

Selecting a Good Stochastic Inventory Policy for the Huge Number of Alternative Policies

Mohammad H. Almomani¹ and Mahmoud H. Alrefaei²

¹Department of Mathematics, Faculty of Science,
The Hashemite University, P.O box 330127, Zarqa 13133, JORDAN
e-mail: mh momani@hu.edu.jo

²Department of Mathematics and Statistics,
Jordan University of Science and Technology, JORDAN
e-mail: alrefaei@just.edu.jo

October 27, 2025

Abstract

In this paper, we present the problem of selecting a good stochastic inventory policy with high probability and minimum total simulation cost when the number of alternative policies is huge. We start with the Ordinal Optimization (*OO*) procedure to select a subset that overlaps with the set of the actual best $m\%$ policies with high probability. Then we use a version of the Optimal Computing Budget Allocation (*OCBA*) to select a small subset that contains the best policy. This follows by a Subset Selection (*SS*) procedure to get a smaller subset that contains the best policy among the subset that is selected before. Finally, the Indifference–Zone (*IZ*) procedure is used to select the best policy among the survivors policies in the previous stage. Numerical test involves with all these procedures show the results for selecting a good stochastic inventory policy

with high probability and a minimum number of simulation samples, when the number of policies is large.

Math. Subject Classification: 62F07, 90C06, 49K30.

Key Words and Phrases: Stochastic Inventory Model, Simulation Optimization; Ranking and Selection; Ordinal Optimization; Optimal Computing Budget Allocation; Subset Selection; Indifference–Zone

1 Introduction

All companies are seeking to achieve a greater percentage profit and gain the largest number of customers through the provision of goods needed by the customers for that resort to the inventory theory. Inventory theory contend with the management of stock levels of goods in order to ensuring that demand for these goods is met. Most designs are modeled to search about answers for two fundamental questions. The first one is when a new order should be placed, and the second is how much the order quantity should be in order to minimize the total cost. The complexity of these questions lies in the assumptions concerning about demand, the cost structure and physical characteristics of the policy. There is no standard solution for all inventory models, since the conditions for each company are unique and contain distinct features and limitations, so the number of possible models are very large. This depends on different assumptions on key parameters such as demand, costs and the physical nature of the policy.

We consider the discrete optimization problem of selecting the best stochastic inventory policy from the top $m\%$ policies, when the number of alternative policies is large. The best stochastic inventory policy is the one that has the minimum total cost. Discrete event systems simulation is a popular tool for analyzing designs and evaluating decision problems since real problems rarely satisfy the assumption of analytical models. However, due to its computational demands, efficiency is an important concern.

In this problem the stochastic simulation is used to infer the total cost of each stochastic inventory policy. However, since simulation output has randomness, the best stochastic

inventory policy cannot be selected with certainty. Therefore, the Probability of Correct Selection ($P(CS)$) is used as a measure of the quality of a selection procedures. The selection procedures identify the best stochastic inventory policy with high probability with the goal of these procedures is to maximize $P(CS)$.

Ranking and Selection (*R&S*) procedures are used usually to select the best stochastic inventory policy or a subset that contains the best policies when the number of alternative policies is relatively small, with a pre specified significance level. The problem arises for a large scale problems because it needs a huge computational time. For each sample of the performance value require one simulation run. Therefor for a large scale problem will need a large number of samples which is very time consuming and may be impossible. In this situation, we would change our objective to finding good stochastic inventory policies rather than estimating accurately the performance value for these policies.

This idea lies in Ordinal Optimization (*OO*) procedure. The objective of the *OO* procedure is to isolate a subset of good policies with high probability and reduce the required simulation time. Since the *OO* procedure can reduce computation cost, the use of Optimal Computing Budget Allocation (*OCBA*) can further reduce the already low cost by allocating the computing budget among different inventory policies instead of equally simulating all policies. In general, to select the best stochastic inventory policy among a finite but large alternative policies, one should take in account the amount of computational time needed for simulating each policy. If a limited computational budget is available to be distributed among the different policies, then instead of distributing these computations evenly, the *OCBA* can be used to distribute this budget in a smart way so as to maximize $P(CS)$, Chen [1].

In this paper, we consider the problem of selecting the best inventory policy when the number of policies is finite but huge. According to Kim and Nelson [2], statistical selection procedures are used to select the best policy from a finite set of policies, when the stochastic simulation is used to specify the performance measure for each alternative in the simulation. Since, the use of simulation has a potential for incorrect selection, we need a measure to determine the quality of selection. The probability of correct selection, where the correct

selection occurs when the selected policy by the selection procedure is the same as the actual best policy.

Selection procedures has been used previously by many authors to select the best design, when the number of solutions is large. For example, Almomani and Abdul Rahman [3] proposed a sequential selection procedure to select a good design when the number of alternatives is large. This procedure starts with the ordinal optimization procedure to select a subset that overlaps with the set of the actual best $m\%$ systems with high probability. Then they use optimal computing budget allocation to allocate the available computing budget in a way that maximizes the probability of correct selection. This is followed by a subset selection procedure to get a smaller subset that contains the best system among the subset that is selected before. Finally, the indifference-zone procedure is used to select the best system among the survivors in the previous stage. Almomani and Alrefaei [4] proposed a sequential procedure for selecting the best subset of simulated systems when the number of alternatives is very large. The procedure consists of two phases; in phase one, the ordinal optimization approach is used to select a relative small subset with a probability of overlapping with the subset that contains the actual best $k\%$ systems is high. In the second phase, the available computing budget using the *OCBA* — m method is allocated to identify all top m systems from the subset that is selected by ordinal optimization in the first phase. More reviews on the sequential procedures can be found in, Al-Salem et al. [5] and Almomani et al. [6].

In this paper, our aim is to select the best stochastic inventory policy with high probability, when the number of policies is large. As a measure of performance, we use $P(CS)$. The best inventory policy will be a policy with a minimum total cost. In many cases, it is hard to accomplish this goal, because it requires long computational time. To estimate the total cost, we use simulation which is costly and needs long computational time especially when we are working with a large number of policies. Initially, we use the *OO* procedure to select a good subset from the set of all policies, with a high probability that this subset contains one of the best $m\%$ alternative policies. Following by *OCBA* technique that distribute the available computing budget on the policies in such a way to maximize the $P(CS)$. Then, a

Subset Selection (SS) procedure is used to screen out the solution set and to eliminate the non-competitive policies and to select a relative small subset from the previous subset, that contains the best policy with high probability. Finally, procedures of Indifference–Zone (IZ) are used to select the best stochastic inventory policy from the previous subset.

The rest of this paper is organized as follows; Section 2, contains the problem statement. In Section 3, we give a background about the stochastic inventory model. In Section 4, we review the *OO* procedure, *OCBA* technique, *SS* procedure and *IZ* procedure. A selection of good stochastic inventory policy for large number of alternatives is presented in Section 5. Finally, concluding remarks are given in Section 6.

2 Problem Statement

The stochastic inventory problem can be considered as the following optimization problem

$$\min_{\vartheta \in \Theta} f(\vartheta) \quad (2.1)$$

where Θ is the feasible solution set and it is an arbitrary, may have no structure, finite but a huge set, f is the expected performance measure of some complex stochastic design $f(\vartheta) = E[L(\vartheta, Y)]$, where ϑ is a vector representing the system design parameters, Y represents all the random effect of the design and L is a deterministic function that depends on ϑ and Y .

The goal of the selection procedure is to identify one of the best n simulated designs, where n is a large number. The best design is defined as the design with the smallest mean, which is unknown and to be inferred from simulation. Suppose that there are n designs, and let Y_{ij} (observation) represent the j^{th} output from the design i , and let $\mathbf{Y}_i = \{Y_{ij}, j = 1, 2, \dots\}$ denote the output sequence from the design i . We assume that Y_{ij} are independent and identically distributed (*i.i.d.*) normal with unknown means $\mu_i = E(Y_{ij})$ and variances $\sigma_i^2 = Var(Y_{ij})$. Also we assume that $\mathbf{Y}_1, \mathbf{Y}_2, \dots, \mathbf{Y}_n$ are mutually independent. In practice σ_i^2 is unknown for all i , so we estimate it using the sample variances S_i^2 for Y_{ij} . Since we assume that a smallest mean is better, therefore if the ordered μ_i -values are denoted by

$\mu_{[1]} \leq \mu_{[2]} \leq \dots \leq \mu_{[n]}$, then the design having mean $\mu_{[1]}$ is referred to as the best design. The Correct Selection (CS) occurs when the design selected by the selection procedure is the same as the actual best design. When the number of alternatives is huge as we assume in this work, it is impossible to simulate all alternatives, therefore we reduce our objective to select one of the best alternative provided that the P (CS) is very large with reasonable computational time.

3 Stochastic Inventory Model

The inventory management problem typically encompasses two questions to be answered; when to order an inventory and how much to order, in such a way to minimize the total cost. The costs consist of the ordering cost, the holding cost and the shortage cost. In the stochastic inventory model (s, S) , it is assumed that the demand is stochastic. When the inventory level I , is dropped below or at a reordering point s , an order of $S - I$ is made, and the cost of doing so is known as the ordering cost. When a demand occurs and I dropped to zero, a shortage occurs, and a cost called the shortage cost is therefore imposed. To reduce overall costs, there must be a balance between the holding cost, the price of keeping a stocked inventory and other costs. The resultant issue is a stochastic optimization issue and simulation can be used to estimate the total cost. The goals of the inventory management issue are to make it more productive to accomplish things quickly but more expensively or slowly but cheaply. Such a plan would include the best inventory management practices, which would lower the total cost of the manufacturing, holding, and inventory shortage.

An excess occurs when there is still some inventory following a request because the demand is lower than the quantity of supply. But if demand exceeds supply, there will be a stock shortfall, leading to incomplete demand; this is referred to as a shortage. An uncertainty about demand will be taken into account in stochastic inventory models. Continual or periodic review is used to divide stochastic inventory models into two categories to restock the inventory the next time frame. It is decided how much to order in the next period after

determining its current inventory level after conclusion of previous session.

Our focus will be on a single item that is subject to periodic review, complete backlog, and random lead times, with K as the ordering cost for set-up and u for cost per item. Backlog is considered to have a set shortage cost of π per item per unit of time and a holding cost of one item per unit of time is h . Per demand, the products has multi-modal distribution, while the time in demand is independent, and the rate λ denotes the identically distributed exponential random variables. Initial inventory level I is where the inventory level begins. According to the demand on the final period, inventories will gradually decrease over time.

The inventory level is recalculated at the conclusion of each period to decide whether or not a fresh order is required at the start of the subsequent period. When a certain demand arises, we compare the demand values to the stock available during the time. If the demand value is greater than the stock quantity, a defect or inventory shortage occurs. Otherwise, there is a surplus rather than a shortage. We place an order at the conclusion of this period if I is below the necessary level; otherwise, there is no need to do so. The inventory model (s, S) is depicted in Figure 1, see Debadyuti [7]. In mathematically notation the order quantity is given by,

$$Z = \begin{cases} S - I & \text{if } I < s, \\ 0 & \text{if } I \geq s \end{cases}$$

where I is the inventory level at the beginning of the period.

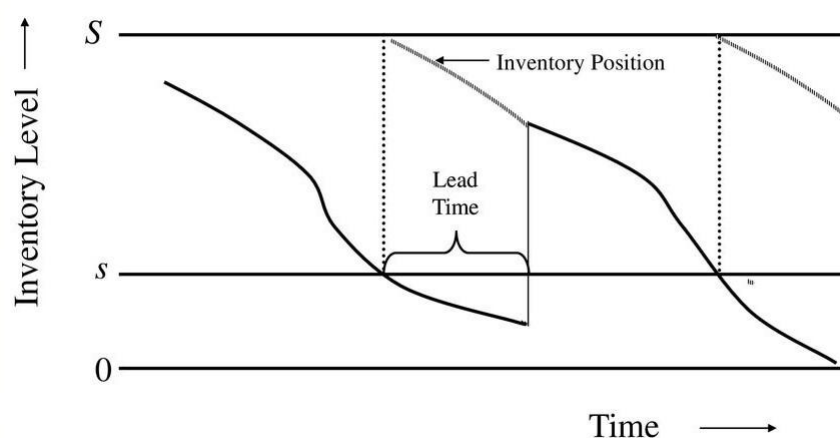


Figure 1: The (s, S) inventory model

The total cost is the sum of the three costs, ordering cost, holding cost, and shortage cost. If Z items are order and then the order cost is $K + iZ$, where K is a fixed cost (setup cost) and i is the incremental cost per item. The shortage cost equal to πI^- , where π is the shortage cost per item per period, and I^- is the area under the curve of the demand when the shortage occurs divided by the length of period. The holding cost also equal to hI^+ , where h is the holding cost per item per period, and I^+ is the area under the curve of demand in the normal mode (when there are items in the inventory level) divided by the length of period. Thus the total cost can be written as

$$TC = (K + iZ) + \pi I^- + hI^+ \quad (3.1)$$

Now, the problem become of minimizing the total cost function in (3.1), see Law and Kelton [8].

The stochastic (s, S) inventory problem can be described as

$$\min_{\vartheta \in \Theta} TC(\vartheta) \quad (3.2)$$

where Θ is the set of all (s, S) policies; $TC(\vartheta)$ is the total cost of the policy ϑ , where ϑ is a vector that represents policy parameters. For more details of inventory models, see Winston and Goldberg [9], Tai et al. [10], Taleizadeh et al. [11], Hajiagha et al. [12], Zhang et al. [13].

4 Background

In this section, we preview the components of the proposed algorithm for solving the stochastic inventory problem.

4.1 Ordinal Optimization

The Ordinal Optimization (OO) procedure has been proposed by Ho et al. [14] to deal with large scale selection problems. The idea of OO is to select a design or a subset from the

feasible solution set that contains some of the best designs with high probability, rather than focusing on accurately estimating the values of the design performance during the optimization process. This means that order should first be considered before values are estimated, or in other words ordinal optimization precedes cardinal optimization. The objective is to isolate a subset of good enough designs with high probability and reduce the required simulation time for the discrete event simulation.

Suppose that the Correct Selection (CS) is to select a subset G of g designs from the feasible solution set Θ , that contains at least one of the top $m\%$ best designs. Since we assume that $|\Theta| = n$ is large, then the probability of correct selection is given by $P(CS) \approx 1 - (1 - \frac{k}{100})^g$. Furthermore, suppose that the correct selection is to select a subset G of g designs that contains at least r of the best s designs. Let S be the subset that contains the actual best s designs, then the probability of correct selection can be obtained using the hyper-geometric distribution as $P(CS) = P(|G \cap S| \geq r) = \sum_{i=r}^g \frac{\binom{s}{i} \binom{n-s}{g-i}}{\binom{n}{g}}$. Since n is large, then $P(CS)$ can be approximated using the binomial distribution as $P(CS) \approx \sum_{i=r}^g \frac{g}{i} \left(\frac{s}{n}\right)^i \left(1 - \frac{s}{n}\right)^{g-i}$, where we assume that $s/n \times 100\% = m\%$. After selecting the set G by using the *OO* procedure, the problem now is to optimize the objective function over the smaller subset G . A comprehensive review of the *OO* procedure can be found in Horng and Lin [15], Shin et al. [16] and Ho et al. [17].

4.2 Optimal Computing Budget Allocation

The Optimal Computing Budget Allocation (*OCBA*) is used to determine the best simulation lengths for each simulation design in order to reduce the total computational time. In fact, this procedure was proposed to improve the performance of *OO* procedure by determining the optimal numbers of simulation samples for each design, instead of simulating equally all the designs.

The goal of *OCBA* is to allocate the total simulation samples from all designs in a way that maximizes the probability of selecting the best design within a given computing budget, see Chen et al.[1] and Chen et al.[18]. Let T be the total sample that are available for solving

the optimization problem given in (2.1). The target is to allocate these computed simulating samples to maximize $P(CS)$. In mathematical notation this can be written as:

$$\begin{aligned} \max_{T_1, \dots, T_n} & P(CS) \\ \text{s.t.} & \sum_{i=1}^n T_i = T \\ & T_i \in \mathbb{N} \quad i = 1, 2, \dots, n \end{aligned}$$

where N is the set of non-negative integers, T_i is the number of samples allocated to design i , and $\sum_{i=1}^n T_i$ denotes the total computational samples. Assume that the simulation times for different designs are roughly the same. To solve this problem Chen et al. [1] proved the following theorem.

Theorem 4.1. Given a total number of simulated samples T to be allocated to n competing designs whose performance is depicted by random variables with means $J(\vartheta_1), J(\vartheta_2), \dots, J(\vartheta_n)$, and finite variances $\sigma_1^2, \sigma_2^2, \dots, \sigma_n^2$ respectively, as $T \rightarrow \infty$, the approximate probability of correct selection can be asymptotically maximized when

$$\begin{aligned} 1. & T_i = \frac{\sigma_i^2 / \delta_{ij}^2}{\sum_{j=1, j \neq i}^n \sigma_j^2 / \delta_{ij}^2}; \text{ where } i, j \in \{1, 2, \dots, n\} \text{ and } i \neq j \neq b. \\ 2. & T = \sigma_b^2 \frac{\sum_{i=1, i \neq b}^n T_i}{\delta_{bb}^2} \end{aligned}$$

where δ_{ij} the estimated difference between the performance of the two designs ($\delta_{ij} = \bar{J}_i - \bar{J}_j$), and $\bar{J}_i \leq \min_{j \neq i} \bar{J}_j$ for all i . Here $\bar{J}_i = \frac{1}{T_i} \sum_{j=1}^{T_i} L(\vartheta_i, \xi_{ij})$, where ξ_{ij} is a sample from ξ_i for $j = 1, \dots, T_i$.

Remark 4.1. If we select a design s following the above theorem where $s \neq b$, we will get

$$N_s = \frac{\sigma_s^2 / \delta_{sb}^2}{\sum_{i=1, i \neq b}^n \sigma_i^2 / \delta_{ib}^2} N_b \text{ where } i = 1, 2, \dots, n. \delta_{sb} = \bar{f}_b - \bar{f}_s \text{ and } s \neq i \neq b.$$

Suppose

$$C_i = \frac{\sigma_i^2 / \delta_{ib}^2}{\sigma_b^2 / \delta_{bb}^2} \text{ then } N_i = C_i N_b \quad (4.1)$$

where $s \neq i \neq b$.

$$\text{Also, } N_b = \frac{\sigma_b^2}{\sigma_b^2} \sum_{i=1, i \neq b}^n \frac{\sigma_i^2}{\delta_{ib}^2} N_{s_i} \text{ and let}$$

$$C_b = \frac{\sigma_b^2}{\sigma_b^2} \sum_{i=1, i \neq b}^n \frac{\sigma_i^2}{\delta_{ib}^2} \frac{1}{C_i} \text{ then } N_b = C_b N_b \quad (4.2)$$

Since $\sum_{i=1}^n N_i = B$ then $N_b \left(\sum_{i=1}^n C_i \right) = B$, where we assume that when $C_b = 1$, this implies that $N_b = \frac{B}{\sum_{i=1}^n C_i}$ and therefore,

$$N_i = \frac{C_i B}{\sum_{j=1}^n C_j} \quad (4.3)$$

Alrefaei et al. [19] discussed how to distribute the available computational budget to the alternative designs in order to maximize the probability of correct selection. Also, they discussed how to approximate the probability of correct selection (*APCS*). More reviews on the *OCBA* procedure can be found in Banks [20], Chen and Lee [21], Huang and Choi [22], and Myoung and Byon [23].

4.3 Subset Selection

Subset Selection (*SS*) procedure screens out the search space and eliminate non-competitive designs and construct a subset I that contains the best design with high probability. The number of designs in subset I could be determined by the user or randomly. This procedure can be used to select a random subset size that contains the actual best design when the number of alternatives is relatively large. For this purpose, it requires that $P(CS) \geq P^*$, where *CS* is the correct selection (i.e. selecting a subset that contains the actual best design) and P^* is a predetermined probability.

The *SS* procedure dating back to Gupta [24], who presented a single stage procedure for producing a subset containing the best design with a specified probability. Extension of this work that relevant to the simulation setting includes Roostapour et al. [25] and Thompson [26].

4.4 Indifference–Zone

The main aim of Indifference–Zone (*IZ*) procedures is selecting the best design among n designs when $n \leq 20$. The *IZ* is defined as the interval $[\mu_{[1]}, \mu_{[1]} + \delta^*]$, where δ^* is a predetermined small positive real number. The interest is in selecting an alternative i^* such that $\mu_{i^*} \in [\mu_{[1]}, \mu_{[1]} + \delta^*]$. Let *CS* be selecting an design whose mean belongs to the indifference zone. We would like that *CS* to take place with high probability, say with a probability not smaller than P^* where $1/n \leq P^* \leq 1$. In mathematical notation, we seek to achieve $P(CS) \geq P^*$ given that $|\mu_{[2]} - \mu_{[1]}| \geq \delta^*$. We should be careful when to specified the value of

δ^* and P^* , because if δ^* is too small, then the number of replication samples that are required to guarantee the achieve $P(CS) \geq P^*$ is expected to be large. Therefore, δ^* can be thought of the smallest difference $\mu_{[2]} - \mu_{[1]}$ and is considered worth detecting. Same situation for P^* because with larger P^* , it may requires a large number of replication samples to meet the designer request.

The *IZ* procedure dating back to Rinott [27] who presented a procedure that is applicable when the data are normally distributed and all designs are simulated independently of each others. This procedure consists of two stages for the case when variances are completely unknown. Another comprehensive review of *IZ* procedure can be found in Bechhofer et al. [28], Kim and Nelson [29], Yu et al. [30] and Yoon and Bekker [31].

5 The Proposed Procedure for Selecting a Good Stochastic Inventory Policy

To select a good stochastic inventory policy from large number of inventory policies, we combine four procedures include *OO*, *OCBA*, *SS* and *IZ*. Initially, using *OO* procedure, a subset G is randomly selected form a feasible solution set that overlaps with the set that contains the actual best $m\%$ policies with high probability. Then *OCBA* procedure is used to allocate the available computing budget to select a smaller subset. This is followed with *SS* procedure to get a smaller subset I with high probability, that contains the best policy among the previous selected subset, where $|I| \leq 20$. Finally, *IZ* procedure is applied to select the best policy from that set I . The algorithm is described as follows:

Algorithm 5.1. :-

Step 0: Determine g and k where $|G| = g$ and $|G'| = k$. Here, G is defined as the selected subset from Θ , that satisfies $P(G \text{ contains at least one of the best } m\% \text{ policies}) \geq 1 - \alpha_1$, and G' is defined as the selected subset from G , where $g \geq k$. Let the number of initial simulation samples $t_0 \geq 2$, and $t = t_{(1-\alpha_2/2)^{\frac{1}{g-1}}, t_0-1}$ from the t -distribution. Consider

also the indifference zone δ^* . Let $T_1^l = T_2^l = \dots = T_g^l = t_0$, here l is the iteration number, and determine the total computing budget T .

Define the confidence level for selecting a good policy using this procedure as: the probability of selecting a subset that contains a good policy using SS procedure $\geq 1 - \alpha_2$, the probability of selecting a good policy using IZ procedure $\geq 1 - \alpha_3$ and the approximate probability of correct selection by OCBA procedure $\geq 1 - \alpha_4$.

Step 1: Select a subset G of size g randomly from Θ . Take a random samples of t_0 observations $TC(\vartheta_{ij})$ ($1 \leq j \leq t_0$) for each policy i in G , where $i = 1, \dots, g$.

Step 2: Calculate the sample mean (total cost mean) $\overline{TC}(\vartheta_i)^{(1)}$ and variances s_i^2 as $\overline{TC}(\vartheta_i)^{(1)} = \frac{\sum_{j=1}^{T_i^l} TC(\vartheta_{ij})}{T_i^l}$ and $s_i^2 = \frac{\sum_{j=1}^{T_i^l} (TC(\vartheta_{ij}) - \overline{TC}(\vartheta_i)^{(1)})^2}{T_i^l - 1}$ where $i = 1, \dots, g$.

Step 3: Order the policies in G according to their sample averages: $\overline{TC}(\vartheta_{[1]})^{(1)} \leq \overline{TC}(\vartheta_{[2]})^{(1)} \leq \dots \leq \overline{TC}(\vartheta_{[g]})^{(1)}$.

Step 4: Select the best k policies from the set G , and represent this subset as G' .

Step 5: If $\sum_{i=1}^g T_i^l \geq T$ select randomly a subset G'' of the $g - k$ alternative policies from $\Theta - G'$.

Step 6: Compute the approximate probability of correct selection (APCS) of OCBA, if $APCS < 1 - \alpha_4$, let $(G = G' \cup G'')$.

Step 6(a): Increase the computing budget by Δ and compute the new budget allocation, $T_1^{l+1}, T_2^{l+1}, \dots, T_g^{l+1}$, using equations (4.1), (4.2) and (4.3).

Step 6(b): Perform additional $\max\{0, T_i^{l+1} - T_i^l\}$ simulations for each policy $i, i = 1, \dots, g$, let $l \leftarrow l + 1$. Go to Step 2.

Step 7: Let $W_{ij} = t \frac{s_i}{T_i} + \frac{s_j}{T_j}^{1/2}$ for all $i \neq j$.

Step 8: Set $I = \{i : 1 \leq i \leq k \text{ and } \overline{TC(\vartheta)}_i^{(1)} \leq \overline{TC(\vartheta)}_j^{(1)} - [W_{ij} - \delta^*]^-, \forall i \neq j\}$, where $[x]^- = x$ if $x < 0$ and $[x]^- = 0$ otherwise.

Step 9: If I contains a single policy, then this policy is the good policy. Otherwise, for all $i \in I$, compute the second sample size $N_i = \max\{T_i, \lceil (\frac{hS_i}{\delta^*})^2 \rceil\}$, where $h = h(1 - \alpha_3/2, t_0, |I|)$ be the Rinott [27] constant.

Step 10: Take additional $N_i - T_i$ random samples of $TC(\vartheta_{ij})$ for each policy $i \in I$, and compute the overall sample means for $i \in I$ as $\overline{TC(\vartheta_i)}^{(2)} = \frac{\sum_{j=1}^{N_i} TC(\vartheta_{ij})}{N_i}$.

Step 11: Select policy $i \in I$ with the smallest $\overline{TC(\vartheta_i)}^{(2)}$ as a good policy.

We need to choose the initial number of simulation replication t_0 and the one time increment Δ . The purpose of taking t_0 simulation replication for each n policy is to get some information about the performance of each policy.

Alrefaei et al. [19] presented a procedure that resembles the OCBA, but it gives an approximation of the probability of correct selection (APCS). Nelson et al. [32] have shown that with probability at least $1 - (\alpha_2 + \alpha_3)$ this procedure selects a good policy in the subset G . Therefore, if G contains at least one of the top $m\%$ policies, then this procedure selects a good policy with probability $1 - (\alpha_2 + \alpha_3)$. Therefore, we have the following lemma:

Lemma 5.1. *If the feasible set Θ is large, the approach given in Algorithm (5.1) will select one of the best $m\%$ policies from the feasible solution set Θ with probability:*

$$\begin{aligned} P(CS^*) &\geq 1 - \left(1 - \frac{m}{100}\right)^g + \alpha_2 + \alpha_3 * (1 - \alpha_4) \\ &\geq 1 - \left(1 - \frac{m}{100}\right)^g + \alpha_2 + \alpha_3 + \alpha_4. \end{aligned}$$

Proof. Let $\alpha_1 = 1 - \frac{m}{100}^g$ then

$$\begin{aligned} P(CS^*) &= (1 - \alpha_1)(1 - \alpha_2)(1 - \alpha_3)(1 - \alpha_4) \\ &= (1 - \alpha_1 - \alpha_2 + \alpha_1\alpha_2)(1 - \alpha_3 - \alpha_4 + \alpha_3\alpha_4) \\ &= 1 - \alpha_1 - \alpha_2 + \alpha_1\alpha_2 - \alpha_3 + \alpha_1\alpha_3 + \alpha_2\alpha_3 - \alpha_1\alpha_2\alpha_3 - \alpha_4 + \alpha_1\alpha_4 \\ &\quad + \alpha_2\alpha_4 - \alpha_1\alpha_2\alpha_4 + \alpha_3\alpha_4 - \alpha_1\alpha_3\alpha_4 - \alpha_2\alpha_3\alpha_4 + \alpha_1\alpha_2\alpha_3\alpha_4 \\ &= (1 - \alpha_1 - \alpha_2 - \alpha_3 - \alpha_4) + \alpha_1\alpha_2 + \alpha_1\alpha_3 + \alpha_2\alpha_3 + \alpha_1\alpha_4 + \alpha_2\alpha_4 + \alpha_3\alpha_4 \\ &\quad - \alpha_1\alpha_2\alpha_3 - \alpha_1\alpha_2\alpha_4 - \alpha_1\alpha_3\alpha_4 - \alpha_2\alpha_3\alpha_4 + \alpha_1\alpha_2\alpha_3\alpha_4 \end{aligned}$$

Now, let $\epsilon = \alpha_1\alpha_2(1 - \alpha_3) + \alpha_1\alpha_3(1 - \alpha_4) + \alpha_2\alpha_3(1 - \alpha_4) + \alpha_1\alpha_4(1 - \alpha_2) + \alpha_2\alpha_4 + \alpha_3\alpha_4 + \alpha_1\alpha_2\alpha_3\alpha_4$

$$\begin{aligned} \Rightarrow P(CS^*) &= (1 - \alpha_1 - \alpha_2 - \alpha_3 - \alpha_4) + \epsilon \\ &= 1 - (\alpha_1 + \alpha_2 + \alpha_3 + \alpha_4) + \epsilon \end{aligned}$$

it is clear that $\epsilon > 0 \Rightarrow P(CS^*) > 1 - (\alpha_1 + \alpha_2 + \alpha_3 + \alpha_4)$

$$\Rightarrow P(CS^*) > 1 - 1 - \frac{m}{100}^g + \alpha_2 + \alpha_3 + \alpha_4$$

□

5.1 Implementation of the proposed procedure for stochastic inventory model

Now we discuss how to implement the proposed procedure for solving the stochastic inventory model. The model is used in our experiment is stochastic inventory model with periodic review (s, S) inventory model. It is a single item under periodic review, full backlogging, and random lead times (uniformly distributed between 0.5 and 1.0 period), with ordering costs (including a fixed set-up cost of \$32 per order and an incremental cost of \$3 per item), on-hand inventory (\$1 per item per period), and backlogging (fixed shortage cost of \$5 per item per period). The time between demands is assumed to be independent and identically distributed exponential random variable and the number of demand items in each demand is distributed as follows: 1, 2, 3, and 4, with probability $1/6$, $1/3$, $1/3$, and $1/6$, respectively. This experiment involves solving inventory model using Algorithm (5.1) with different parameter settings to estimate the probability of correct selection $P(CS)$.

The objective is to locate the policy with the minimum total cost, using Algorithm (5.1), for selecting the best inventory policy. We utilize $P(CS)$ as the measure of the performance of the procedure. Also, we pursue to maximizes $P(CS)$. We implement Algorithm (5.1) for solving this problem under some parameter settings.

Example 1:

In the first experiment, assume the number of stochastic inventory policies is $n = 150$. We run simulation for each 150 policies to select the best policy that has the minimum total cost. Then, we apply Algorithm (5.1) on these policies to select a good one. Suppose we want to select one of the best $m\% = 4\%$ policies. Then we are selecting the best 6 policies, which are (20, 190), (20, 40), (30, 190), (20, 60), (20, 50), and (20, 150). We consider the policy that is selected by Algorithm (5.1) as a Correct Selection (CS) if it is one of these 6 policies. Set the parameters here as; $g = 20$, $t_0 = 10$, $k = 10$, $\Delta = 10$, $\alpha_2 = \alpha_3 = 0.005$, $\delta^* = 0.05$, $\alpha_4 = 0.06$, and $T = 4000$. Then, we can evaluate the approximated probability of the CS as

$$P(CS^*) \geq 1 - \frac{1}{100} + 0.005 + 0.005 + 0.06 \geq 0.49.$$

Table 1 contains the results of this experiment, with 10 replications for selecting one of the best 4% policies. From the table, “*Good Policy*” means that the policy that has been selected by Algorithm (5.1) as the good policy, which it has the minimum total cost in that replication, “*Total Cost*” means that the total cost of the policy that has been selected by Algorithm (5.1) in that replication, $\sum_{i=1}^g T_i$ is the total sample size used in Step 5 in algorithm, $\sum_{i \in I} N_i$ is the total sample size used in Step 9 in algorithm, and $|I|$ is the number of policies from Step 8 in algorithm. Note that in Algorithm (5.1) we required $|I| < 20$, and this is hold in this experiment. Clearly, in each replication with star (*), the Algorithm (5.1) selects one from the best 4% policies, also from Table 1 that the proposed algorithm made the correct selection in 5 out of 10 replication (i.e. $P(CS) = 0.5$).

Replications	Good Policy	Total Cost	$\sum_{i=1}^g T_i$	$\sum_{i \in I} N_i$	$ I $
1	(50, 120)	272.681561136988	Σ	Σ	8
2*	(30, 190)	181.255587654684	3897 3564	3658 2998	9
3	(40, 100)	242.865879588456	3054	3568	5
4	(70, 170)	296.125689784556	4981	3257	7
5*	(20, 50)	185.796381809419	4754	3871	8
6*	(20, 40)	178.675838192989	3551	3412	8
7*	(20, 190)	176.199889754628	3115	2358	9
8	(20, 200)	211.954812122354	2828	1541	9
9*	(30, 190)	181.255587654684	3354	2871	6
10	(50, 120)	272.681561136988	3388	3118	7

Table 1: The numerical illustration for $n = 150$, $g = 20$, $t_0 = 10$, $k = 10$, $\Delta = 10$, $T = 4000$, $m\% = 4\%$.

Now, change the parameters settings in this experiment to be: $m\% = 8\%$, $g = 40$,

$t_0 = 20$, $k = 15$, $\Delta = 20$, $T = 10000$, $\alpha_4 = 0.03$ and keep $\alpha_2 = \alpha_3 = 0.005$, $\delta^* = 0.05$.

Since we want to select one of the best $m\% = 8\%$ policies. Then we are shooting for the best 12 policies (20, 190), (20, 40), (30, 190), (20, 60), (20, 50), (20, 150), (30, 150), (20, 70), (30, 70), (20, 170), (30, 140), (30, 50). We consider the policy that is selected by Algorithm (5.1) as a CS if it is one of these 12 policies. The approximated probability of the CS here is $P(CS^*) \geq 1 - \frac{1}{1 - \frac{8 \cdot 40}{100} + 0.005 + 0.005 + 0.03} \geq 0.92$. The results are shown in Table 2 for 10 replications. Obviously, in each replication with star (*), the Algorithm (5.1) select one from the best 8% policies. Again the algorithm performs well one and made the correct selection in 9 out of 10 (i.e. $P(CS) = 0.9$).

Replications	Good Policy	Total Cost	$\sum_{i=1}^g T_i$	$\sum_{i \in I} N_i$	$ I $
1*	(20, 170)	189.758746215487	Σ	Σ	11
2*	(30, 50)	191.628714794078	7125 9312	3334 998	10
3*	(30, 50)	190.567714794078	8801	2155	13
4*	(20, 60)	174.287564862776	7849	3315	12
5	(100, 190)	336.035689745589	6994	3985	7
6*	(20, 190)	166.089889757315	9875	8879	10
7*	(30, 140)	190.015487645425	8197	4235	9
8*	(20, 70)	187.230578488971	9578	3254	12
9*	(30, 150)	178.348796214357	8447	2985	14
10*	(20, 190)	174.055507654357	7540	5540	10

Table 2: The numerical illustration for $n = 150$, $g = 40$, $t_0 = 20$, $k = 15$, $\Delta = 20$, $T = 10000$, $m\% = 8\%$.

These two different settings are repeated for 100 replications, and we summarize the result in Table 3, $\overline{\sum_{i=1}^g T_i}$ = the average number of the total sample size in Step 5, and $\overline{\sum_{i \in I} N_i}$ = the average number of the total sample size in Step 9. From the results, it is clear that the $P(CS)$ is high and closed to the approximate value of $P(CS^*)$ with number of simulation sample size relative small.

$m\%$	T	t_0	Δ	$\sum_{i=1}^g T_i$	$\sum_{i \in I} N_i$	Algorithm $P(CS)$	Approximate $P(CS^*)$
4%	4000	10	10	Σ	Σ	0.48	0.49
				3968	3691		

Table 3: The performance of algorithm for selecting one of a good policy from 150 policies.

Example 2:

In the second experiment, assume the number of policies $n = 500$. We run simulation for each policy in 500 policies, to select the best policy that has the minimum total cost. Then, we apply the Algorithm (5.1) on these 500 policies to select the best policy. Suppose we want to select one of the best $m\% = 2\%$ policies. Then we are shooting for policies (40, 100), (50, 70), (20, 130), (20, 60), (20, 80), (30, 90), (30, 180), (20, 310), (20, 200), (40, 250). We consider the policy that is selected by Algorithm (5.1) as a CS if it is belongs to these 10 policies. Note that, We take $g = 50$, $t_0 = 5$, $k = 10$ and $\Delta = 10$. Let $\alpha_2 = \alpha_3 = 0.005$, $\alpha_4 = 0.05$, $\delta^* = 0.05$, and the total budget $T = 5000$. We can evaluate the approximate probability of the correct selection as $P(CS^*) \geq 1 - \frac{1}{200^{50}} + 0.005 + 0.005 + 0.05 \geq 0.58$. Table 4 contains the results of this experiment, with 10 replications for selecting one of the best 2% policies. Clearly, in each replication with star (*), the Algorithm (5.1) selects one from the best 2% policies. So $P(CS) = 0.6$.

Now, take the parameters in this experiment as, $m\% = 5\%$, $g = 80$, $t_0 = 10$, $k = 20$, $\Delta = 20$, $T = 10000$, $\alpha_4 = 0.02$ and keep $\alpha_2 = \alpha_3 = 0.005$, $\delta^* = 0.05$. Since we want to select one of the best $m\% = 5\%$ policies, then we are shooting for policies (40, 100), (50, 70), (20, 130), (20, 60), (20, 80), (30, 90), (30, 180), (20, 310), (20, 200), (40, 250), (20, 110), (20, 160), (40, 60), (30, 50), (30, 140), (30, 220), (40, 180), (50, 180), (40, 280), (50, 140), (20, 240), (40, 130), (30, 250), (30, 330), (60, 80). We, consider the policy that is selected by Algorithm (5.1) as a correct selection if it is belongs to these 25 policies. The approximate probability of the correct selection here as $P(CS^*) \geq 1 - \frac{1}{300^{80}} + 0.005 + 0.005 + 0.02 \geq 0.95$. The results are shown in Table 5 for 10 replications. Obviously, in each replication with star (*), the

Algorithm (5.1) select one from the best 5% policies, and $P(CS) = 1$.

			$i=1$	$i \in I$	
1*	(20, 80)	187.254687809419	Σ	Σ	8
2	(30, 270)	235.512489453484	4897 5048	3981 4457	6
3*	(40, 250)	249.586556987648	4294	4357	9
4	(40, 60)	221.985311931174	4882	4512	6
5*	(50, 70)	259.893261280024	5187	4345	8
6	(50, 110)	282.681525466988	4915	4473	9
7*	(40, 100)	249.901254388456	5216	4689	9
8*	(20, 310)	233.954812122342	4992	4635	8
9*	(20, 130)	260.550837481201	4368	4194	7
10	(110, 130)	350.121852456789	4112	3685	5

Table 4:
The
numerical

illustration for $n = 500$, $g = 50$, $t_0 = 10$, $k = 10$, $\Delta = 10$, $T = 5000$, $m\% = 2\%$

			$i=1$	$i \in I$		
1*	(30, 180)	232.658223586425	Σ	$\sum_i$		15
2*	(60, 80)	257.522119836398	R_e	8475 9125	3224 2358	18
3*	(20, 130)	251.112437481201	p_j	7554	4055	17
4*	(30, 90)	216.70137183658	i	8842	2997	16
5*	(20, 200)	201.224812122354	c_a	7887	3351	15
6*	(30, 220)	253.521489453484	t	8547	4289	19
7*	(50, 180)	288.024787698588	i_o	7245	4882	11
8*	(40, 60)	217.341521931174	n	8335	2395	14
9*	(30, 90)	219.219471836518	s	7731	4221	17
Good Pol						

Good Policy

Total

$$N_i \quad |I|$$

Table 5: The numerical illustration for $n = 500$, $g = 80$, $t_0 = 10$, $k = 20$, $\Delta = 20$, $T = 10000$, $m\% = 5\%$

These two settings are repeated for 100 replications, and we summarize the result in Table

6. From the results, it is clear that the P (CS) is very high and closed to the approximate value of P (CS*) with number of simulation sample size is small.

		$i=1$		$i \in I$			
2%	5000	5	10	Σ	Σ	0.55	0.58
5%	10000	10	20	4691 8853	3976 3425	0.96	0.95

$m\%$ T t_0 Δ $\overline{\sum_{i=1}^g T_i}$ $\overline{\sum_{i \in I} N_i}$ Algorithm $P(CS)$ Approximate $P(CS^*)$

Table 6: The performance of algorithm for selecting one of a good policy from 500 policies

Example 3:

To study the effect of the simulation parameters on the $P(CS)$, we apply Algorithm (5.1) on a stochastic inventory problem base on some parameters such as the total budget, simulation samples, the size of the set of the actual best policies $m\%$, initial sample size and the increment in simulation samples. In doing so, we apply the algorithm on 150 (s, S) stochastic inventory policies with $g = 40$, $k = 15$ and different values of $t_0 = 20$, $\Delta = 20$, $m\% = 8\%$, and T . Figure 2 shows the relation between the total budget T and the $P(CS)$ for a 100 replication. Clearly, it shows that when we increase the total budget T the $P(CS)$ will increase, and it reaches to the maximum value of 1, when the total budget, T are 12000 and more.

Figure 3 shows the relation between the total budget T and the number of total simulation samples that represents the sum of $\overline{\sum_{i=1}^g T_i}$ and $\overline{\sum_{i \in I} N_i}$, for a 100 replication. It is clear that, by increasing the total budget T , it increases the number of total simulation samples $\overline{\sum_{i=1}^g T_i}$ and $\overline{\sum_{i \in I} N_i}$, that are used in Algorithm (5.1).

To see the effect of size of the set of the actual best policies m , on the performance of the Algorithm (5.1), we repeat the experiment for different values of the good enough solutions $m\%$, where we assume $m\% = 2\%$, $m\% = 4\%$, $m\% = 8\%$, $m\% = 10\%$, and $m\% = 14\%$. Figure 4 shows the effect of m on the $P(CS)$ over 100 replication. As expected when $m\%$ is increased, then $P(CS)$ is increased and reach to 1.

We implement Algorithm (5.1) on three different values of initial sample size t_0 , including

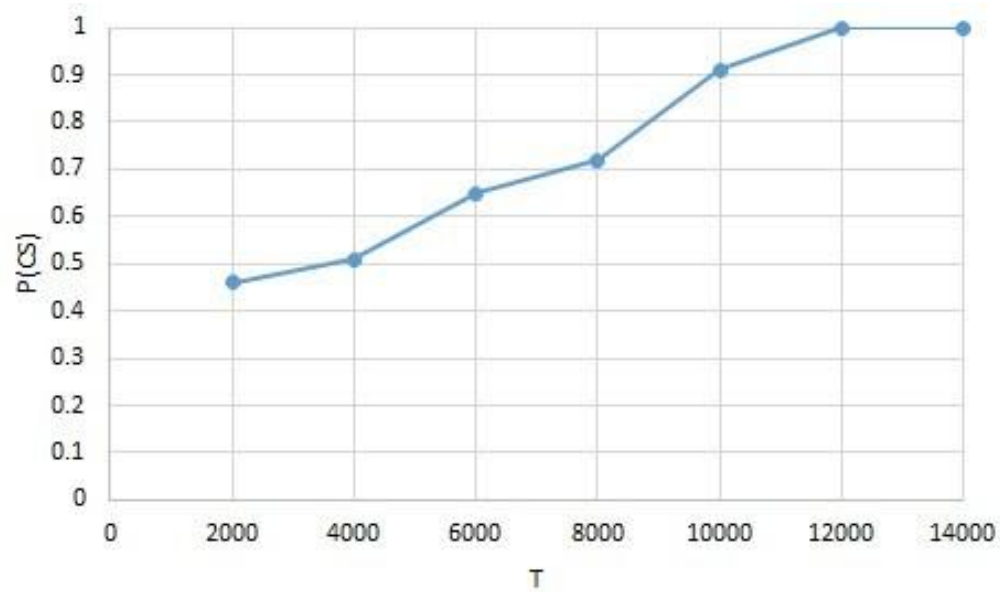


Figure 2: The relationship between the total budget T and $P(CS)$ when $n = 150$, $g = 40$, $t_0 = 20$, $k = 15$, $\Delta = 20$, $m\% = 8\%$ over 100 replications

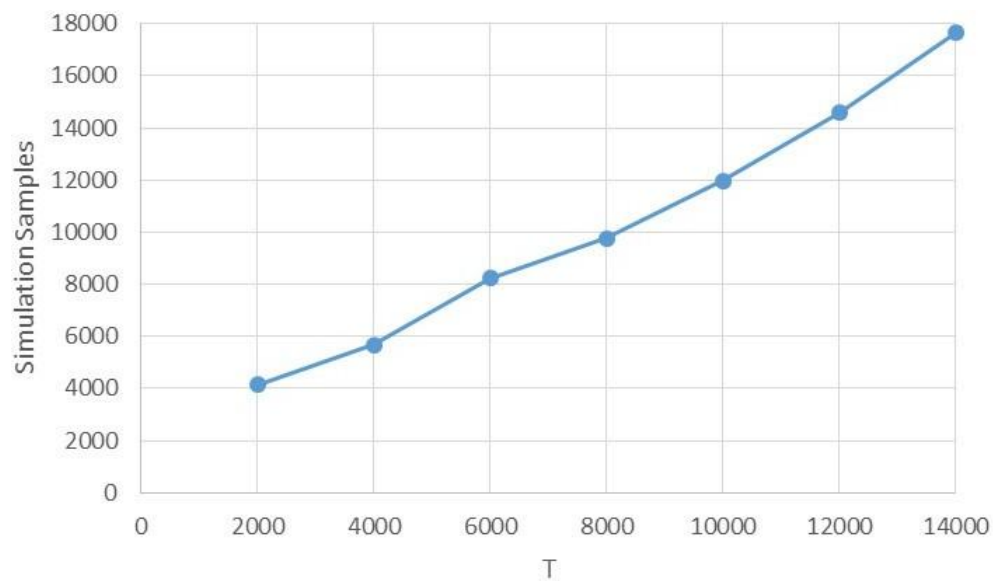


Figure 3: The relationship between the total budget T and the total simulation samples when $n = 150$, $g = 40$, $t_0 = 20$, $k = 15$, $\Delta = 20$, $m\% = 8\%$ over 100 replications

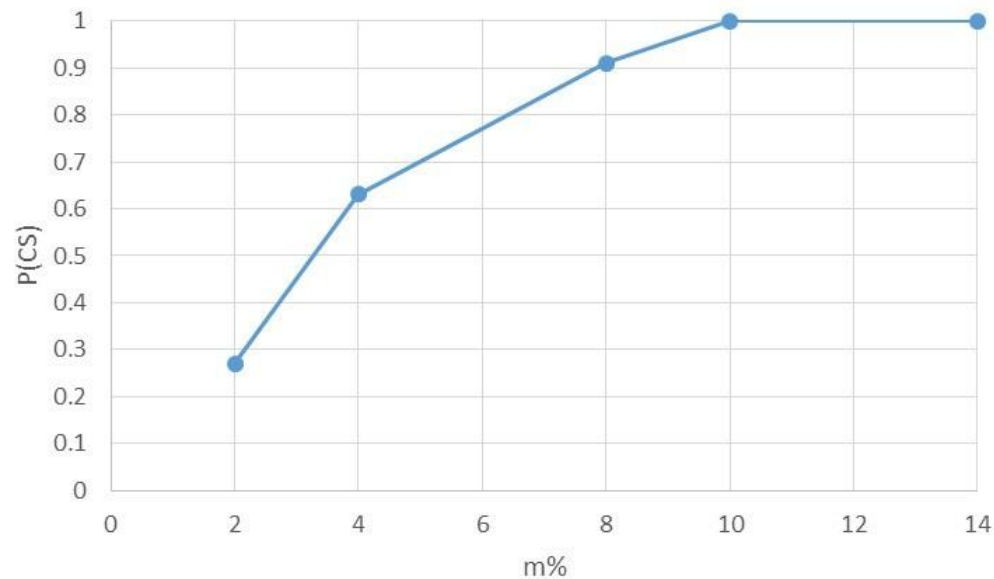


Figure 4: The relationship between the $m\%$ and $P(CS)$ when $n = 150$, $g = 40$, $t_0 = 20$, $k = 15$, $\Delta = 20$, $T = 15000$ over 100 replications.

$t_0 = 5$, 10 and 20. The results are depicted in Figure 5. It is clear that $t_0 = 5$ gives better performance than the other two values of t_0 , and $t_0 = 10$ gives better performance than $t_0 = 20$. This is because the algorithm focuses on the alternative policies that are likely to be selected rather than spending times on the ones that are not of interest.

We implement Algorithm (5.1) on the experiment using three different selection of the value of Δ , these include $\Delta = 10$, 20 and 40. We have used $t_0 = 5$ since it gives better performance. The results are depicted in Figure 6. It is clear that the algorithm gives better performance when the value of $\Delta = 10$ and $\Delta = 20$ and when $\Delta = 20$, the algorithm behaviors better than when $\Delta = 40$. This is because the *OCBA* distribute the available budget on alternative policies that are not likely to be the best policy.

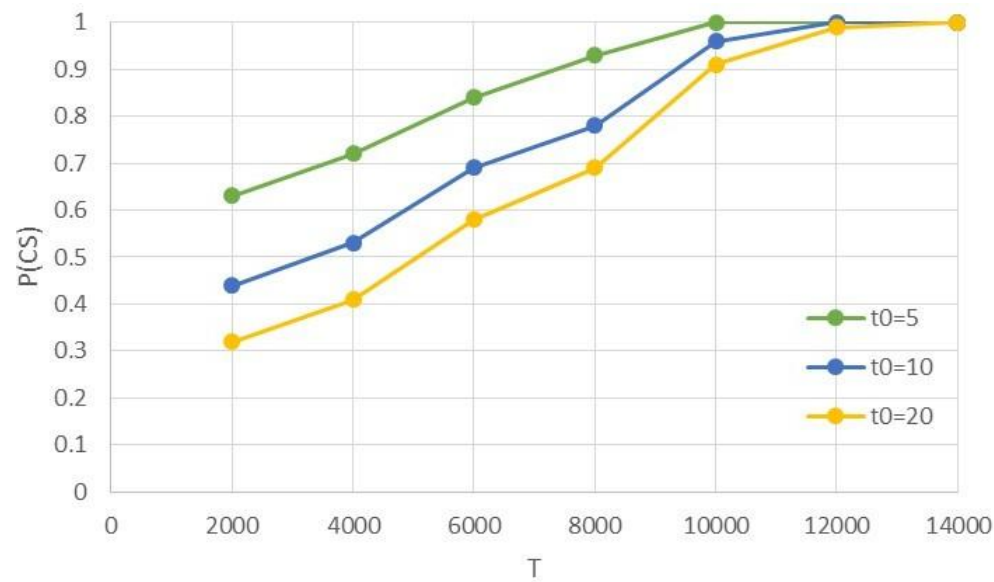


Figure 5: The relationship between the t_0 and $P(CS)$ when $n = 150$, $g = 40$, $k = 15$, $\Delta = 20$, $m\% = 8\%$ over 100 replications.

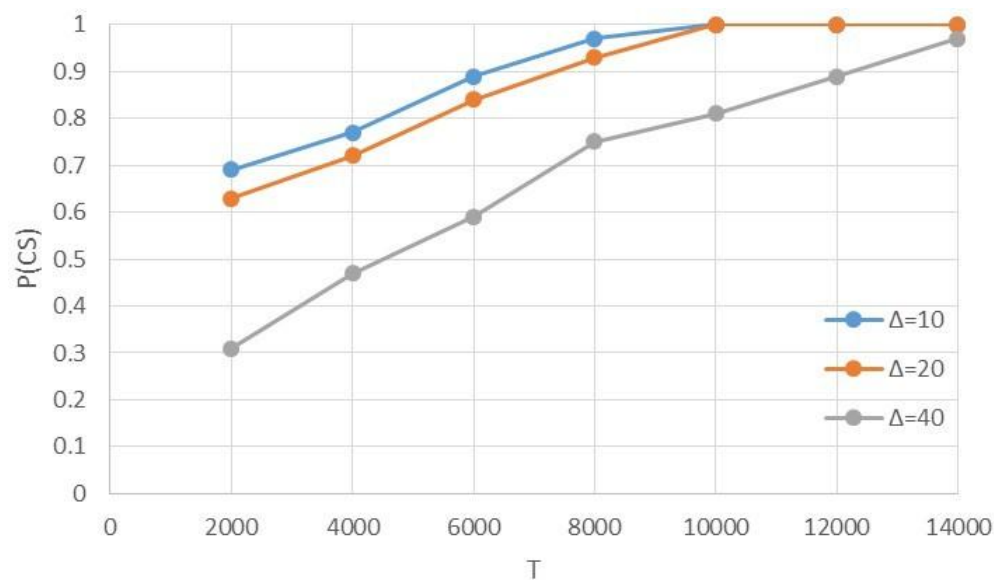


Figure 6: The relationship between the Δ and $P(CS)$ when $n = 150$, $g = 40$, $t_0 = 5$, $k = 15$, $m\% = 8\%$ over 100 replications.

6 Conclusion

In this study, we discussed how to select the best policy with minimum total cost of a stochastic inventory management problem. The total cost includes the ordering, handling and shortage cost. The stochastic (s, S) inventory models is discussed at which an order is placed when the inventory level I falls at or below the reorder point s , in this case an $S - I$ items are ordered. To select the stochastic inventory policy that has the minimum total cost, we have considered a sequential statistical algorithm for selecting a good stochastic inventory policy when the number of alternative policies is large. In this algorithm we start with OO procedure to select relatively a small subset whose intersection with the subset that contains the actual best $m\%$ policies is high. Next we allocate the available computing budget using the $OCBA$ method to achieve the maximum $P(CS)$. Then using the SS procedure, a smaller subset contains the best policy are identified. Finally the IZ procedure is used to select the best policy from the survivors in the previous stage which results in high probability of selecting an alternative policy that belongs to the top $m\%$ policies. This study uses $P(CS)$ as a measure of correct selection. Since we assume the computational budget is available by the algorithm we distribute the available budget in a way that maximizes the $P(CS)$. It is clear that the algorithm is robust and achieves high $P(CS)$ with a small numbers of replication samples.

References

- [1] C.H. Chen, E. Yücesan and S.E. Chick, Simulation budget allocation for further enhancing the efficiency of ordinal optimization, *Discrete Events Dynamic Systems*, **10**, (2000), 251–270.
- [2] S.H. Kim and B.L. Nelson. Selecting the Best System. In *Handbooks in Operations Research and Management Science*. Chapter 17, Elsevier, (2006), 501–534.

- [3] M.H. Almomani and R.Abdul Rahman, Selecting a good stochastic system for the large number of alternatives, *Communications in Statistics-Simulation and Computation*, **41**, (2012), 222–237.
- [4] M.H. Almomani and M.H. Alrefaei, Ordinal optimization with optimal computing budget for selecting an optimal subset, *Asia-Pacific Journal of Operational Research*, **33**, No 2 (2016), 1–17.
- [5] M. Al-Salem, M.H. Almomani, M.H. Alrefaei and A. Diabat, On the optimal computing budget allocation problem for large scale simulation optimization, *Simulation Modelling Practice and Theory*, **71**, (2017), 149–159.
- [6] M.H. Almomani, M.H. Alrefaei and M.H. Almomani. Stopping rules for selecting the optimal subset. *Nonlinear Dynamics and Systems Theory*, **21**, No 1 (2021), 13–26.
- [7] D. Debadyuti. Inventory Models (II) Under SC Environment, [Online]. Available from World Wide Web: <https://slideplayer.com/slide/12877200/> [Accessed 14 July 2022].
- [8] M. Law and W.D. Kelton. Simulation Modeling and Analysis, Mcgraw Hill (2000).
- [9] W.L. Winston and J.B. Goldberg. Operations Research: Applications and Algorithms. Thomson/Brooks/Cole (2004).
- [10] P.D. Tai, P.P.N. Huyen and J. Buddhakulsomsiri. A novel modeling approach for a capacitated (S,T) inventory system with backlog under stochastic discrete demand and lead time. *International Journal of Industrial Engineering Computations*, **12**, (2021), 1-14.
- [11] A.A. Taleizadeh, K. Tafakkori and P. Thaichon. Resilience toward supply disruptions: A stochastic inventory control model with partial backordering under the base stock policy, *Journal of Retailing and Consumer Services*, **58**, (2021), 102291.
- [12] S.H. Razavi Hajiagha, M. Daneshvar, and J. Antucheviciene. A hybrid fuzzy-stochastic multi-criteria ABC inventory classification using possibilistic chance-constrained programming, *Soft Comput* **25**, (2021), 1065-1083.
- [13] S. Zhang, K. Huang and Y. Yuan. Spare parts inventory management: a literature review, *Sustainability*, **13**, No 5, (2021), 2460.
- [14] Y.C. Ho, R.S. Sreenivas, P. Vakili, Ordinal optimization of DEDS, *Discrete Events Dynamic Systems*, **2**, (1992), 61–88.
- [15] S.C. Horng, S.S. Lin, Apply Ordinal Optimization to Optimize the Job-Shop Scheduling Under Uncertain Processing Times, *Arabian Journal for Science and Engineering*, **47**, (2022), 9659-9671.
- [16] D. Shin, M. Broadie, A. Zeevi, Practical Nonparametric Sampling Strategies for

Quantile-Based Ordinal Optimization, *INFORMS Journal on Computing*, **34**, No 2 (2021), 752–768.

- [17] Y.C. Ho, Q.C. Zhao, Q.S. Jia, *Ordinal Optimization: Soft Optimization for Hard Problems*, Springer, (2008).
- [18] C.H. Chen, C.D. Wu and L. Dai, Ordinal comparison of heuristic algorithms using stochastic optimization, *IEEE Transaction on Robotics and Automation*, **15** No 1 (1999), 44–56.
- [19] M.H. Alrefaei, M.H. Almomani and S.N. Alabed Alhadi. Asymptotic approximation of the probability of correctly selecting the best systems, *1st International Conference on Mathematical and Related Sciences (ICMRS 2018): AIP Conference Proceedings*, Vol. 1991, Issue 1, (2018) pp. 020022-1 - 020022-5.
- [20] J. Banks, *Handbook of simulation*, John Wiley, (1998).
- [21] C.H. Chen and Lee, L.H., *Stochastic simulation optimization: an optimal computing budget allocation* (Vol. 1), World scientific, (2011).
- [22] X. Huang and S.H. Choi, An Efficient Simulation-Based Policy Improvement with Optimal Computing Budget Allocation Based on Accumulated Samples, *Electronics*, **11**, No 7 (2022), 1141.
- [23] Y. Myoung Ko, E. Byon, Optimal budget allocation for stochastic simulation with importance sampling: Exploration vs. replication, *IIE Transactions*, **54**, No 9 (2022), 881–893.
- [24] S.S. Gupta, On some multiple decision (selection and ranking) rules, *Technometrics*, **7**, (1965), 225–245.
- [25] V. Roostapour, A. Neumann, F. Neumann, T. Friedrich, Pareto optimization for subset selection with dynamic cost constraints, *Artificial Intelligence*, **302**, (2022), 103597.
- [26] R. Thompson, Robust subset selection, *Computational Statistics & Data Analysis*, **169**, (2022), 107415.
- [27] Y. Rinott, On two-stage selection procedures and related probability-inequalities, *Communications in Statistics–Theory and Methods*, **A7**, (1978), 799–811.
- [28] R.E. Bechhofer, T.J. Santner, D.M. Goldsmanf, *Design and Analysis of Experiments for Statistical Selection, Screening, and Multiple Comparisons*, Wiley, New York, (1995).
- [29] S.H. Kim, B.L. Nelson, Selecting the best system. *Operations Research and Management Science*, **Chapter 17**, (2006), 501–534.

- [30] C. Yu, A. Matta, Q. Semeraro, Group decision making in manufacturing systems: An approach using spatial preference information and indifference zone, *Journal of Manufacturing Systems*, **55**, (2020), 109–125.
- [31] M. Yoon, J. Bekker, Bayesian-based indifference–zone multi-objective ranking and selection procedures, *Computers & Industrial Engineering*, **167**, (2022), 108007.
- [32] B.L. Nelson, J. Swann, D. Goldsman, W. Song, Simple procedures for selecting the best simulated system when the number of alternatives is large, *Operations Research*, **49**, (2001), 950–963.