theta)$. Conditional on the simulation outputs, $Y_{i1}(\theta), Y_{i2}(\theta), \dots, Y_{iN_i^n(\theta)}(\theta)$, at (i, θ) , the posterior mean and variance, $\mu_{i,n}(\theta)$ and $\sigma_{i,n}^2(\theta)$, of $\eta_i(\theta)$ are updated as

$$\sigma_{i,n}^2(\theta) = \left(\frac{1}{\sigma_{i,0}^2(\theta)} + \frac{N_i^n(\theta)}{\lambda_i^2(\theta)} \right)^{-1}, \quad \mu_{i,n}(\theta) = \sigma_{i,n}^2(\theta) \left(\frac{\mu_{i,0}(\theta)}{\sigma_{i,0}^2(\theta)} + \frac{1}{\lambda_i^2(\theta)} \sum_{r=1}^{N_i^n(\theta)} Y_{ir}(\theta) \right). \quad (3)$$

Assuming noninformative prior $\sigma_{i,0}^2(\theta) = \infty$ for all (i, θ) , (3) becomes $\sigma_{i,n}^2(\theta) = \lambda_i^2(\theta)/N_i^n(\theta)$ and $\mu_{i,n}(\theta) = \sum_{r=1}^{N_i^n(\theta)} Y_{ir}(\theta)/N_i^n(\theta)$. Further, let $\mathcal{E}_n$ denote the information set consisting of n simulation outputs.

3 PROBLEM FORMULATION

Song and Nelson (2019) point out that solving the plug-in version of (1) in which θ^c is replaced with θ introduces input model risk as the plug-in optimum is suboptimal to (1) in general. Instead, our goal is to solve the following problem:

$$i^*(\pi_m) := \arg \max_{1 \leq i \leq k} \mathbb{P}_{\pi_m} (y_i(\theta) - \min_{j \neq i} y_j(\theta) \leq \delta) \quad (4)$$

where $\delta \geq 0$ is an *error tolerance*. Formulation (4) aims to find a solution that has the highest probability of being δ -optimal considering all realizations of θ from the posterior density π_m . Clearly, $i^*(\pi_m)$ may fail to be equal to i_0 with finite $\mathcal{L}_m$. Nevertheless, with a limited amount of data, we argue that our formulation is robust to input uncertainty as it maximizes the posterior probability of selecting a correct δ -optimum.

In this paper, we focus on the case when $\delta = 0$; the case of $\delta > 0$ is currently under investigation. In the former, (4) returns a solution that has the highest probability of being a minimizer with respect to the posterior density π_m . Thus, we refer to $i^*(\pi_m)$ as *the most probable best* under π_m .

The definition of the most probable best exploits the common input data (CID) effect, which refers to the dependence among $y_1(\theta), y_2(\theta), \dots, y_k(\theta)$ caused by common θ estimated from data (Song and Nelson 2019). If the CID effect induces positive correlations among the conditional means, pairwise comparisons among them become sharper, which increases the posterior probability of $i^*(\pi_m)$ being optimal.

As a first step, we approximate the posterior distribution, $\pi_m(\theta)$, by an empirical distribution constructed from size B sample $\{\theta_1, \dots, \theta_B\} \sim \pi_m(\theta)$. Then, (4) can be rewritten as

$$i^*(\pi_m) = \arg \max_i \frac{1}{B} \sum_{b=1}^B \mathbf{1} \left\{ y_i(\theta_b) = \min_j y_j(\theta_b) \right\} = \arg \max_i \frac{1}{B} \sum_{b=1}^B \mathbf{1} \left\{ y_i(\theta_b) = y_{i^b(\pi_m)}(\theta_b) \right\}, \quad (5)$$

where $i^b(\pi_m)$ is a *conditional optimum* at θ_b , i.e., $i^b(\pi_m) := \arg \min_i y_i(\theta_b)$. Since we fix the input data size and its corresponding posterior, we omit π_m for notational convenience in the remainder of the paper, i.e., $i^* = i^*(\pi_m)$ and $i^b = i^b(\pi_m)$. Further, i^* and i^b for each θ_b are assumed to be unique to simplify the analysis.

Problem (5) has a nested structure, which distinguishes it from the classical R&S or the contextual R&S problems. At the inner level, the goal is to correctly identify i^b at each b . At the outer level, i^* is determined by aggregating the inner level decisions. We present a high-level algorithmic scheme for selecting the most probable best in the following page.

The goal of this paper is to develop a novel sampling strategy to be adopted in Step 4 of the algorithm above so that i^* can be learned efficiently. We measure efficiency by computing the *probability of correct selection* (PCS), $P(CS) = P(i_n^* = i^*)$, or equivalently, by computing the *probability of false selection* (PFS), $P(FS) := P(i_n^* \neq i^*) = 1 - P(CS)$. Many R&S algorithms aim to provide a static or dynamic sampling

Selection of Most Probable Best

- 1: **for** $n = 0, 1, 2, \dots$ **do**
 - 2: Inner problem : For each b , find $i_n^b := \arg \min_i \mu_{i,n}(\theta_b)$.
 - 3: Outer problem : Find $i_n^* = \arg \max_i \frac{1}{B} \sum_{b=1}^B \mathbf{1} \left\{ \mu_{i,n}(\theta_b) = \mu_{i_n^b,n}(\theta_b) \right\}$.
 - 4: Make the next sampling decision (i, θ_b) depending on some criteria given current information.
 - 5: Run simulation at (i, θ_b) and update the belief according to the update rule (3).
 - 6: **end for**
-

framework that maximizes the PCS (or minimize the PFS). However, this objective does not have a closed-form expression in general. To overcome this issue, we discuss a large deviation perspective for our problem in Section 4.

4 OPTIMAL COMPUTING BUDGET ALLOCATION

The false selection event becomes a rare event as simulation budget n increases. The OCBA scheme, first proposed by Chen et al. (2000), aims to find the optimal sampling rule that minimizes the PFS. Using the large deviation theory, Glynn and Juneja (2004) further show that

$$\lim_{n \rightarrow \infty} -\frac{1}{n} \log P(FS) = G(\boldsymbol{\alpha}), \quad \text{if } \alpha_i > 0, 1 \leq i \leq k$$

for some function $G(\cdot)$ where $\boldsymbol{\alpha} = (\alpha_1, \alpha_2, \dots, \alpha_k)$ is a vector of fixed sampling proportions at the k solutions. The quantity $G(\boldsymbol{\alpha})$ is often referred to as the large deviation rate (LDR) of PFS. Thus, maximizing $G(\boldsymbol{\alpha})$ is equivalent to minimizing PFS. Hence, OCBA can be formulated as the following program.

$$\max_{\boldsymbol{\alpha}=(\alpha_1, \dots, \alpha_k)} G(\boldsymbol{\alpha}) \text{ subject to } \sum_{i=1}^k \alpha_i = 1, \quad \alpha_i \geq 0, 1 \leq i \leq k. \quad (6)$$

In general, the rate function G depends on unknown quantities, which makes it difficult to solve (6) exactly. Instead, many algorithms that learn G sequentially from sample statistics have been proposed in the hope to achieve near-optimal sampling ratios. For instance, we refer the readers to Pasupathy et al. (2014), Shin et al. (2018), and Chen and Ryzhov (2019) that take this approach in solving their own problems.

In the following, we establish large deviation theory for Problem (5) and an optimal sampling strategy based on the theory. We define the following quantities for further development:

$$G_i(\theta_b) := \frac{(y_i(\theta_b) - y_{i^b}(\theta_b))^2}{2(\lambda_i^2(\theta_b)/\alpha_i(\theta_b) + \lambda_{i^b}^2(\theta_b)/\alpha_{i^b}(\theta_b))}, \quad i \neq i^b,$$

$$d_i := \sum_{b=1}^B \mathbf{1} \left\{ i^* = i^b \right\} - \sum_{b=1}^B \mathbf{1} \left\{ i = i^b \right\}.$$

The quantity $G_i(\theta_b)$ is the LDR of false selection under Model (2) (Glynn and Juneja 2004) for the conditional problem given θ_b ; if $\alpha_{i,n}(\theta_b) := N_i^n(\theta_b)/n \rightarrow \alpha_i(\theta)$, then we have $\lim_{n \rightarrow \infty} -\frac{1}{n} \log P(\mu_{i,n}(\theta_b) < \mu_{i^b,n}(\theta_b)) = G_i(\theta_b)$. The d_i measures dominance of i^* over the i th solution.

One can easily observe that $\max_{i \neq i^*} P(i_n^* = i) \leq P(i_n^* \neq i^*) \leq \sum_{i \neq i^*} P(i_n^* = i)$. This directly implies that

$$\liminf_{n \rightarrow \infty} -\frac{1}{n} \log P(i_n^* \neq i^*) = \min_{i \neq i^*} LDR_{i,i^*}, \text{ where } LDR_{i,i^*} := \liminf_{n \rightarrow \infty} -\frac{1}{n} \log P(i_n^* = i). \quad (7)$$

In words, (7) implies that characterizing the target LDR reduces to finding the LDR of marginal false selection for each i . The marginal false selection, $\{i_n^* = i\}$, occurs when some $i \neq i^b$ is selected as the

conditional optimum at some θ_b . We define a collection of sets of solution-parameter pairs for each $i \neq i^*$:

$$\mathcal{A}_i := \left\{ I \subset I_{\text{tot}} \mid i \text{ performs equally or better than } i^* \text{ if } \forall (j, \theta_b) \in I, j \text{ is misspecified to be } i^b, \right. \\ \left. \text{while at any } \theta_{b'} \text{ not included in } I, i^{b'} \text{ is specified correctly} \right\}, \quad (8)$$

where $I_{\text{tot}} := \{(i, \theta_b) \mid 1 \leq i \leq k, 1 \leq b \leq B\}$ is the total index set. Clearly, the set $\mathcal{A}_i$ is a subset of the power set $2^{I_{\text{tot}}}$. Since we assume the conditional optimum at each θ_b to be unique, each $I \in \mathcal{A}_i$ may include at most one solution- θ_b pair for each b . In fact, such $\mathcal{A}_i$ is not unique; in the subsequent development, we define $\mathcal{A}_i$ to be the largest set that satisfies (8). Note that $|\mathcal{A}_i|$ is upper bounded by $|2^{I_{\text{tot}}}| = 2^{kB}$ (loosely).

To illustrate $\mathcal{A}_i$, consider the example in Figure 1, where $k = 4$ and $B = 8$ and i^b is marked with a solid circle for $1 \leq b \leq 8$. Clearly, $i^* = 1$. As an example, let us focus on $\mathcal{A}_4$. Suppose $i_n^3 = 4$ and $i_n^b = i^b$ for all $b \neq 3$, i.e., the solution 4 is incorrectly selected as the optimum at θ_3 while the other conditional optima are correctly specified after n simulation runs. Then, solutions 1 and 4 have the same performance measure values leading to an incorrect selection. Therefore, $\{(4, \theta_3)\} \in \mathcal{A}_4$. Similarly, one can confirm that $\{(3, \theta_1), (3, \theta_2)\}$ and $\{(4, \theta_4), (4, \theta_5)\}$ are elements of $\mathcal{A}_4$. Of course, these are only a few examples; $\mathcal{A}_4$ contains more elements.

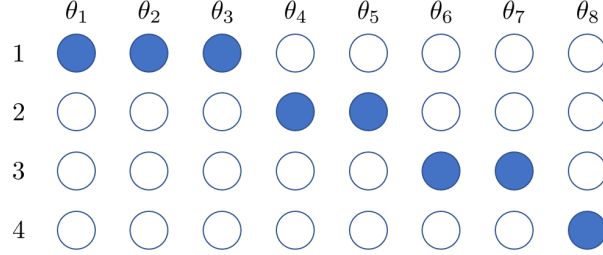


Figure 1: An example with $k = 4$ and $B = 8$. Conditional optima are marked with solid circles.

The following theorem stipulates the LDR of PFS for any fixed sampling scheme α .

Theorem 1 Given α , $\liminf_{n \rightarrow \infty} -\frac{1}{n} \log P \{i_n^* \neq i^*\} = \min_{i \neq i^*} LDR_{i,i^*}$, where

$$LDR_{i,i^*} = \min_{I \in \mathcal{A}_i} \sum_{(j, \theta_b) \in I} G_j(\theta_b). \quad (9)$$

Theorem 1 shows that the LDR of marginal false selection is determined by the slowest rate at which $i \neq i^*$ is falsely selected to be optimum among all possible misspecification cases in $\mathcal{A}_i$. Therefore, constructing $\mathcal{A}_i$ and finding the corresponding slowest rate for each i are central parts to determining the LDR of false selection for Problem (5). We close this section by presenting the OCBA framework for Problem (5):

$$LDR^* := \max_{\alpha} \min_{i \neq i^*} LDR_{i,i^*} \quad \text{subject to} \quad \sum_{i=1}^k \sum_{b=1}^B \alpha_i(\theta_b) = 1, \alpha_i(\theta_b) \geq 0, 1 \leq i \leq k. \quad (10)$$

5 APPROXIMATE SOLUTION TO OCBA FORMULATION

Computational complexity of characterizing $\mathcal{A}_i$ increases combinatorially in k and B . To derive LDR_{i,i^*} , however, it is not necessary to find $\mathcal{A}_i$ exactly. To see this, consider J_1 and $J_2 \in \mathcal{A}_i$ such that $J_1 \subset J_2$. Then, LDR_{i,i^*} remains unchanged when $\mathcal{A}_i$ is replaced with $\mathcal{A}_i - \{J_2\}$ in (9). In other words, we may remove all elements of $\mathcal{A}_i$ that are supersets of other elements without affecting the value of LDR_{i,i^*} . This approach is top-down, i.e., it can be carried out once $\mathcal{A}_i$ is fully characterized. Nevertheless, this observation motivates us to devise a computationally efficient bottom-up method to approximate LDR_{i,i^*} closely without fully characterizing $\mathcal{A}_i$.

Table 1: Five disjoint subsets of Ξ ; each subset belongs to $V_{i,\ell}$ with the corresponding value of ℓ .

	$i^b = i$	$i^b \neq i \text{ \& } i^b \neq i^*$	$i^b = i^*$
$j = i$	$\ell = 0$	$\ell = 1$	$\ell = 2$
$j \neq i$		$\ell = 0$	$\ell = 1$

To this end, suppose we define $\Xi = \{(j, \theta_b) : j \neq i^* \text{ and } j \neq i^b \text{ at } \theta_b\}$, and we have a realized sample mean $\mu_{i,n}(\theta_b)$. In words, Ξ includes all solution-parameter pairs, where the solution is neither i^* nor the conditional best at the paired parameter. For $i \neq i^*$, let $V_{i,\ell}$ be a subset of Ξ such that $(j, \theta_b) \in V_{i,\ell}$ if and only if $i_n^b = j$ and $i_n^{b'} = i^{b'}$ results in $\ell = (d_i - \tilde{d}_{i,n})^+$ for nonnegative integer ℓ , where $\tilde{d}_{i,n} = \sum_{b=1}^B \mathbf{1}\{i^* = i_n^b\} - \sum_{b=1}^B \mathbf{1}\{i = i_n^b\}$ is the estimated dominance of i^* over i . Then, simple arithmetic shows that $\{V_{i,0}, V_{i,1}, V_{i,2}\}$ constitutes a partition of Ξ . This is evident in Table 1, which classifies all (j, θ_b) in Ξ into five disjoint subsets and each subset is marked with the corresponding $V_{i,\ell}$ it belongs to. For the example in Figure 1, $(j, \theta_b) = (4, \theta_1)$ belongs to $V_{4,2}$ as $j = i = 4$ and $i^1 = 1$, which corresponds to the upper right corner of Table 1. Proceeding similarly, one can show that $V_{4,1} = \{(2, \theta_1), (2, \theta_2), (2, \theta_3), (3, \theta_1), (3, \theta_2), (3, \theta_3), (4, \theta_4), (5, \theta_4), (6, \theta_4), (7, \theta_4)\}$, $V_{4,2} = \{(4, \theta_1), (4, \theta_2), (4, \theta_3)\}$, and $V_{4,0} = \Xi \setminus (V_{4,1} \cup V_{4,2})$.

From the definition of $\mathcal{A}_i$, any $I \in \mathcal{A}_i$ is a subset of Ξ , and thus can be written as $I = I_0 \cup I_1 \cup I_2$, where $I_\ell := I \cap V_{i,\ell}$ for $\ell = 0, 1, 2$. As elements in $V_{i,0}$ do not cause $\tilde{d}_{i,n}$ to be smaller than d_i when misspecified, any $I \in \mathcal{A}_i$ such that $I_0 \neq \emptyset$ can be replaced with $I_1 \cup I_2$ without affecting the value of LDR_{i,i^*} . Note that $|I_1| + 2|I_2|$ is equal to $(d_i - \tilde{d}_{i,n})^+$ when all solution-parameter pairs in $I_1 \cup I_2$ are misspecified, while at any parameter not included in the pairs in $I_1 \cup I_2$, the conditional optimum is correctly specified. Therefore, $|I_1| + 2|I_2| \geq d_i$ is a necessary condition for I to be in $\mathcal{A}_i$. Now, for any I , the lower bound on

$$\sum_{(j, \theta_b) \in I_1} G_j(\theta_b) + \sum_{(j, \theta_b) \in I_2} G_j(\theta_b)$$

can be obtained by replacing $G_j(\theta_b), \forall (j, \theta_b) \in I_1$ with $\min_{(j, \theta_b) \in V_{i,1}} G_j(\theta_b)$ and $G_j(\theta_b), \forall (j, \theta_b) \in I_2$ with $\min_{(j, \theta_b) \in V_{i,2}} G_j(\theta_b)$, respectively. Expanding this argument and exploiting the necessary condition mentioned above, we can further obtain an easy-to-compute lower bound on LDR_{i,i^*} as follows.

Proposition 1 Let $\underline{LDR}_{i,i^*} := d_i \min \left\{ \min_{(j, \theta_b) \in V_{i,1}} G_j(\theta_b), \frac{1}{2} \min_{(j, \theta_b) \in V_{i,2}} G_j(\theta_b) \right\}$. Then, $LDR_{i,i^*} \geq \underline{LDR}_{i,i^*}$.

Intuitively, a correct specification of an element in $I_{i,2}$ is more important since misspecifying it reduces d_i twice as much as misspecifying an element in $I_{i,1}$. Our lower bound reflects this intuition by giving a half weight to $V_{i,2}$. Accordingly, an approximate version of (10) is formulated as below.

$$\underline{LDR}^* := \max_{\alpha} \min_{i \neq i^*} \underline{LDR}_{i,i^*} \quad \text{subject to} \quad \sum_{i=1}^k \sum_{b=1}^B \alpha_i(\theta_b) = 1, \alpha_i(\theta_b) \geq 0. \quad (11)$$

Proposition 1 implies that $\underline{LDR}^* \leq LDR^*$. Since (11) is a convex program, the Karush-Kuhn-Tucker conditions provide simple and interpretable optimality conditions for (11) as stated in Theorem 2.

Theorem 2 (Balance conditions for (11)) Any allocation rule $\alpha = \{\alpha_i(\theta_b)\}_{1 \leq i \leq k, 1 \leq b \leq B}$ is optimal for (11), if and only if, α satisfies the following system of equations;

- $\alpha_{i^*}(\theta_b) = 0$ for all b such that $i^b \neq i^*$.
- (Pairwise balance condition) For all $(i, \theta_b), (j, \theta_{b'}) \in \Xi$,

$$W_i(\theta_b) G_i(\theta_b) = W_j(\theta_{b'}) G_j(\theta_{b'}), \quad (12)$$

where $W_i(\theta_b)$ is defined as

$$W_i(\theta_b) = \begin{cases} \min \left(\min_{j \neq i^*} d_j, \frac{d_i}{2} \right), & \text{if } i^b = i^* \\ d_i, & \text{if } i^b \neq i^* \end{cases}. \quad (13)$$

- (Global balance condition) For all $1 \leq b \leq B$, $\frac{\alpha_{i^b}^2(\theta_b)}{\lambda_{i^b}^2(\theta_b)} = \sum_{i \neq i^b} \frac{\alpha_i^2(\theta_b)}{\lambda_i^2(\theta_b)}$.

The first condition states that the asymptotically optimal sampling ratio for i^* at any parameter that does not have i^* as its conditional optimum is 0. This may be surprising as the optimal sampling ratios for the classical R&S problem obtained by solving (6) are strictly positive (Glynn and Juneja 2004). This stark difference has implications in our problem setting. Suppose i^* is correctly identified as the conditional optimum at all θ_b such that $i^b = i^*$. Then, for other θ_b s, it only matters whether the best among the solutions other than i^* is correctly identified not to advantage $i \neq i^*$ when determining the most probable best. Even if i^* is incorrectly identified as the conditional optimum at those θ_b s, correct selection still occurs.

The pairwise balance condition provides the weight $W_i(\theta_b)$ of $G_i(\theta_b)$ at each (i, θ_b) , which we refer to as *balance weight*. The balance weight for each $G_i(\theta_b)$ depends on the dominance factor, d_i . The smaller d_i is, the more important to simulate solution i is, since it is more likely to outperform i^* due to simulation error. Also, notice that $W_i(\theta_b)$ is smaller for θ_b whose conditional optimum is i^* as it is important to correctly find the conditional optimum at such θ_b to make correct selection. The following $k \times B$ matrix shows the balance weights computed according to (13) for the example in Figure 1

$$\mathbf{W} := \begin{bmatrix} \infty & \infty & \infty & \infty & \infty & \infty & \infty & \infty \\ 0.5 & 0.5 & 0.5 & \infty & \infty & 1 & 1 & 1 \\ 0.5 & 0.5 & 0.5 & 1 & 1 & \infty & \infty & 1 \\ 1 & 1 & 1 & 2 & 2 & 2 & 2 & \infty \end{bmatrix}.$$

We fill the elements corresponding to (i^b, θ_b) and (i^*, θ_b) for all b with ∞ since these pairs are not included in Ξ , and thus do not appear in (12).

Let $\mathbf{G}$ be a $k \times B$ matrix whose (i, b) th entry corresponds to $G_i(\theta_b)$. For two matrices A and B with the same size, the Hadamard product, $A \circ B$, is defined as $(A \circ B)_{ij} = A_{ij}B_{ij}$. Condition (12) implies that the elements of $\mathbf{W} \circ \mathbf{G}$ corresponding to all $(i, \theta_b) \in \Xi$ are identical.

Lastly, the global balance condition ensures that the conditional best at each θ_b is correctly identified.

6 SEQUENTIAL LEARNING PROCEDURE

Since $y_i(\theta)$ is unknown in advance, we need a dynamic sampling policy that simultaneously learns the mean surface and optimal allocation. We present a novel sequential sampling algorithm based on the balance weights derived in Theorem 2 in this section.

Given $\mathcal{E}_n$, we approximate i^* with i_n^* , and $\underline{LDR}_{i, i^*}$ with

$$\underline{LDR}_{i, i_n^*}(\mathcal{E}_n) = d_{i,n} \min \left\{ 2 \min_{(j, \theta_b) \in \hat{V}_{i,1}} G_{j,n}(\theta_b), \frac{1}{2} \min_{(j, \theta_b) \in \hat{V}_{i,2}} G_{j,n}(\theta_b) \right\},$$

respectively. Note that $\hat{V}_{i,\ell}$ is a plug-in version of $V_{i,\ell}$ constructed in the same way as $V_{i,\ell}$ by replacing i^* with i_n^* , and i^b with i_n^b for all b , respectively, from Table 1. The quantities, $G_{i,n}$ and $d_{i,n}$, are plug-in versions of G_i and d_i , respectively, defined by replacing $y_i(\theta_b)$, $\alpha_i(\theta)$, i^b , and i^* with $\mu_{i,n}(\theta_b)$, $\alpha_{i,n}(\theta)$, i_n^b , and i_n^* , respectively, which are all $\mathcal{E}_n$ -measurable. Similarly, we denote the plug-in versions of $\mathbf{W}$ and $\mathbf{G}$ by $\mathbf{W}_n$ and $\mathbf{G}_n$, respectively. All plug-in quantities are constructed based on simulation results up to the n th replication, and therefore $\mathcal{E}_n$ -measurable. Moreover, these plug-in quantities depend on the sampling policy deployed up to the n th replication, although we do not explicitly denote the dependence due to notational convenience. Recall that we denote the fraction of replications allocated to (i, θ_b) up to the n th replication by $\alpha_{i,n}(\theta_b)$. The following proposition stipulates sufficient conditions for $\{\alpha_{i,n}(\theta_b)\}$ to satisfy in order for the resulting $\underline{LDR}_{i, i_n^*}(\mathcal{E}_n)$ to converge to $\underline{LDR}_{i, i^*}$ for each i as n increases.

Proposition 2 For all $1 \leq i \leq k$, $1 \leq b \leq B$, suppose $N_i^n(\theta_b) = n\alpha_{i,n}(\theta_b) \rightarrow \infty$ and $\alpha_{i,n}(\theta_b) \rightarrow \alpha_i(\theta_b)$ almost surely. Then, $|\underline{LDR}_{i, i_n^*} - \underline{LDR}_{i, i^*}(\mathcal{E}_n)| \rightarrow 0$ as $n \rightarrow \infty$ with probability 1.

Thanks to Proposition 2, it suffices to find a dynamic sampling policy satisfying $N_i^n(\theta_b) \rightarrow \infty$ and $\alpha_{i,n}(\theta_b) \rightarrow \alpha_i(\theta_b)$ almost surely to achieve $\underline{LDR}^*$ in the limit. To accomplish this, we exploit the balance conditions in Theorem 2. Namely, we find the solution-parameter pair corresponding to the smallest entry of $\mathbf{W}_n \circ \mathbf{G}_n$, say (i, θ_b) , and select either (i, θ_b) or (i_n^b, θ_b) to simulate based on the global balance condition.

Before we present our main algorithm, we briefly explain how to handle a case when there is a tie in identifying i_n^* . Although we only consider the case when i^* is unique, i_n^* may not be for finite n . Let us visit the example in Figure 1 once again for exposition. Suppose after sampling n solution-parameter pairs, i^b is incorrectly specified to be solution 2 instead of solution 3 while all other i^b s are selected correctly. This scenario is depicted in Figure 2, where i_n^b for all θ_b are boxed. As a result, Solutions 1 and 2 are tied given $\mathcal{E}_n$. Suppose a tie-breaking rule (e.g. random selection) is applied and solution 1 is picked as i_n^* . Then, the resulting balance weight matrix is

$$\mathbf{W}_n := \begin{bmatrix} \infty & \infty & \infty & \infty & \infty & \infty & \infty & \infty \\ 0 & 0 & 0 & \infty & \infty & \infty & 0 & 0 \\ 0 & 0 & 0 & 2 & 2 & 2 & \infty & 2 \\ 0 & 0 & 0 & 2 & 2 & 2 & 2 & \infty \end{bmatrix}.$$

Notice that some balance weights are zero since $d_2 = 0$. As a consequence, if we choose the next solution-parameter pair to simulate by finding the smallest element of $\mathbf{W}_n \circ \mathbf{G}_n$, then several elements of the matrix have zero entries making them indistinguishable from each other. To resolve such a deadlock, we modify $\mathbf{W}_n$ by substituting zeroes with ones and non-zeros with ∞ should a tie occurs when determining i_n^* . We describe the intuition for this modification in the following.

One of the reasons for the tie to occur is when $i_n^* \neq i^*$. In the example discussed above, both i_n^* and i^* happen to be Solution 1, however, this is not always the case. Assuming $i_n^* \neq i^*$, the tie can be broken by correcting a misspecification at θ_b where i_n^* is selected as the conditional optimum or at a θ_b where a tied solution is deemed (conditional) suboptimal. This strategy corresponds to sampling a solution-parameter pair among those with 0 entries in $\mathbf{W}_n$. If the sampling algorithm allocates infinitely many replications to all (i, θ_b) , the strong law of large numbers implies that the event of tie happens only finitely many times with probability 1, so the modification to $\mathbf{W}_n$ does not affect the limiting behavior of our algorithm.

Based on these observations, we propose Algorithm 1. Algorithm 1 is easy to implement as we only need to compute $\mathbf{G}_n$, $\mathbf{W}_n$, and $(N_i^n(\theta_b)/\lambda_i(\theta_b))^2$ at each iteration. We cannot enjoy this convenience with (10) due to complexity of $\mathcal{A}_i$.

We conjecture that $\alpha_{i,n}(\theta_b)$ allocated by Algorithm 1 converges to $\alpha_i(\theta_b)$ satisfying Theorem 2 as n increases, which is the second part of the sufficient condition for Proposition 2 to hold. However, Algorithm 1 cannot guarantee the first part of the condition: $N_i^n(\theta_b) \rightarrow \infty$ for all (i, θ_b) . To see why, let us define $\Theta^* := \{\theta_b : i^* = i^b\}$. Algorithm 1 stops assigning replications to $\{(i_n^*, \theta_b) : i_n^b \neq i_n^*\}$ once it correctly specifies the minimum number of $\{b : i^* = i^b\}$ needed to distinguish i_n^* from the second best. For instance, suppose that $B = 50$, $|\Theta^*| = 15$, and the second best solution is the conditional optima at 10 θ_b s. Then, it

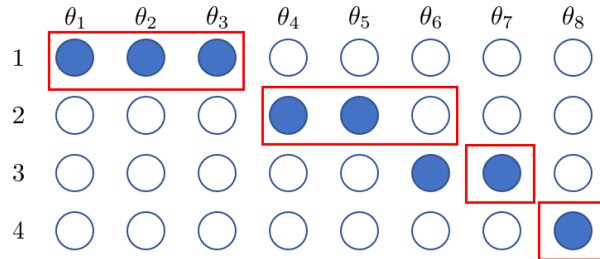


Figure 2: The same example as in Figure 1. The solutions in boxes are the selected conditional optima given a sample path.

Algorithm 1 Rate Optimal Sequential Sampling Algorithm

-
- 1: Warm-start by allocating n_0 simulation budgets for all (i, θ_b) . Let $n = n_0 k B$ and update $\mu_{i,n}(\theta_b)$ and $G_{i,n}(\theta_b)$.
 - 2: **while** simulation budget remains **do**
 - 3: Form a balance weight matrix $\mathbf{W}_n$ and a LDR matrix $\mathbf{G}_n$. Find $(i, \theta_b) = \arg \min \mathbf{W}_n \circ \mathbf{G}_n$.
 - 4: **if** $\left(N_{i^b}^n(\theta_b)/\lambda_{i^b}(\theta_b)\right)^2 < \sum_{j \neq i^b} \left(N_j^n(\theta)/\lambda_j(\theta_b)\right)^2$ **then**
 - 5: Run a replication at (i^b, θ_b) .
 - 6: **else**
 - 7: Run a replication at (i, θ_b) .
 - 8: **end if**
 - 9: Update $\mu_{i,n}(\theta_b)$, $G_{i,n}(\theta_b)$ and $d_{i,n}$ at (i, θ_b) . Update $\hat{V}_{i,1}$ and $\hat{V}_{i,2}$ at all i . Let $n \leftarrow n + 1$.
 - 10: **end while**
-

is enough to correctly specify i^* as the conditional optima at 13 out of 15 parameters in Θ^* to make correct selection provided that the conditional optima are correctly specified at all other parameters not in Θ^* . Then, even if at 2 remaining parameters in Θ^* the second best is incorrectly specified as the conditional optima, correct selection occurs.

As this behavior may impede performance of the algorithm, we modify Algorithm 1 to guarantee that $N_{i^*}^n(\theta_b)$ increases sublinearly in n for all $(i^*, \theta_b), i^b \notin \Theta^*$, which in turn results in $N_{i^*}^n(\theta_b) \rightarrow \infty$ and $\alpha_{i^*,n}(\theta_b) \rightarrow 0$, the exact sufficient condition of Proposition 2 Algorithm 1 fails to achieve.

The modification involves a notable sampling criterion, *expected improvement* (EI), first introduced by Jones et al. (1998). The *information valuation function* is defined as $f(x) = x\Phi(x) + \phi(x)$, where ϕ and Φ are probability density and cumulative distribution functions of standard normal random variable, respectively. Given θ_b , the EI of (i, θ_b) with respect to (i^b, θ_b) , is computed as

$$v_{i,n}(\theta_b) = \mathbb{E} \left[\left(\mu_{i^b,n}(\theta_b) - \eta_i(\theta_b) \right)^+ \middle| \mathcal{E}_n \right] = \sigma_{i,n}(\theta_b) f \left(-\frac{|\mu_{i,n}(\theta_b) - \mu_{i^b,n}(\theta_b)|}{\sigma_{i,n}(\theta_b)} \right).$$

The EI has been widely applied in Bayesian optimization and best-arm identification problems. For classical R&S, Ryzhov (2016) shows that sequentially sampling the largest-EI solution at each iteration allocates simulation budget to a suboptimal solution at the rate of $O(\log n)$ as the total budget n increases. Although this behavior leads to a suboptimal LDR for the traditional R&S problem, it is precisely what we need to guarantee for all $(i^*, \theta_b), i^b \notin \Theta^*$ in our problem. Algorithm 2 below presents our modification.

Algorithm 2 Modified Rate Optimal Sequential Sampling Algorithm

-
- 1: Choose (i, θ_b) according to Algorithm 1 without simulating it.
 - 2: **if** $i_n^b \neq i_n^*$ and $\alpha_{i^b,n}(\theta_b)^{1/2} v_{i^b,n}(\theta_b) < v_{i^*,n}(\theta_b)$ **then**
 - 3: Run a replication at (i^*, θ_b) .
 - 4: **else**
 - 5: Run a replication at (i, θ_b) .
 - 6: **end if**
-

Once (i, θ_b) is selected in Step 1, Step 2 compares the (scaled) EI of (i^b, θ_b) with the EI of (i^*, θ_b) to select the next pair to simulate. If it turns out the EI of (i^b, θ_b) is relatively small compared to (i^*, θ_b) , then we simulate (i^*, θ_b) . Otherwise, we simulate (i, θ_b) selected in Step 1. Notice that we slightly modify the EI criterion for (i^b, θ_b) by multiplying $\alpha_{i^b,n}(\theta_b)^{1/2}$ to $v_{i^b,n}(\theta_b)$. Without the multiplier, we have $N_{i^*}^n(\theta_b) = O(\log N_{i^b}^n(\theta_b))$ following the result in Ryzhov (2016) as $N_{i^b}^n(\theta_b)$ is the amount allocated to the conditional best at θ_b . If k and B are large, then $N_{i^b}^n(\theta_b)$ becomes small for finite n . As a result, Step 3 of

Algorithm 2 rarely occurs. With the multiplier, we have

$$\alpha_{i_n^b, n}(\theta_b)^{1/2} v_{i_n^b, n}(\theta_b) = \left(N_{i_n^b}^n(\theta_b)/n \right)^{1/2} v_{i_n^b, n}(\theta_b) = n^{-1/2} \lambda_{i_n^b}(\theta_b) f(0). \quad (14)$$

Note that the right-hand side of (14) would be the EI of (i_n^b, θ_b) if n replications were allocated to (i_n^b, θ_b) . Therefore, our modification makes Algorithm 2 allocate the effort to (i_n^*, θ_b) as if n replications are allocated to (i_n^b, θ_b) . As a result, we expect Step 3 to occur at the rate of $O(\log n)$ as desired.

Asymptotic properties of Algorithms 1 and 2 remain to be investigated further in future research.

7 EMPIRICAL ANALYSIS

In this section, we present a numerical performance analysis using a synthetic example with $k = 10$ and $B = 50$. We compare Algorithms 1 and 2 with two existing dynamic sampling policies described below.

- Equal allocation (EA) : This policy allocates the simulation budget uniformly over $\{(i, \theta_b)\}$.
- Contextual R&S allocation (C-OCBA) : Proposed by Gao et al. (2019), it aims to maximize the worst-case probability of correct selection (worst-case PCS over all contexts), which is defined as $\min_b P(i_n^b = i^b)$ for our problem if θ_b s represent contexts. Finding i^b for each b is equally important under this framework.

We assume $\lambda_i^2(\theta_b) = 5^2$ for all (i, θ_b) . For each θ_b , we fix the conditional optimum to be

$$i^b = \begin{cases} j, & \text{if } 5j - 4 \leq b \leq 5j, \text{ for some } 1 \leq j \leq 7, \\ 8, & \text{if } 36 \leq b \leq 41, \\ 10, & \text{if } 42 \leq b \leq 50. \end{cases}.$$

This implies $i^* = 10$ as its posterior probability of being optimal is $9/50 = 0.18$. All solutions except for solution 9 is a conditional optimum at some θ_b . For each $1 \leq b \leq 50$, we set $y_{i^b}(\theta_b) = 1$ and fill in $\{y_i(\theta_b)\}_{1 \leq i \leq k, i \neq i^b}$ randomly without replacement from $\{2, 3, \dots, k\}$ for each macro run. In Algorithms 1 and 2, we apply the smallest index rule as the tie-breaking rule; this makes i^* never selected as i_n^* when tied.

Figure 3 presents empirical performances of four algorithms we compare. All results are averaged from 10,000 macro runs and the initial sample size, n_0 , is set to 5 for all algorithms. The left-hand side panel of Figure 3 shows that the logarithmic PFS decreases as the simulation budget increases for each algorithm. Notice that our algorithms outperform other methods; EI slightly improves the convergence rate.

The right-hand side panel of Figure 3 displays the number of parameters at which $i_n^* = i^*$ is selected as the conditional optima; we have $|\Theta^*| = 9$ for this problem. Notice that C-OCBA outperforms Algorithms 1 and 2 in this measure. Recall that the goal of our algorithm is to achieve the faster convergence of PCS. For this purpose, misspecifying i^b for one b is allowed since the second best ($i = 8$) is the conditional optimum at 6 parameters, i.e., $\min_{j \neq i^*} d_j = 3$. For this reason, Algorithm 1 tends not to characterize all conditional optima correctly once it correctly specifies 8 conditional optima.

We ran longer simulations to observe the long-run behavior of Algorithm 1 and 2 in estimating Θ^* as presented in Figure 4. As we conjectured, for Algorithm 1, the size of estimated Θ^* converges to 8. In contrast, for Algorithm 2, the same statistic inches toward 9 albeit more slowly than C-OCBA. This result confirms that Algorithm 2 serves the exact purpose it is designed for; that is, to ensure increasing number of replications are allocated to (i_n^*, θ_b) at all θ_b in Θ^* .

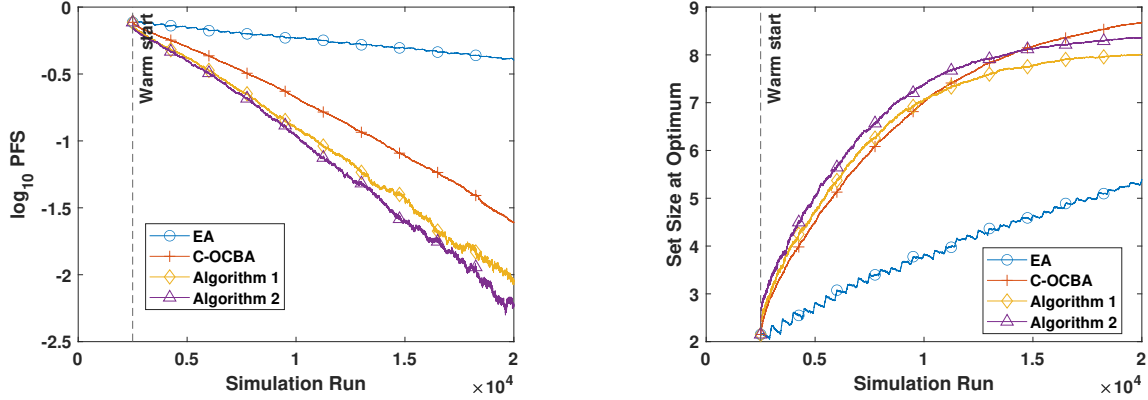


Figure 3: The logarithmic PFS (left) and the estimated size of Θ^* (right) averaged over 10000 macro runs.

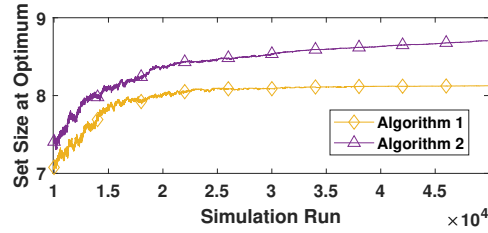


Figure 4: Long-run behaviors of Algorithms 1 and 2 in estimating Θ^* .

ACKNOWLEDGEMENT

Song's work is partially supported by the United States National Science Foundation under grant DMS-1854659. The work by K. Kim is supported by the Basic Science Research Program through the National Research Foundation of Korea funded by the Ministry of Science and ICT (NRF-2019R1A2C1003144).

REFERENCES

- Chen, C.-H., J. Lin, E. Yücesan, and S. E. Chick. 2000. "Simulation budget allocation for further enhancing the efficiency of ordinal optimization". *Discrete Event Dynamic Systems* 10(3):251–270.
- Chen, Y., and I. O. Ryzhov. 2019. "Balancing optimal large deviations in ranking and selection". 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, 3368–3379. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Corlu, C., 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, T. M. K. R. I. Moon, C. Macal, and M. D. Rossetti, 437–446. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Fan, W., L. J. Hong, and X. Zhang. 2020. "Distributionally robust selection of the best". *Management Science* 66(1):190–208.
- Gao, S., J. Du, and C.-H. Chen. 2019. "Selecting the optimal system design under covariates". In *2019 IEEE 15th International Conference on Automation Science and Engineering (CASE)*, 547–552. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Gao, S., H. Xiao, E. Zhou, and W. Chen. 2017. "Robust ranking and selection with optimal computing budget allocation". *Automatica* 81:30–36.
- Glynn, P., and S. Juneja. 2004. "A large deviations perspective on ordinal optimization". In *Proceedings of the 2004 Winter Simulation Conference*, edited by R. G. Ingalls, M. D. Rossetti, J. S. Smith, and B. A. Peters. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.

- Jones, D. R., M. Schonlau, and W. J. Welch. 1998. "Efficient global optimization of expensive black-box functions". *Journal of Global optimization* 13(4):455–492.
- Kim, S.-H., and B. L. Nelson. 2006. "Selecting the Best System". In *Simulation*, edited by S. G. Henderson and B. L. Nelson, Volume 13 of *Handbooks in Operations Research and Management Science*, 501–534. Elsevier.
- Li, H., H. Lam, Z. Liang, and Y. Peng. 2020. "Context-Dependent Ranking and Selection under a Bayesian Framework". In *Proceedings of the 2020 Winter Simulation Conference*, edited by K.-H. G. Bae, B. Feng, S. Kim, S. Lazarova-Molnar, Z. Zheng, T. Roeder, and R. Thiesing, 2060–2070. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Pasupathy, R., S. R. Hunter, N. A. Pujowidianto, L. H. Lee, and C.-H. Chen. 2014. "Stochastically constrained ranking and selection via SCORE". *ACM Transactions on Modeling and Computer Simulation* 25(1):1–26.
- Pearce, M., and J. Branke. 2017. "Efficient expected improvement estimation for continuous multiple ranking and selection". 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, 2161–2172. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Ryzhov, I. O. 2016. "On the convergence rates of expected improvement methods". *Operations Research* 64(6):1515–1528.
- Shen, H., L. J. Hong, and X. Zhang. 2021. "Ranking and selection with covariates for personalized decision making". *INFORMS Journal on Computing*. In press.
- Shin, D., M. Broadie, and A. Zeevi. 2018. "Tractable sampling strategies for ordinal optimization". *Operations Research* 66(6):1693–1712.
- Song, E., and B. L. Nelson. 2019. "Input–output uncertainty comparisons for discrete optimization via simulation". *Operations Research* 67(2):562–576.
- Song, E., B. L. Nelson, and L. J. Hong. 2015. "Input Uncertainty and Indifference-zone Ranking & Selection". In *Proceedings of the 2015 Winter Simulation Conference*, edited by L. Yilmaz, W. K. V. Chan, T. M. K. R. I. Moon, C. Macal, and M. D. Rossetti, 414–424. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Ungredda, J., M. Pearce, and J. Branke. 2020. "Bayesian Optimisation vs. Input Uncertainty Reduction". <https://arxiv.org/abs/2006.00643>.
- 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., 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.
- Xie, W., and E. Zhou. 2015. "Simulation Optimization when Facing Input Uncertainty". In *Proceedings of the 2015 Winter Simulation Conference*, edited by L. Yilmaz, W. K. V. Chan, T. M. K. R. I. Moon, C. Macal, and M. D. Rossetti, 3714–3724. Piscataway, New Jersey: Institute of Electrical and Electronics Engineers, Inc.
- Zhu, H., T. Liu, and E. Zhou. 2020. "Risk Quantification in Stochastic Simulation under Input Uncertainty". *ACM Transactions on Modeling and Computer Simulation* 30(1):1–24.

AUTHOR BIOGRAPHIES

Kyoung-Kuk Kim is an associate professor in the School of Management Engineering at Korea Advanced Institute of Science and Technology. His research interests include simulation analytics as well as financial engineering and risk management. His email address is kkim3128@kaist.ac.kr.

Taeho Kim is the corresponding author and currently a Ph.D. candidate in the Department of Mathematical Sciences at Korea Advanced Institute of Science and Technology. His research focuses on the data-driven simulation modeling of complex system and simulation optimization. His email address is thk5594@kaist.ac.kr and his webpage is at <https://sites.google.com/view/taehokim>.

Eunhye Song is the Harold and Inge Marcus Early Career assistant professor in Industrial and Manufacturing Engineering at Penn State. Her research interests include simulation model risk analysis, robust simulation optimization, and large-scale discrete simulation optimization. Her email address is eus358@psu.edu and her website can be found at eunhyesong.info.