

2019

Practice-driven solutions for inventory management problems in data-scarce environment

<https://hdl.handle.net/2144/36014>

"Downloaded from OpenBU. Boston University's institutional repository."

BOSTON UNIVERSITY
QUESTROM SCHOOL OF BUSINESS

Dissertation

**PRACTICE-DRIVEN SOLUTIONS FOR INVENTORY
MANAGEMENT PROBLEMS IN DATA-SCARCE
ENVIRONMENTS**

by

LE WANG

B.S., Shanghai Jiao Tong University, 2012
M.Eng., Boston University, 2014

Submitted in partial fulfillment of the
requirements for the degree of
Doctor of Philosophy

2019

© 2019 by
LE WANG
All rights reserved

Approved by

First Reader

Nitin R. Joglekar, Ph.D.
Dean's Research Scholar
Associate Professor of Operations and Technology Management

Second Reader

Sean P. Willems, Ph.D.
Haslam Chair in Supply Chain Analytics
Professor of Business Analytics & Statistics
University of Tennessee, Knoxville

Third Reader

Erol A. Peköz, Ph.D.
Professor of Operations and Technology Management
Boston University, Questrom School of Business

Professor of Systems Engineering
Boston University, College of Engineering

Acknowledgments

My deepest gratitude goes to my advisor Prof. Sean Willems whose enthusiasm and attention to details are motivating. This thesis could not have been possible without his support throughout my five years at Questrom. Sean is one of the most energetic and hard-working people I've ever known. He demonstrates the philosophy of operations management not only in research but also in every corner of life. It is my fortune to have him as my advisor and role model in almost all perspectives: as a researcher, as a teacher, and as a mentor.

Besides my advisor, I would like to thank Prof. Nitin Joglekar and Prof. Erol Pekoz for being my committee members. The work could not have been done without their insightful comments and invaluable advice. I am thankful to Prof. Nitin Joglekar for all the support throughout the years, in both life and research. I am thankful to Prof. Erol Pekoz, for bringing in inspirations during the discussions which eventually resulted in the usage of Metropolis-Hastings sampling technique in the thesis.

I would also like to thank my fellow Ph.D classmates, Emre Guzelsu, Kyung Min Lee, Elnaz Karimi, Temidayo Amaju Adepoku, Arshya Feizi, Aykut Turkoglu, Minje Park for their great companionship. I've learned so much from each of them. A special debt of gratitude is due to my friend Emre Guzelsu, for supporting me in tough times and sharing my joy in good times.

Last but not the least, this thesis is dedicated to my parents, Linan Zhang and Yong Wang who sacrificed so much to provide me with the best educational resources possible. This dissertation would not have been possible without their warm love, continued patience, and endless support.

PRACTICE-DRIVEN SOLUTIONS FOR INVENTORY MANAGEMENT PROBLEMS IN DATA-SCARCE ENVIRONMENTS

LE WANG

Boston University, Questrom School of Business, 2019

Major Professor: Nitin R. Joglekar, Ph.D.

Dean's Research Scholar, Associate Professor of
Operations and Technology Management

ABSTRACT

Many firms are challenged to make inventory decisions with limited data, and high customer service level requirements. This thesis focuses on heuristic solutions for inventory management problems in data-scarce environments, employing rigorous mathematical frameworks and taking advantage of the information that is available in practice but often ignored in literature. We define a class of inventory models and solutions with demonstrable value in helping firms solve these challenges.

Contents

1	Introduction	1
2	Setting Inventory Targets for Low-Demand Products with Discrete Compound Process	4
2.1	Introduction	4
2.2	Literature Review	6
2.3	Model Assumptions	8
2.3.1	Repeated Newsvendor Model with Compound Demand	8
2.3.2	Observed Data on Demand History and Customer Arrivals	9
2.4	Choosing Newsvendor Quantity with General Distribution	10
2.4.1	Integer Patterns under $\mathbf{z}_T$, $\mathbf{d}_T$, and BI	11
2.4.2	Approach I—Maximum Likelihood Estimation	14
2.4.3	Approach II—Metropolis-Hastings Sampler	17
2.4.4	Convergence of the Metropolis-Hastings Sampler	22
2.4.5	MLE vs. MH Sampler	23
2.5	Numerical Experiment	24
2.5.1	Study Design	24
2.5.2	Results for Numerical Experiment	27
2.6	Conclusion	32
3	Set Inventory Targets for Low-Demand SKUs with Integer Pattern Sampling Heuristic	34
3.1	Problem Description	34

3.1.1	Challenges at the W Facility	37
3.2	Integer Pattern Sampling (IPS) Solution for Setting Inventory Targets	40
3.2.1	Uniform Sampler on Θ_{BI_0}	42
3.2.2	Uniform Sampler on Θ_{BI}	43
3.2.3	Estimating Inventory Targets with IPS Heuristic	45
3.3	Numerical Experiment	46
3.4	Implementation at the W Facility	50
3.5	Conclusion	51
4	Identifying Inventory Reduction and Process Improvement Opportunities in Supply Chain	52
4.1	Introduction	52
4.2	Company M's ANT Supply Chain and Current Operations	55
4.2.1	Problem with Only Considering Safety Stock	60
4.3	Maximum Forward Coverage Heuristic	64
4.4	Full Implementation and Discussion	68
5	Conclusions	71
5.1	Summary of the thesis	71
A	Recursion for Calculating $\tau(\delta)$	74
B	Proof Lemma 1	75
C	Proof Theorem 1	77
D	Statistics for Test Cases	78
E	Table appendix	80
F	Generation of Δ_t	84
G	Calculation of $\tau(\delta_t)$	85

H	Generation of All Feasible z_T	86
I	Company M's Approach	87
J	MFC Approach	88
K	Simulation-based Numerical Test on Frederic Facility	91
	Curriculum Vitae	96

List of Tables

2.1	Five types of distribution used in generating a synthetic dataset . . .	24
2.2	Average Optimality Gap between MLE and MH with tight bounds and benchmarks when service requirement is 98%.	28
2.3	Percentage of Underestimation/Optimality/Overestimation of each benchmark algorithm for the 98% quantile.	29
2.4	Percentage of Underestimation/Optimality/Overestimation of MLE, MH-IC, and MH-MHR with tight bounds for the 98% quantile. . . .	29
2.5	Acceptance of an example test case, showing that MH-MHR has a higher acceptance rate than MH-IC.	32
3.1	Mean and standard deviation (in paranthesis) of Optimality Cost Gap of solutions under high service requirements.	49
4.1	Statistics of number of changes in safety stock parameters over 52-week horizon at 5 major DCs in ANT supply chain.	61
4.2	Disguised results of the simulation based test at the Maryland facility. All numbers have been disguised to protect company confidentiality. .	67
4.3	Top 5 category of root causes of excessive inventory identified by MFC.	70
D.1	Statistics of 25 test cases. Each generates 40 random sample paths for the numerical experiment in Section 2.5	79
E.1	Average OCG for MLE, MH, and all benchmark algorithms over 1,000 synthetic test cases.	81

E.2	Standard deviation of OCG for MLE, MH, and all benchmark algorithms over 1,000 synthetic test cases.	82
E.3	Percentage of Under-Estimation/Optimality/Over-Estimation of each benchmark algorithm with service level 90%, 95%, 98%, 99%, 99.5%.	83

List of Figures

2·1	MLE Solution for example above.	15
2·2	Example showing that likelihood function of MLE is not necessarily concave.	16
2·3	Candidate $\mathbf{q}$ moves in the three-dimensional probability simplex: in the first iteration, $i = 0$, and $j = 2$ are chosen, and in the second iteration, $i = 1$, and $j = 2$ are chosen.	20
2·4	MH samples from posterior distribution as more periods with $d_t = 2$ and $z_t = 1$ are observed.	23
2·5	Average Optimality Gap between MLE and MH with tight bounds and benchmarks when service requirement is 98%	27
2·6	Sensitivity of Bounds for MLE and MH when target service level is 98%, $\gamma = 2$ for MH-IC sampler with self-regulating upper bounds. . .	31
2·7	Distribution of y^{MH-IC} and y^{MH-MHR} under different number of iterations.	32
3·1	Repacking process at Company M's W facility. Repacking lead time for varies between 24-48 hours. Procurement lead times are longer depending on specific sources.	36
3·2	Illustration of generating Θ_{BI_0} for $\mathbf{d} = (0, 1, 2, 3)$ and $\tilde{z} = 4$, with wall insertion method.	43
3·3	Flow diagram of how φ -quantile is estimated by IPS solution.	46

3.4	Distribution of mean, coefficient of variation, skewness, and kurtosis of the true distribution of 1,000 synthetic test cases	47
3.5	Under estimation, over estimation, and optimality of each heuristic when target service level is 98%.	50
4.1	Manufacturing and distribution supply chain for ANT product family, which are either made in Company M's facilities in Frederick, Maryland and Eugene, Oregon, or outsourced to OEM.	56
4.2	Usage profile for a ANT SKU example, which is manufactured at Oregon MFG and stored in Maryland DC, showing that most inventory is stored at Maryland DC.	58
4.3	Two SKU examples showing that frequencies of safety stock levels adjustment vary between SKUs.	59
4.4	Left: distribution of weeks of supply of safety stock parameter snapshot for all ANT SKU-L; right: service level equivalence for all ANT SKU-L under Normal assumption.	62
4.5	Examples of three scenarios indicating that the metric "safety stock" is used differently for different SKUs.	63
4.6	Left: distribution of weeks of supply of latest safety stock level snapshot for all ANT SKU-L; right: service level equivalent for all ANT SKU-L under Normal assumption.	64
4.7	Different portions of inventory are identified as reduction opportunities by the M Approach and MFC.	67
4.8	Longitudinal impact on inventory identified by MFC at Company M.	69
J.1	Example of the M Approach with parameter $p = 0.75$, and $p = 0.9$. .	88

List of Abbreviations

ANT		Antibodies
BI		Boundary Information
BOM		Bill of Materials
CAVE		Concave, Adaptive Value Estimation
DC		Distribution Center
DM		Decision Maker
E1		EnterpriseOne Requirement Planning
ERP		Enterprise Resources Planning
FED		Full Empirical Distribution
FG		Finished Goods
IC		Independence Chain
IPS		Integer Pattern Sampling
MFG		Manufacturing Sites
MH		Metropolis-Hastings (sampler)
MHR		Modified-Hit-and-Run
MLE		Maximum Likelihood Estimator
MRP II		Manufacturing Resources Planning
MRP		Material Requirements Planning
OCG		Optimality Cost Gap
OEM		Original Equipment Manufacturer
P.m.f.		Probability mass function
SAA		Sample Average Approximation
SCM		Supply Chain Management
SKU		Stock Keeping Units
SKU-L		Stock Keeping Units at a Location
WIP		Work-in-process

Chapter 1

Introduction

Managing inventory under demand uncertainty remains a core problem in operations management. Parametric inventory models make assumption regarding the functional form of the underlying demand distribution in order to derive inventory targets. However, the parametric paradigm can not provide satisfactory results when the assumed distributions fail to capture the shape of demand realizations in practice. An alternative paradigm is the non-parametric approach, which makes no assumption on the functional form and bases decisions on data. Of course, there is no such thing as free lunch. Non-parametric inventory models avoid restricting themselves to a distribution family at the cost of reliance on data. Moreover, the higher customer service requirements are, the more data is required to capture the tail behavior of demand. In practice, firms often have to make inventory decisions when data are scarce and service requirement is high, in which none of the existing techniques can be implemented directly. Motivated by the collaboration with a major biotechnology company on several inventory management projects, this thesis focuses on non-parametric heuristics that set inventory targets and identify opportunities for reducing stock level in data-scarce environments, where high service requirements are the top priority. Meanwhile, the proposed heuristics also extract information from data on order arrivals, which is collected by most companies but often ignored by the literature on non-parametric inventory models.

Chapter 2 lays the foundation of a theoretical framework. We consider a repeated newsvendor problem for a single product. Instead of employing a single random variable, we model demand as a compound random variable which consists of order arrivals and order quantities, where the decision maker only sees periodic demand and the number of orders in each period. By exploiting the integer nature of demand, and allowing the decision maker to include boundary information on demands, we develop two approaches that employ Maximum Likelihood and Metropolis-Hastings sampler, respectively. In numerical experiments, both approaches outperform benchmarks based only on periodic demand realizations. The approach based on Metropolis-Hastings sampler even outperforms the case in which order quantities are known to the decision maker.

Chapter 3 further relaxes the model in Chapter 2 and assumes that, in addition to periodic demand, the decision maker only knows the total number of orders over the same interval. This reflects a practical challenge at one of the regional distribution centers of the biotechnology company we collaborated with. Despite having limited data, the firm still maintains a high customer service target. We propose a uniform sampling technique that generates random feasible combinations that are consistent with realized observations, as well as optional input on boundary information. The approach was implemented at the facility and set inventory targets for more than 29,000 SKUs.

Chapter 4 takes a fresh look at the objective function of an inventory reduction project at the biotechnology company we collaborated with. It is impractical to apply a new inventory policy because the company's planning infrastructure has already determined the production policy and automated the majority of replenishment orders.

We take a holistic view of the system's constraints to develop a heuristic that rank orders SKUs in terms of risk-free inventory reduction opportunities. The heuristic was developed into a toolbox on the firm's ERP platform and helped identify potential inventory reduction. With the help of this tool, inventory planners at the firm have made more than 22,000 comments regarding the root causes of excessive inventory. Multi-million dollar worth of inventory was reduced from the network between 2017 Q2 and Q3.

Chapter 2

Setting Inventory Targets for Low-Demand Products with Discrete Compound Process

2.1 Introduction

One of the most challenging tasks in inventory management is to characterize the uncertainty of demand. A classical paradigm is to assume a functional form of demand and estimate parameters with data, i.e. past observations. Consequently, the performance of any parametric approach depends heavily on how well the selected distribution matches the “true” distributions. Alternative paradigms include using partial information and empirical distribution function from data. However, such approaches tend to make conservative decisions or require a fair amount of data to obtain satisfying results. When desired service level is higher, the amount of data also increase since more observations are needed to capture the tail behavior.

In practice, firms often confront tougher variants of the problem, in which direct implementation of either parametric or non-parametric models is hard. In this paper, we consider a case that features three common challenges in practice: (1) limited historical data, (2) low demand, and (3) high service requirements, each of which adds to the difficulty of making inventory decisions with available information. Although modern supply chain software has made it possible for firms to track demand

and usage in a timely fashion, it is still a common practice that firms set inventory targets based on weekly demand data, where 52 observations cover all usable data collected over a year-long history. The requirement for quick response also makes it impossible for firms to wait for a long time to collect “enough” information before making decisions. Thus, setting inventory targets under limited data is a practical challenge for firms. In this chapter, we consider the case where firms need to make decisions with less than 12 data points, which mimics a three-month length of data if observations are stored at a weekly interval.

Inventory literature occasionally models demand as a discrete integer random variable. The characteristic of having low demand does not prevent implementation of non-parametric approaches. However, high service requirements is usually problematic for non-parametric models, since they need to wait for long enough time for a small-chance event to show up in observations to account for them. In our model, we assume that the product of interest possesses a high underage cost, and thus a high service level is desired. A key challenge is to exploit the available data and come up with an estimate of the quantile of underlying distribution that is close to one. Motivated by Feeney and Sherbrooke (1966), we model the demand as a general compound process that is the outcome of two independent finite discrete stochastic processes—customer arrival and order quantities. Allowing their probability mass function (p.m.f.) to take any value in a finite dimensional probability simplex, we consider two different approaches—Maximum Likelihood Estimator (MLE) and Metropolis-Hastings (MH) Sampler to estimate inventory target for high service levels ($\geq 95\%$). Both approaches require data on periodic demand and customer arrival, and consider the combinatorial structure of observed data. They also allow the decision maker to include boundary information on order quantities as a part of model input. Through

numerical experiments, we show that MLE has a higher expected cost, yet it is more robust under various tightness of bounds. Metropolis-Hastings sampler, in comparison, is able to obtain lower newsvendor cost when the boundary information is close to true distribution.

The rest of this chapter is organized in the following manner: literature is reviewed in section 2.2; model assumptions are discussed in section 2.3; two algorithms are provided in section 2.4; results of numerical experiments are shown in section 2.5; section 2.6 concludes the chapter.

2.2 Literature Review

Inventory management literature with demand uncertainty falls in one of two existing paradigms: parametric or non-parametric models. Parametric models either assume that both functional form and parameters are known to the decision maker, or the unknown parameters are updated based on the data observed (Scarf, 1959; Iglehart, 1964; Azoury, 1985). Despite the fact that parametric models facilitate the tractability of problems, there always exists the concern that assumed distributions do not fully capture the uncertainty in practice. One solution within the parametric framework is to use a more general distribution family. Numerous studies adopt distributions with more parameters for less restrictive shapes. For instance, Kottas and Lau (1980) show that using two-parameter distributions is insufficient and instead use four-parameter Pearson distribution and the Schmeiser-Deutsch system. More recently, Akcay et al. (2011) applies the Johnson Translation System, a parameterized family of distribution that can match any finite first four moments, and quantify the uncertainty of the estimation by expected total operating cost (Hayes, 1969).

One stream of research assumes that only partial information, instead of the functional form, are available to the decision maker. Scarf (1957) considers the case where only mean and standard deviation are known, based on which the worst-case expected cost is maximized for all possible distributions. Later papers extend Scarf's work and solve similar problems under different settings (Gallego and Moon, 1993; Moon and Gallego, 1994; Gallego et al., 2001). More recently, Perakis and Roels (2008) consider a wider range of possible partial information, and derive order quantities that minimize the newsvendor's maximum regret. Saghaian and Tomlin (2016) develop a maximum-entropy-based technique which takes advantage of both moment and tail information to form a belief-updating procedure for an ambiguity set of possible distributions.

Another stream of research replaces the parametric functional form with the empirical distribution derived from the dataset. Bookbinder and Lordahl (1989) use bootstrap procedure for estimating reorder points. Godfrey and Powell (2001) develop Concave, Adaptive Value Estimation (CAVE) that directly estimates the value function using sample gradient. Kleywegt et al. (2002) introduce the framework of Sample Average Approximation (SAA) that solves empirical counterpart of a stochastic discrete optimization. Levi et al. (2007) establish bounds on the number of samples needed by SAA in a newsvendor setting for desired closeness to optimality. Levi et al. (2015) obtain a tighter bound of SAA on the newsvendor problem with regard to relative regret. The non-parametric newsvendor model with censored sales data rather than demand observations has also raised attention. Huh and Rusmevichientong (2009) propose an adaptive policy under censored demand. Huh et al. (2011) consider a similar problem and develop an adaptive policy with Kaplan-Meier Estimator.

Our paper is different from the existing literature in several aspects. First, we focus on the average performance under small dataset and high service level for low-demand products rather than their limiting behavior. Second, we model the demand as a compound process rather than a single random variable. Third, we apply Metropolis-Hastings sampler in generating posterior distribution for a general form of demand distribution. To our best knowledge, this is the first paper that models demand as a general compound process and proposes algorithms that prescribe inventory targets in a data-scarce environment.

2.3 Model Assumptions

2.3.1 Repeated Newsvendor Model with Compound Demand

Consider a single-item newsvendor problem over discrete time horizon $t \in \mathcal{T} := \{1, 2, \dots\}$ with compound demand process. The decision maker (referred as DM hereafter) starts each period with zero on-hand inventory and needs to decide how many units y to stock prior to demand realization. All unsold units are scrapped at the end of each period. In period t , Z_t customers arrive, with the k -th customer ordering $W_{t,k}$ units of products. Z_t Are discrete random variables with left support 0. $W_{t,k}$ are identically distributed discrete random variables that are mutually independent of one another and independent of Z_t . The demand in period t , denoted by D_t , is the sum of customer orders,

$$D_t := \begin{cases} 0 & \text{if } Z_t = 0 \\ \sum_{k=1}^{Z_t} W_{t,k} & \text{if } Z_t > 0 \end{cases} \quad (2.1)$$

After observing demand in period $1, \dots, T$, the DM's goal is to pick y_{T+1} so that the

total expected cost in period $T + 1$ is minimized, i.e.

$$\min_{y \geq 0} E[L(y, D_{T+1})] = [C_o(y - D_{T+1})^+ + C_u(D_{T+1} - y)^+] \quad (2.2)$$

where C_o is the per-unit overage cost for each unsold unit at the end of each period, C_u is the per-unit underage cost for lost sales penalty, and $x^+ = \max\{x, 0\}$. In this standard newsvendor setting, the optimal order quantity satisfies a critical fractile result,

$$y^{NV}(\varphi) := \inf\{y : F_D(y) \geq \varphi\} \quad (2.3)$$

where F_D is the c.d.f. of D_t for all $t = 1, 2, \dots$, and $\varphi = \frac{C_u}{C_u + C_o}$ is the critical fractile. In practice, it is usually difficult to quantify C_u and C_o . Instead, firms directly set φ as their service target and make stocking decisions accordingly. Therefore, instead of specifying C_u and C_o , we can write the expected newsvendor cost as $C_D(y, \varphi) = E[L(y, D_{T+1})]$. Equivalently, we consider φ as the sole parameter in the problem.

2.3.2 Observed Data on Demand History and Customer Arrivals

When approaching the decision of ordering quantity, the DM has collected the following information:

- ***Demand History*** The DM collects each period's realized demand realization d_t . Vector $\mathbf{d}_T = (d_1, \dots, d_T)$ stores all d_t realizations by the end of period T . If the DM only observes $\mathbf{d}_T$ but not $\mathbf{z}_T$, then the problem is a standard repeated newsvendor model.

- **Customer Arrivals** In each period, the DM also collects the number of customers z_t that arrive in period t . Such information can be obtained by simply counting the number of transactions within a period or distinct order numbers. Let vector $\mathbf{z}_T = (z_1, \dots, z_T)$ stores all z_t realizations by the end of period T .
- **Boundary Information on Order Quantity** In addition to the actual observed data, the DM may also have some prior knowledge, BI , regarding customer purchase behavior. We assume that $BI = \{\underline{w}, \bar{w}\}$, where $0 \leq \underline{w} \leq W_{t,k} \leq \bar{w}$ represents the DM's belief of the support of order quantity. If the dataset stores customer arrivals with no purchase, then $\underline{w} = 0$, otherwise, $\underline{w} = 1$ unless other lowerbound is assumed. Since BI only reflects DM's belief of order distribution, we do not require BI to match the exact bound of the true distribution. However, BI must be consistent with the data, i.e. $z_t \underline{w} \leq d_t \leq z_t \bar{w}$ for all $t = 1, \dots, T$. Finally, let $BI_0 = \{\underline{w} = 0, \bar{w} = \infty\}$ be the non-informative boundary information indicating that DM does not have any useful prior knowledge. We assume that the DM does not have any prior knowledge besides $\mathbf{d}_T$, $\mathbf{z}_T$, and BI , nor does the DM assume any functional form of the underlying distributions.

2.4 Choosing Newsvendor Quantity with General Distribution

For any given period, the compound random variable D has the following c.d.f.:

$$F_D(d) = \sum_{z=1}^{\infty} P\{W_1 + \dots + W_z \leq d | Z = z\} P\{Z = z\} \quad (2.4)$$

If the DM knows the values of $P\{Z\}$ and $P\{W\}$, he can numerically derive $F_D(d)$ using Equation (2.4). Based on the assumption that Z and W_k are independent, $P\{Z\}$ and $P\{W\}$ can be estimated separately. In this model formulation, $\mathbf{z}_T$ is available to the DM. Without other prior knowledge on Z_t , the best estimate for $P\{Z\}$, in terms of relative entropy, is its empirical counterpart. In this sections, we propose two approaches to estimate $P\{W\}$. To estimate the φ quantile of F_D , we evaluate Equation (2.4) by substituting $P\{Z\}$ by its empirical counterparts, and $P\{W\}$ by its estimates from each approach. Let $q_j := P\{W_{t,k} = j\}$ be the probability that a customer orders j units of product, and let $\mathbf{q} = (q_{\underline{w}}, \dots, q_{\bar{w}})$ be the probability mass function (p.m.f.) of $W_{t,k}$, both approaches start with investigating feasible integer combinations given observed data of the compound process.

2.4.1 Integer Patterns under $\mathbf{z}_T$, $\mathbf{d}_T$, and BI

Although $w_{t,k}$ are not observed in dataset, the integer nature of $\mathbf{d}_T$ and $\mathbf{z}_T$ reveals information regarding the feasible combinations of order quantities, while the low-demand nature keeps the size of all combinations within a tractable amount. Order quantities in period t must satisfy $\sum_{k=1}^{z_t} w_{t,k} = d_t$ and $\underline{w} \leq w_{t,k} \leq \bar{w}$. For instance, when $z_t = 1$ and $d_t > 0$, the only possibility is that all d_t units are demanded by a single customer, with $w_{t,1} = d_t$. In general, there are fewer feasible combinations when $\bar{w} - \underline{w}$ is tighter. In order to represent all feasible combinations, we define

$$\Delta_t := \left\{ (w_{t,1}, \dots, w_{t,z_t}) \left| \sum_{k=1}^{z_t} w_{t,k} = d_t, \underline{w} \leq w_{t,1} \leq \dots \leq w_{t,k} \leq \bar{w}, \forall k = 1, \dots, z_t \right. \right\} \quad (2.5)$$

as the set that contains all feasible combinations of $w_{t,k}$ in period t , which are consistent with z_t , d_t and BI . For instance, when $\mathbf{d}_3 = (0, 3, 5)$, $\mathbf{z}_3 = (1, 2, 3)$, and

non-informative boundary information BI_0 , then the corresponding Δ 's are:

$$\begin{aligned}\Delta_1 &= \{(0)\}, & \Delta_2 &= \{(0, 3), (1, 2)\}, \\ \Delta_3 &= \{(0, 0, 5), (0, 1, 4), (0, 2, 3), (1, 1, 3), (1, 2, 2)\}\end{aligned}\tag{2.6}$$

Note that BI influences the outcome of Δ_t significantly. For instance, if $BI = \{\underline{w} = 0, \overline{w} = 3\}$ instead of BI_0 , then $(0, 0, 5)$ and $(0, 1, 4)$ will be removed from Δ_3 . The ordering of $w_{t,k}$ is introduced to reduce redundant representations with the same conditional probability under z_t and $\mathbf{q}$. For example, the unordered $(0,1)$ and $(1, 0)$ have the same probability of q_0q_1 conditioned on z_t and $\mathbf{q}$. Let $\delta_t^* \in \Delta_t$ denote elements in Δ_t , with superscript $*$ denoting that there can be multiple feasible combinations in Δ_t . δ_t^* can be generated iteratively as follow: first, set a candidate δ_t^* to be an empty vector and initiate the iteration by setting the first order quantity $w_{t,1}$ to be $\underline{w}, \dots, \min\{\overline{w}, d_t\}$. Then $w_{t,2}, \dots, w_{t,z_t-1}$ are considered in turns, with each $w_{t,k} \geq w_{t,k-1}$. Similar to a branch and bound algorithm, feasibility check is performed upon adding $w_{t,k}$ to an evolving δ_t^* . Infeasible δ_t^* s are removed from the iteration. Eventually, w_{t,z_t} is set to be $d_t - \sum_{k=1}^{z_t-1} w_{t,k}$. If w_{t,z_t} satisfies $w_{t,z_t-1} \leq w_{t,z_t} \leq \overline{w}$, then a feasible δ_t^* is added to the final Δ_t output. The pseudo-code for generating Δ_t is as follow:

Algorithm 1 Iteration for generating Δ_t under d_t , z_t , and BI

```

1: procedure GENDELTA( $d_t, z_t, \underline{w}, \bar{w}$ )
2:   Initiate  $\Delta_t \leftarrow \{(\underline{w}), (\underline{w} + 1), \dots, (\min\{\bar{w}, d_t\})\}$ 
3:   for  $k \leftarrow 2 : z_t$  do
4:     for  $\delta_t^* \in \Delta_t$  do
5:        $\Delta_t \leftarrow \Delta_t \setminus \{\delta_t^*\}$ 
6:       if  $k < z_t$  then
7:         for  $w_{t,k}^* \leftarrow w_{t,k-1} : \min\{\bar{w}, d_t - \mathbf{1}^T \delta\}$  do
8:           if  $w_{t,k}^* \leq (d_t - \sum_{s=1}^{k-1} w_{t,s}^\delta - w_{t,k}^*) / (z_t - k) \leq \bar{w}$  then
9:              $\Delta_t \leftarrow \Delta_t \cup \{(w_1^\delta, \dots, w_{k-1}^\delta, w_k^*)\}$ 
10:          end if
11:        end for
12:      else if  $k = z_t$  and  $w_{t,z_t-1}^\delta \leq d_t - \mathbf{1}^T \delta \leq \bar{w}$  then
13:         $w_k^* \leftarrow (d_t - \mathbf{1}^T \delta)$ 
14:         $\Delta_t \leftarrow \Delta_t \cup \{(w_1^\delta, \dots, w_{z_t-1}^\delta, w_k^*)\}$ 
15:      end if
16:    end for
17:  end for
18:  Return  $\Delta_t$ .
19: end procedure

```

Let $\tau(\delta)$ denote the number of unordered permutations of δ . In the example (2.6), we have $\tau((0, 0, 5)) = 3$ and $\tau((0, 2, 3)) = 6$. For the general case, $\tau(\delta_t^*)$ can be obtained by a simple iteration (Appendix A). Let $f(\cdot)$ denote the mapping between δ_t^* and its permutations to the sum of their conditional probability (condition on z_t). Under p.m.f. $\mathbf{q}$,

$$f(\delta_t^* | \mathbf{q}, z_t, BI) = \tau(\delta_t^*) \prod_{k=1}^{z_t} q_{w_{t,k}} \quad (2.7)$$

Let $g(\cdot)$ be the mapping of Δ_t to its conditional probability (also conditioned on z_t). $g(\Delta_t | \mathbf{q}, z_t, BI)$ is the sum of all $f(\delta_t^* | \mathbf{q}, z_t, BI)$,

$$g(\Delta_t | \mathbf{q}, z_t, BI) = \sum_{\delta_t^* \in \Delta_t} f(\delta_t^* | \mathbf{q}, z_t, BI) \quad (2.8)$$

Using Δ_t as a building block, we can define $\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI} = \Delta_1 \times \dots \times \Delta_T$ to denote

the set of all combinations from period 1 up to T . Each element $\theta_T^* \in \Theta_{\mathbf{d}_T, \mathbf{z}_T, BI}$ represents combination possibilities $\theta_T^* = (\delta_1^*, \dots, \delta_T^*)$ and their permutations, with probability

$$P\{\theta_T^* | \mathbf{q}, \mathbf{z}_T, BI\} = \prod_{t=1}^T f(\delta_t^* | \mathbf{q}, z_t, BI) \quad (2.9)$$

And the probability of all feasible combinations can be expressed by

$$\begin{aligned} P\{\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI} | \mathbf{q}, \mathbf{z}_T, BI\} &= \sum_{\theta_T^* \in \Theta_{\mathbf{d}_T, \mathbf{z}_T, BI}} P\{\theta_T^* | \mathbf{q}, \mathbf{z}_T, BI\} \\ &= \sum_{\theta_T^* \in \Theta_{\mathbf{d}_T, \mathbf{z}_T, BI}} \left[\prod_{t=1}^T f(\delta_t^* | \mathbf{q}, z_t, BI) \right] \\ &= \prod_{t=1}^T \left[\sum_{\delta_t^* \in \Delta_t} f(\delta_t^* | \mathbf{q}, z_t, BI) \right] \\ &= \prod_{t=1}^T g(\Delta_t | \mathbf{q}, z_t, BI) \end{aligned} \quad (2.10)$$

Armed with the combinatorial structure and the capability of calculating $\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI}$, we propose two approaches to estimate $\mathbf{q}$, from which we can also estimate any quantile of D_t .

2.4.2 Approach I—Maximum Likelihood Estimation

One intuitive way of estimating $\mathbf{q}$ is to use Maximum Likelihood Estimation (MLE) and directly find the parameters that maximize the log-likelihood function of the observed data. MLE is known for being the consistent estimator with the lowest asymptotic mean squared error (Cramer-Rao lower bound) when sample size goes to

infinity. In this case, the log-likelihood function sums over all feasible combinations, i.e. $\log(P\{\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI} | \mathbf{q}, \mathbf{z}_T, BI\})$. MLE solves the nonlinear optimization numerically over the probability simplex $\{\mathbf{q} | \sum_{j=\underline{w}}^{\bar{w}} q_j = 1, q_j \geq 0, \forall j = \underline{w}, \dots, \bar{w}\}$.

$$\begin{aligned} \max_{\mathbf{q}} \quad & \log(P\{\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI} | \mathbf{q}, \mathbf{z}_T, BI\}) \\ \text{subject to} \quad & \sum_{j=\underline{w}}^{\bar{w}} q_j = 1 \\ & q_j \geq 0, \quad \forall j = \underline{w}, \dots, \bar{w} \end{aligned} \tag{2.11}$$

For instance, when $\mathbf{d}_3 = (0, 3, 5)$, $\mathbf{z}_3 = (1, 2, 3)$, and $BI = \{\underline{w} = 0, \bar{w} = 2\}$, we have $\Delta_1 = \{(0)\}$, $\Delta_2 = \{(1, 2)\}$, $\Delta_3 = \{(1, 2, 2)\}$ with corresponding $\tau(\delta_t^*)$, Equation (2.11) becomes

$$\begin{aligned} \max_{\mathbf{q}} \quad & \log(q_0) + \log(2q_1q_2) + \log(3q_1q_2^2) \\ \text{subject to} \quad & q_1 + q_2 + q_3 = 1 \\ & q_j \geq 0, \quad j = 1, 2, 3 \end{aligned}$$

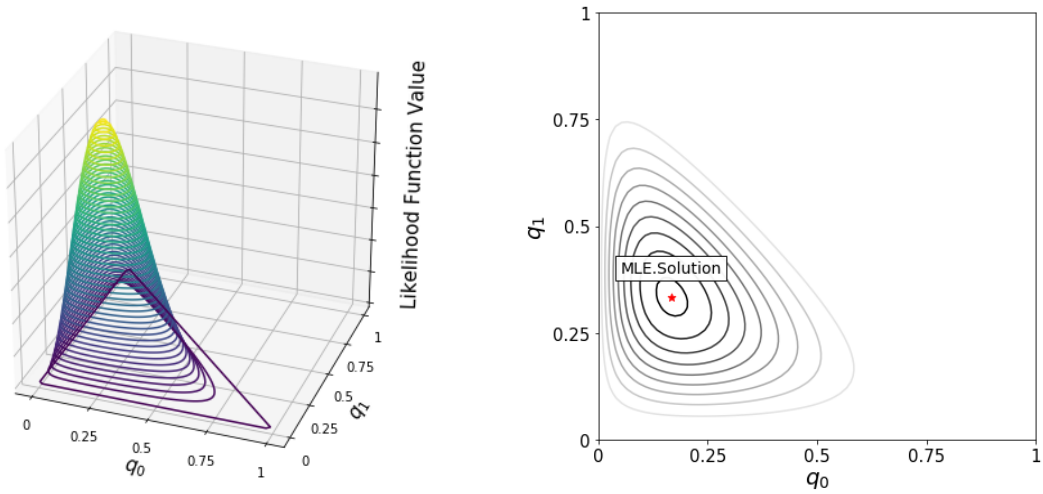


Figure 2.1: MLE Solution for example above.

Figure 2.1 shows the contour plot of the likelihood function (instead of the log-likelihood function for better illustration) of the 3-period example above. Sequential Least Squares Programming Algorithm (SLSQP) is used to numerically solve the non-linear programming problem in Equation (2.11) with multiple start points to avoid being trapped in local maxima. Due to the equality constraint, the actual dimensionality of the non-linear programming is $\bar{w} - \underline{w}$, with the last dimension expressed as $q_{\bar{w}} = 1 - \sum_{j=\underline{w}}^{\bar{w}-1} q_j$.

Although the log-likelihood and likelihood functions seem to be concave in the above example, it is not necessarily always true. In fact, we can construct a simple 2-period example with the same $BI = \{\underline{w} = 0, \bar{w} = 2\}$ whose objective function is not concave in the probability simplex. For instance, let $\mathbf{d}_2 = (0, 2)$, and $\mathbf{z}_2 = (1, 3)$, the contour plot of the likelihood function is shown in Figure 2.2.

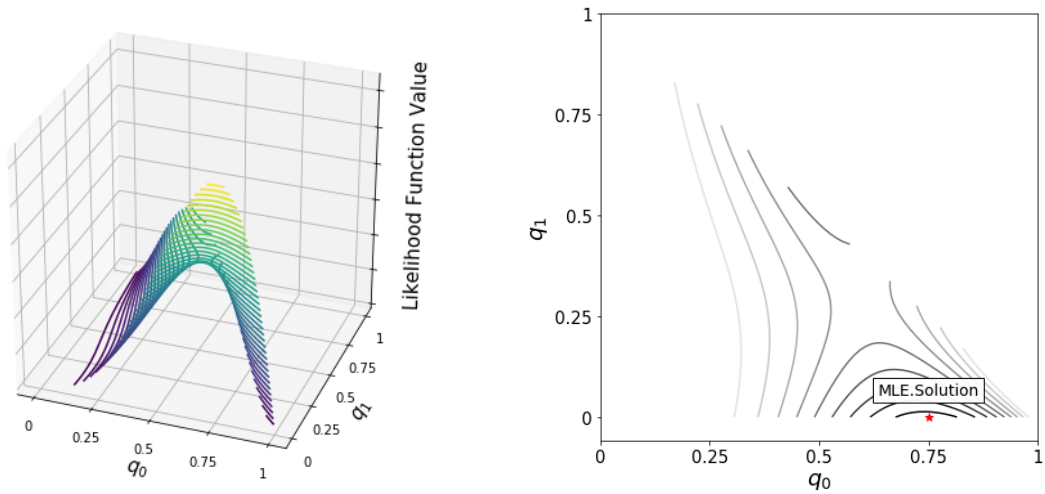


Figure 2.2: Example showing that likelihood function of MLE is not necessarily concave.

For given $\mathbf{d}_T$ and $\mathbf{z}_T$, the MLE approach finds $\mathbf{q}^{MLE}$ that maximizes the probability of observing the data by solving Equation (2.11). To obtain an estimate

of F_D , the MLE evaluates Equation (2.4) by substituting $P\{W = w\}$ with q_w^{MLE} and $P\{Z = z\}$ with its empirical counterparts. Let $\hat{F}_D^{MLE}$ denote the estimated c.d.f. by the MLE. y^{MLE} is obtained by taking the φ quantile of $\hat{F}_D^{MLE}$.

2.4.3 Approach II—Metropolis-Hastings Sampler

The second approach takes a Bayesian view of the problem and generates a posterior distribution after observing $\mathbf{d}_T$ and $\mathbf{z}_T$, achieved by the Metropolis-Hastings (MH) sampler. MH sampler is a Markov Chain Monte Carlo (MCMC) method that generates a multidimensional random samples $\mathbf{X}^{(1)}, \mathbf{X}^{(2)}, \dots$, whose distribution converges to a full joint distribution $\pi(\mathbf{x})$. It was first proposed by Metropolis et al. (1953) in statistical physics and later generalized by Hastings (1970) to accommodate general applications. By carefully constructing a time-reversible Markov chain, the MH sampler is often used to simulate distributions that are hard to sample directly (see Chib and Greenberg, 1995). The MH sampler has three components: (1) a *proposal probability* $\phi(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand})$ that generates a candidate sample $\mathbf{x}^{cand}$ given a previously generated sample $\mathbf{x}^{(i-1)}$; (2) an *acceptance function* $\alpha(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand})$ that sets a probability threshold with which the candidate sample will be accepted or rejected, based on both the proposal probability and the desired full joint distribution $\pi(\mathbf{x})$. The general MH sampler has an acceptance function:

$$\alpha(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand}) = \begin{cases} \min \left\{ 1, \frac{\phi(\mathbf{x}^{cand}, \mathbf{x}^{(i-1)})\pi(\mathbf{x}^{cand})}{\phi(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand})\pi(\mathbf{x}^{(i-1)})} \right\} & , \text{ if } \phi(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand})\pi(\mathbf{x}^{(i-1)}) > 0 \\ 1 & , \text{ if } \phi(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand})\pi(\mathbf{x}^{(i-1)}) = 0 \end{cases} \quad (2.12)$$

(3) a *random decision* to accept candidate sample, with probability $\alpha(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand})$, or reject it with probability $1 - \alpha(\mathbf{x}^{cand}, \mathbf{x}^{(i-1)})$. In the case of

estimating order quantities, the p.m.f. of order quantity $\mathbf{q}$ is the state of interest. We use MH sampler to generate samples from the posterior distribution of $\mathbf{q}$ after observing $\mathbf{d}_T$ and $\mathbf{z}_T$. The boundary information BI is an optional input.

Proposal Probability To implement MH Sampling efficiently, the proposal probability should be carefully chosen so that the state space can be efficiently explored with finite number of samples, and the rejection rate $1 - \alpha(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand})$ should be low to save computational efficiency. Many proposal probabilities have been proposed in the literature. Three of them are commonly used—*Symmetric Probabilities*, *Independence Chain*, and *Gibbs Sampler*. *Symmetric Probabilities*, or often referred as the *Metropolis Update*, uses a symmetric proposal probability $\phi(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand}) = \phi(\mathbf{x}^{cand}, \mathbf{x}^{(i-1)})$. The symmetry of the proposal probability removes its calculation from the acceptance function since they appear in both the nominator and demoninator in Equation (2.12). A well-known example of such probability is the Random Walk, in which $\mathbf{x}^{cand} = \mathbf{x}^{(i-1)} + \mathbf{e}$. To arrange the symmetry property, $\mathbf{e}$ is chosen to be stochastically independent of $\mathbf{x}^{(i-1)}$ and has a symmetrical distribution around $\mathbf{0}$. Normal distributions with zero mean and symmetic uniform distributions about zero are often selected in this group (see Metropolis et al., 1953; Müller, 1991; Hastings, 1970). *Independence chain* applies $\phi(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand}) = \phi_{IC}(\mathbf{x}^{cand})$, where $\mathbf{x}^{cand}$ is drawn from a distribution $\phi_{IC}(\cdot)$, which is independent of $\mathbf{x}^{(i-1)}$. (see Hastings, 1970; Tierney, 1994). *Gibbs sampler*, as a special case of MH sampler, is proposed by Geman and Geman (1987) and uses the conditional joint distribution of one component given the rest of the components. As a result, Gibbs Sampler requires the capability to sample from the conditional joint distribution, which is not met in this problem setting. One motivating approach is the “Hit-and-Run” Algorithm (see Bélisle et al., 1993; Geman and Geman, 1987), in which a random direction is

generated and a random point is chosen along the direction. Hit-and-Run is often used to generate samples when the state space is bounded. Since $\mathbf{q}$ is bounded by the probability simplex, we use the following two proposal distribution to generate candidates— $\phi_{IC}(\mathbf{q})$ for independence chain (*IC*) and $\phi_{MHR}(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand})$ for Modified Hit-And-Run (*MHR*):

- $\phi_{IC}(\mathbf{q}^{cand})$ Generates a sample from the Dirichlet distribution with shape parameter $\mathbf{1}$, in which each candidate point is selected with probability $\mathbf{B}(\mathbf{1}) = 1/\Gamma(\bar{w} - \underline{w} + 1) = 1/(\bar{w} - \underline{w})!$. It is equivalent to sampling the probability simplex uniformly.
- $\phi_{MHR}(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand})$ Generates $\mathbf{q}^{cand}$ based on $\mathbf{q}^{(i-1)}$. In each iteration, two indices i and j , where $\underline{w} \leq i < j \leq \bar{w}$ are randomly picked. Meanwhile, a random number V is drawn from $Unif(0, 1)$. The candidate state $\mathbf{q}^{cand}$ is generated by redistributing probabilities between the chosen indices according to V , while keeping other components in $\mathbf{q}$ unchanged, or formally,

$$\begin{aligned} q_i^{cand} &= V(q_i^{(i-1)} + q_j^{(i-1)}) \\ q_j^{cand} &= (1 - V)(q_i^{(i-1)} + q_j^{(i-1)}) \\ q_k^{cand} &= q_k^{(i-1)}, \quad \forall k \neq i, j \end{aligned} \tag{2.13}$$

Instead of choosing a direction along one of the dimensions in the classical “Hit-and-Run”, we choose a direction parallel to an “edge” of the probability simplex. Due to the symmetry of V , ϕ_{MHR} is also symmetric, $\phi_{MHR}(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) = \phi_{MHR}(\mathbf{q}^{cand}, \mathbf{q}^{(i-1)})$. Figure 2.3 shows how candidate $\mathbf{q}$ moves along the three-dimensional probability simplex under ϕ_{MHR} .

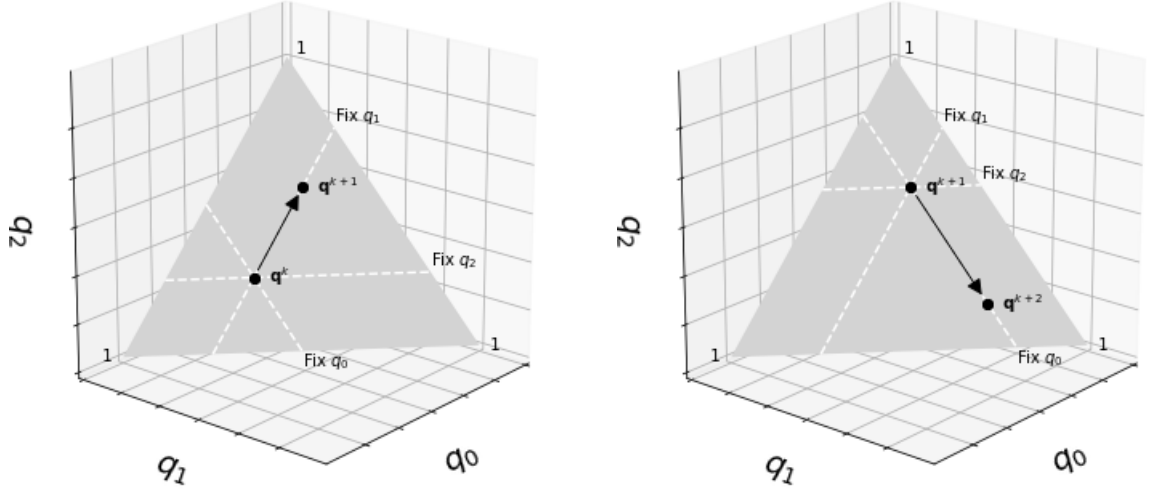


Figure 2.3: Candidate $\mathbf{q}$ moves in the three-dimensional probability simplex: in the first iteration, $i = 0$, and $j = 2$ are chosen, and in the second iteration, $i = 1$, and $j = 2$ are chosen.

Full Joint Distribution We want to simulate samples that follow the posterior distribution of $\mathbf{q}$ given observed data, $\mathbf{d}_T$, $\mathbf{z}_T$ and BI . Let $\pi_T(\mathbf{q})$ be the full joint distribution we want to simulate after observing T periods of data, we have $\pi_T(\mathbf{q}) \equiv p(\mathbf{q}|\mathbf{d}_T, \mathbf{z}_T, BI)$. Let $p(\mathbf{q})$ be a prior of $\mathbf{q}$, we have

$$\pi_T(\mathbf{q}) \propto P(\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI} | \mathbf{q}, \mathbf{z}_T) p(\mathbf{q}) \quad (2.14)$$

Acceptance Function Based on the above proposal probabilities and full joint distribution, we have the acceptance function using $\phi_S, S \in \{IC, MHR\}$ after T periods as:

$$\alpha_{S,T}(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) = \min \left\{ 1, \frac{\phi_S(\mathbf{q}^{cand}, \mathbf{q}^{(i-1)}) \pi_T(\mathbf{q}^{cand})}{\phi_S(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) \pi_T(\mathbf{q}^{(i-1)})} \right\} \quad (2.15)$$

Note that $\phi_S(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) = \phi_S(\mathbf{q}^{cand}, \mathbf{q}^{(i-1)})$ for both transition probabilities. Thus, for either proposal probability, we have the same acceptance function:

$$\begin{aligned}
\alpha_T(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) &= \min \left\{ 1, \frac{\pi_T(\mathbf{q}^{cand})}{\pi_T(\mathbf{q}^{(i-1)})} \right\} \\
&= \min \left\{ 1, \frac{P(\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI} | \mathbf{q}^{cand}, \mathbf{z}_T) p(\mathbf{q}^{cand})}{P(\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI} | \mathbf{q}^{(i-1)}, \mathbf{z}_T) p(\mathbf{q}^{(i-1)})} \right\}
\end{aligned} \tag{2.16}$$

For each sample in the posterior distribution, a φ quantile of D_t can be calculated by Equation (2.4), with $P\{Z_t = z_t\}$ substituted by its empirical counterpart $\hat{P}\{Z_t = i\} = \frac{1}{T} \sum_{t=1}^T \mathbb{I}_{[z_t=i]}$. Let $y = h(\mathbf{q}, \varphi; \mathbf{z}_T)$ be the mapping from $\mathbf{q}$ to a φ quantile of under $\hat{P}\{Z_t = i\}$, M be the number of samples generated, the pseudo-code of the MH sampler is described as follow:

Algorithm 2 MH Sampler for $\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI}$

```

1: procedure MH SAMPLING( $\mathbf{d}_T, \mathbf{z}_T, BI, \phi_S, p(\mathbf{q}), \varphi, \mathbf{q}^0, M$ )
2:   Initiate  $\Upsilon \leftarrow \{\}$ 
3:   Generate  $\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI}$ 
4:   for  $m \leftarrow 1 : M$  do
5:     (i)  $\mathbf{q}^{cand} \leftarrow \phi_S(\mathbf{q}^{cand} | \mathbf{q}^{(i-1)})$ 
6:     (ii) Calculate  $\alpha_T(\mathbf{q}^{(m-1)}, \mathbf{q}^{cand})$  with Equation (2.16)
7:     (iii) Generate  $U$  from Uniform(0,1)
8:     if  $U < \alpha_T(\mathbf{q}^{(m-1)}, \mathbf{q}^{cand})$  then
9:       Set  $\mathbf{q}^{(m)} \leftarrow \mathbf{q}^{cand}$ .
10:    else
11:      Set  $\mathbf{q}^{(m)} \leftarrow \mathbf{q}^{(m-1)}$ .
12:    end if
13:    (iv) Calculate  $y^{(m)} = h(\mathbf{q}^{(m)}, \varphi; \mathbf{z}_T)$ 
14:    Add new sample  $\Upsilon \leftarrow \Upsilon \cup \{y^{(m)}\}$ 
15:  end for
16:  Return  $y^{MH} = \frac{1}{M} \sum_{m=1}^M y^{(m)}$ .
17: end procedure

```

We first show the convergence properties of the MH sampler and then the intuition behind MH sampler in setting inventory target.

2.4.4 Convergence of the Metropolis-Hastings Sampler

Let $K_S(\cdot, \cdot)$ denote the transition kernel of the MH sampler, we now show that the constructed MH sampler converges to the posterior distribution of $W_{t,k}$. Proofs of lemma and theorem are included in Appendix chapter B and chapter C.

Lemma 1. *(i) Proposal probability ϕ_{IC} and ϕ_{MHR} form irreducible and aperiodic Markov Chains on the probability simplex. (ii) For history of finite length T , the Metropolis-Hastings sampler with proposal distribution ϕ_{IC} , ϕ_{MHR} , and acceptance function α_T is π_T -irreducible and aperiodic.*

Remark 1. In order to show that the designed MH sampler converges to the desired full joint distribution, we need to show that its transition kernel is π_T -irreducible. However, the irreducibility and aperiodicity of ϕ_S on the probability simplex is not enough. Roberts and Smith (1994) has shown a counterexample where the transition kernel is π -irreducible over a bigger domain but the corresponding MH sampler is not. Lemma 1 shows that ϕ_{IC} and ϕ_{MHR} are actually both π_T -irreducible and aperiodic, which guarantees its convergence to π_T .

Theorem 1. *For history of finite length T , (i) the MH sampler with proposal distribution ϕ_{IC} , ϕ_{MHR} , and acceptance function α_T has an invariant distribution π_T , and*

$$\frac{1}{M} \sum_{m=1}^M h(\mathbf{q}^{(m)}, \varphi; \mathbf{z}_T) \xrightarrow{a.s.} \int h(\mathbf{q}, \varphi; \mathbf{z}_T) \pi_T(\mathbf{q}) d\mathbf{q} \quad (2.17)$$

(ii) the MH sampler with independence chain is uniformly ergodic on the probability simplex, i.e.

$$\|K_S^m(\mathbf{q}, \cdot) - \pi_T(\cdot)\| \leq r(m) \quad (2.18)$$

where $\|\cdot\|$ is the total variation norm, and $r(m) \rightarrow 0$ as $m \rightarrow \infty$. (see Mengersen et al., 1996, Thm 1.3)

Remark 2. Roberts and Smith (1994) shows that if the transition kernel of a MH sampler is π -irreducible (π , a general target joint distribution) and aperiodic on the domain, then the Markov chain of the MH sampler converges to π . In addition, Mengersen et al. (1996) shows that when proposal distribution is an independence chain, the convergence is uniform for any state, i.e. the distance between a m -step

transition and π are bounded by a decreasing function that is independent of the initial state. Unfortunately, Mengersen et al. (1996) also shows that such uniform ergodicity does not hold for a symmetric proposal distribution. We show the performance of ϕ_{IC} and ϕ_{MHR} through numerical examples.

2.4.5 MLE vs. MH Sampler

Despite that MLE and MH sampler both take advantage of the integer characteristic of the problem, they take different routes towards the problem. The MLE approach assumes that there is a true distribution and its parameters, in this case $\mathbf{q}$, can be estimated by maximizing the log-likelihood function. As a result, MLE reduces some q_j to zero in order to maximize the log-likelihood function. The MH sampler, as a way to simulate the posterior distribution, almost never assigns zero probability to a point in the probability simplex unless the prior is zero. With more data observed, the posterior probability for unlikely events diminishes to zero, as samples in its neighborhood region are accepted with smaller chances. When the prior is uniform, all points in the probability simplex remain positive in the posterior, and are considered when taking numerical convolution in estimating the φ quantile.

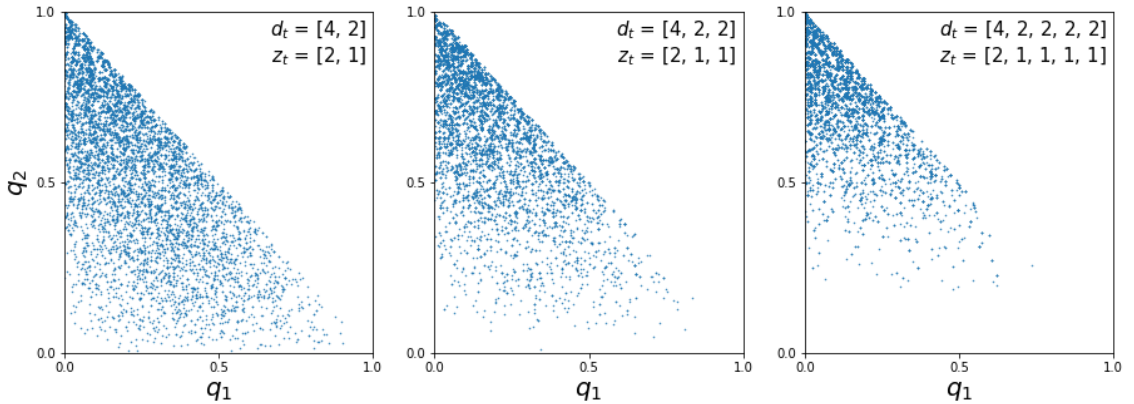


Figure 2.4: MH samples from posterior distribution as more periods with $d_t = 2$ and $z_t = 1$ are observed.

Consider a case where $\mathbf{d}_3 = (4, 2, 2)$ and $\mathbf{z}_T = (2, 1, 1)$, and $BI = \{\underline{w} = 1, \bar{w} = 3\}$. It does not matter how many periods with $d_t = 2, z_t = 1$ are observed, MLE always assigns probability 1 to q_2 since it maximizes the log-likelihood function. Meanwhile, MH sampler keeps $q_3 > 0$ and gradually updates its belief if more periods with $d_t = 2$ and $z_t = 1$ are observed. Figure 2-4 shows the MH samples as observations evolves.

2.5 Numerical Experiment

2.5.1 Study Design

The performance of MLE and MH sampler is demonstrated through a numerical experiment on synthetic data. The *true distributions* for Z_t and $W_{t,k}$ are each from one of five types: Uniform, Increasing, Decreasing, Normal-like, and U-shape. All distributions have left support of 0 and right support of 4, with different probabilities given to each outcome. Table 2.1 shows the probability for each test case. The compound random variables forms 25 different distributions for D_t , with mean range from 1.2 to 8.4, coefficient of variation range from 0.5 to 1.5, skewness range from -0.3 to 1.7, and kurtosis range from 1.6 to 5.9. Table D.1 shows the detail moment information for each test case.

r.v.	Probability				
	0	1	2	3	4
Uniform	0.20	0.20	0.20	0.20	0.20
Increasing	0.05	0.10	0.15	0.26	0.42
Decreasing	0.42	0.26	0.15	0.10	0.05
Normal-like	0.09	0.18	0.45	0.18	0.09
U-shape	0.33	0.13	0.07	0.13	0.33

Table 2.1: Five types of distribution used in generating a synthetic dataset

For each combination of (Arrival Process, Order Distribution), 40 sample paths

are simulated. The length of each sample path is set at 12 periods to resemble a typical three-month or a year-long history (depending whether weekly or monthly observations are used). In the experiment, each algorithm takes in the same sample path as input and calculates inventory target corresponding to φ -quantile. The length of observed history vary among $\{4, 6, 8, 10, 12\}$, representing typical number of observations under a data-scarce environment. Target service level φ is chosen from $\{90\%, 95\%, 98\%, 99\%, 99.5\%\}$ to reflect high service requirements. MLE, MH sampler, and benchmark algorithms receive input from the same sample path and each prescribes a inventory target. We measure their performance by the Optimality Cost Gap (OCG), calculated by $OCG = (C_D(y^{alg}; \varphi) - C_D(y^*; \varphi)) / C_D(y^*; \varphi)$, in which y^* is the optimal newsvendor quantity under the true distribution F_D .

Benchmark Algorithms In comparison, we include five benchmark algorithms: four approaches based on periodic demand history only (two parametric, two non-parametric), and the other one takes full empirical distribution function (e.d.f.) of both Z_t and $W_{t,k}$ assuming that the DM has access to more granular data. The benchmark algorithms are:

- *Poisson*: $y^{Poisson} = \inf\{y; \sum_{k=0}^y \frac{\hat{\lambda}}{k!} e^{-\hat{\lambda}} \geq \varphi\}$, where $\hat{\lambda}$ is the sample mean.
- *Normal*: $y^{Normal} = \hat{\mu} + \Phi^{-1}(\varphi)\hat{\sigma}$, where $\hat{\mu}$ and $\hat{\sigma}$ are estimates from $\mathbf{d}_T$, and Φ is the c.d.f. of a standard normal distribution.
- *SAA*: $y^{SAA} = \hat{F}_D^{-1}(\varphi)$, where $\hat{F}_D$ is the empirical distribution of $\mathbf{d}_T$
- *Max Approach*: $y^{Max} = \max\{d_t\}$
- *Full Empirical Distribution (FED)*: $y^{FED} = \tilde{F}_D^{-1}(\varphi)$, where $\tilde{F}_D$ is the empirical counterpart of Equation (2.4).

y^{SAA} is the well-known approach by Kleywegt et al. (2002) and is guaranteed to converge to y^* in the long run for any given φ . y^{Max} resembles a “naive” reaction when observations are limited but the desired service target is high, under which the DM simply stock up to the highest demand observation in history. y^{Max} is an unbiased estimate of the right support for D and should be fairly close to the true φ quantile when demand is right bounded and φ is close to 1. y^{Normal} and $y^{Poisson}$ are the two most commonly used parametric assumptions in practice. y^{FED} assumes that the DM can observe not only z_t , but also $w_{t,k}$. The DM can obtain both e.d.f. of arrival process and order distributions. y^{FED} is set to at φ quantile of the outcome distribution derived by numerically convolving e.d.f. of Z_t and $W_{t,k}$. The independence of Z_t and $W_{t,k}$ ensures that y^{FED} converges to y^* in the long run for any given φ . We first show the comparison of average OCG between all benchmark algorithms and proposed algorithms with different bounds.

For MLE and MH sampler, the input boundary information are $BI_1 = \{0, 4\}$, $BI_2 = \{0, 6\}$, and $BI_3 = \{0, 8\}$. BI_1 is the tight bound, representing that the DM knows the accurate right support. BI_2 relaxes the right support by 50%, representing the DM’s less accurate knowledge of the right support, and BI_3 relaxes the right support by 100%, representing a lack of customer knowledge. We also include a self-regulating BI_{SR} , where $\underline{w} = 0$ and $\bar{w} = \max_t \left\{ \left\lceil \frac{d_t}{z_t} \right\rceil, \left\lceil \gamma \frac{\sum d_t}{\sum z_t} \right\rceil \right\}$. The first part of the $\bar{w}$ in self-regulating bounds is the consistency requirement, in which $z_t \bar{w} \geq d_t$ must satisfy for all t . The second part is a multiplier of average order size, constraining that even the most extraordinary orders should be within certain bound of the average population. For the sake of a fair comparison, the MH sampler applies a non-informative prior, i.e. assuming all $\mathbf{q}$ are uniformly distributed in the probability simplex. Length of simulation M is chosen from $\{2, 500, 5,000, 10,000\}$.

2.5.2 Results for Numerical Experiment

Average Performance Figure 2-5 shows the mean and standard deviation of OCG between the best of benchmark results based on $\mathbf{d}_T$ only, MLE, and MH sampler with tight bound ($\underline{w} = 0, \bar{w} = 4$). The figure shows the case when target service level is 98%, in which MH sampler with tight bounds outperform all benchmark algorithms by a significant margin for all tested T . Table 2.2 shows that MH sampler even outperforms FED by a non-trivial margin. The OCG gap between MH sampler and FED is about $\sim 36\%$ for $T = 4$ and decreases to $\sim 5\%$ for $T = 12$. Note that the DM actually observes quantity in each order in FED, this shows that the DM can benefit from slight “relax” observations. MLE also outperforms all benchmark algorithms in the numerical tests, with average OCG very close to FED. Moreover, MH sampler with tight bounds shows a more stable performance, achieving the lowest standard deviation OCG. In fact, the performance gap and stability rank between MLE, MH sampler, and benchmark algorithms is consistent for all tested service requirement. Detailed averages and standard deviations of OCG for all test cases are reported in Appendix E.

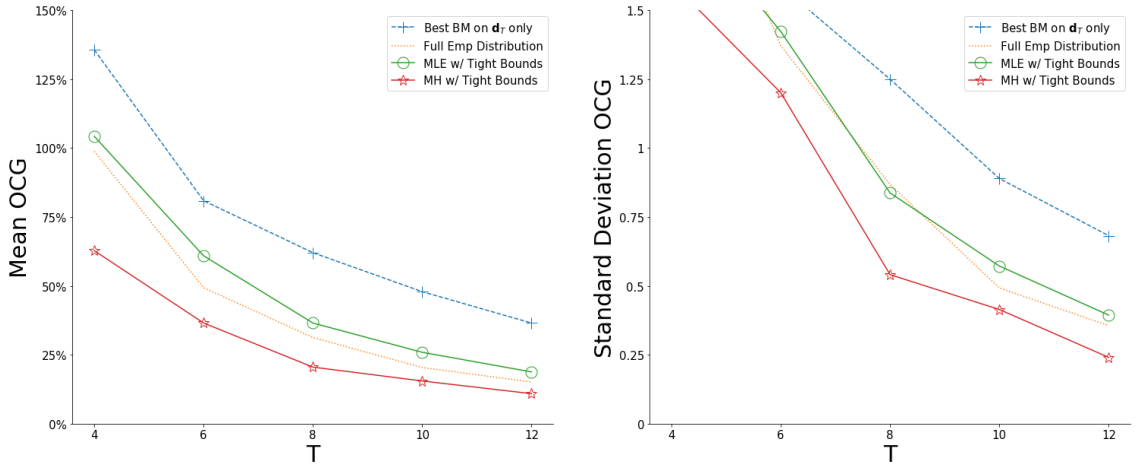


Figure 2-5: Average Optimality Gap between MLE and MH with tight bounds and benchmarks when service requirement is 98%

SL	T	Par. Alg		Non-Par. Alg		MLE	MH		FED
		Poisson	Normal	Max	SAA		ϕ_{IC}	ϕ_{MHR}	
98%	4	164.7%	135.6%	228.6%	228.6%	104.3%	62.7%	62.8%	98.8%
	6	126.6%	80.9%	134.8%	134.9%	60.9%	36.6%	36.7%	49.4%
	8	115.6%	62.1%	95.2%	95.1%	36.5%	20.6%	20.7%	31.3%
	10	105.9%	47.9%	69.5%	69.6%	25.9%	15.5%	15.6%	20.4%
	12	99.1%	36.5%	49.8%	49.9%	18.8%	10.9%	11.3%	15.1%

Table 2.2: Average Optimality Gap between MLE and MH with tight bounds and benchmarks when service requirement is 98%.

When service requirement is high, over-estimation is preferred to under-estimation since one lost sale costs much more than one unit of exceeded inventory. Nevertheless, it is hard for any algorithm to detect a small chance event with such limited data. For instance, to observe a demand realization within the top 5% quantile, it takes 20 periods on average. The probability of observing such d_t within a 6 period $\mathbf{d}_T$ is about 30%, and drops to 6% if the service requirement rises to 99%. The pure data-driven algorithm SAA is limited by actual history and never prescribes inventory targets higher than the maximum among past observations. Parametric algorithms, despite that they may prescribe inventory targets higher than observations, rely on the similarity between assumed demand and the true distribution. Since most test cases have a heavier tail than Poisson, under-estimation is more likely to happen. Table 2.3 shows the percentage of under-estimation, optimality, and over-estimation for all benchmark algorithms when target service level is 98%. Table for other tested service requirements is included in Appendix E Table E.3. When more data is observed, all benchmark algorithms except for Poisson show decrease in under-estimation, increase in over-estimation, as well as increase in the number of cases where $y^{alg} = y^*$.

On the other hand, MLE and MH sampler demonstrate a very different behavior. Even for high service level as 98%, both MLE and MH sampler present higher

SL	T	Par. Alg		Non-Par. Alg		Full
		Poisson	Normal	MaxApp	SAA	Emp. Dist.
98%	4	80.3/8.7/11.0	66.5/10.3/23.2	89.4/6.7/3.9	89.4/6.7/3.9	58.1/15.6/26.3
	6	82.0/8.4/9.6	61.8/13.4/24.8	82.7/10.2/7.1	82.7/10.3/7.0	49.7/19.1/31.2
	8	84.7/8.6/6.7	61.0/12.2/26.8	77.5/12.5/10.0	77.8/12.3/9.9	46.3/22.3/31.4
	10	85.3/9.7/5.0	59.7/14.8/25.5	73.2/14.6/12.2	74.5/14.3/11.2	44.5/25.3/30.2
	12	87.1/7.9/5.0	57.9/17.1/25.0	68.7/17.0/14.3	70.4/17/12.6	41.2/29.6/29.2

Table 2.3: Percentage of Underestimation/Optimality/Overestimation of each benchmark algorithm for the 98% quantile.

chances of obtaining the newsvendor quantity of the true distribution. The major performance difference between MLE and MH is that, although both algorithms exploit the combination structure beneath $\mathbf{d}_T$ and $\mathbf{z}_T$, MLE tends to allocate weights to certain q_j with low j . Since larger orders are less prevalent in $\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI}$, MLE tends to “ignore” δ_t by driving corresponding q_j for higher j to zero. In contrast, the MH sampler keeps all possibilities unless that there is strong evidence from data indicating that large orders are unlikely. Thus, MH has a stronger tendency to over-estimate, which results in more conservative inventory targets. Table 2.4 shows the results for MLE and MH sampler with tight bounds and $M = 10,000$. The proposal distribution used in the MH sampler does not influence the outcome since they both generate samples from the same posterior distribution.

$BI = \{\underline{w} = 0, \overline{w} = 4\}$				
SL	T	MLE	MH-IC	MH-MHR
0.98	4	62.6/11.3/24.8	46.5/14.2/38	46.2/14.4/38.1
	6	54.4/16.8/28.6	38.9/18.3/42.6	39.1/18.3/42.4
	8	51.6/19.7/28.5	34.5/21.5/43.8	34.9/21.1/43.8
	10	49.8/22.0/28.2	32.8/23.2/44.0	32.4/23.2/44.4
	12	46.8/24.1/29.1	28.6/26.3/45.1	29.7/24.7/45.6

Table 2.4: Percentage of Underestimation/Optimality/Overestimation of MLE, MH-IC, and MH-MHR with tight bounds for the 98% quantile.

Sensitivity of Boundary Information The tightness of bounds measures the DM's knowledge on order distribution. MLE is insensitive to the boundary information since it searches for the local maximum. Consider the case where the DM may hold two different boundary information, $BI = \{\underline{w}, \bar{w}\}$ and $BI' = \{\underline{w}, \bar{w} + 1\}$. As long as both BI are consistent with data, the MLE solution for $\mathbf{q}$ is less likely to change unless the extra δ added to $\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI'}$ significantly alters the expression of Equation (2.10). The tightness of bound has a greater impact on MH sampler, which requires more data to obtain a tight posterior distribution. The self-regulating bounds, on the other hand, efficiently fix the impact of loose bounds as long as γ is chosen reasonably. Figure 2-6 shows the average optimality cost gap of MLE and MH sampler under bounds with different tightness. For service requirement of 98%, the MLE solutions with different bounds are within 1.4% in terms of average OCG between the best and worst performance. Most MH samplers outperform best benchmark algorithms for 98% service requirement, except for the one with 200% bounds. The lower the service requirement, the smaller gap between MH sampler and benchmark algorithms. MH sampler with self-regulating bounds outperforms FED in most cases when service requirement is greater than or equal to 95%. Detailed numerical results are shown in Appendix E.

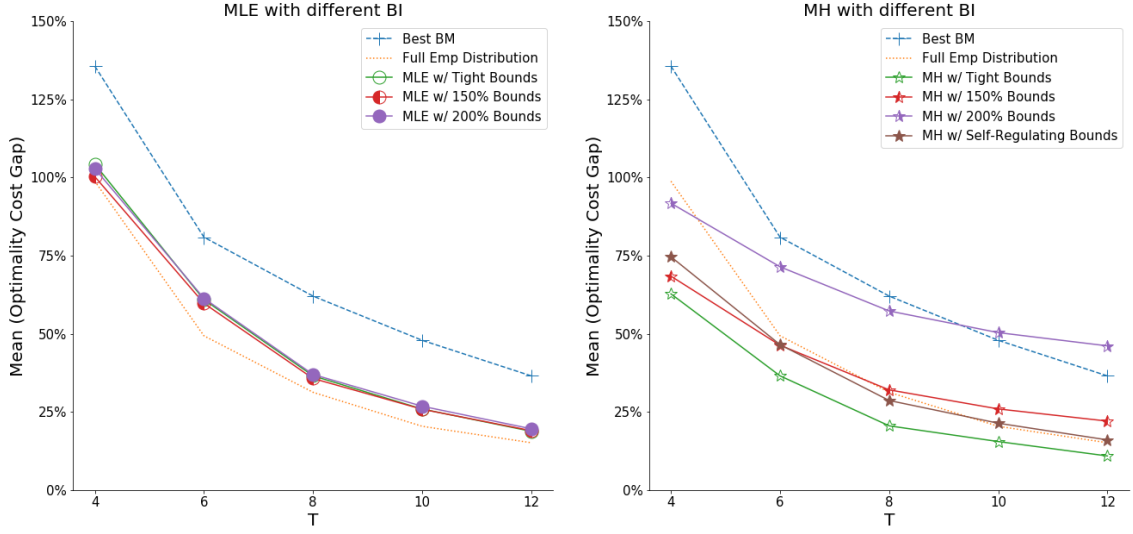


Figure 2.6: Sensitivity of Bounds for MLE and MH when target service level is 98%, $\gamma = 2$ for MH-IC sampler with self-regulating upper bounds.

Acceptance Rate of and Coverage Rate of MH Sampler Despite the fact that independence chain and modified hit-and-run simulate samples from the same posterior distribution, both have trade-offs in computation. In each iteration, ϕ_{IC} generates a candidate $\mathbf{q}$ from a Dirichlet distribution, which is independent from the previous sample. It explores the probability simplex more evenly at the cost of recalculating Equation (2.10) entirely in every iteration. When actual demand is higher, such computational burden can grow exponentially. ϕ_{MHR} , in contrast, only alters values p_i and p_j in each iteration, and the conditional probability $f(\delta_t^*)$ remain unchanged if δ_t^* does not involve $w_{t,k} = i$ nor $w_{t,k} = j$. So computation of $P\{\Theta_{\mathbf{d}_T, \mathbf{z}_T, BI} | \mathbf{q}^{cand}, \mathbf{z}_T\}$ can be reduced by keeping some of $f(\delta_t^*)$ unchanged. Samples generated by MH-MHR have higher correlations and tend to have higher acceptance rates. But the sampler should also run for more iterations for a more stationary distribution. The acceptance rate also decreases for both proposal distributions when more data is available. Table 2.5 shows the acceptance rate of MH-IC and MH-MHR in an example test case. Figure 2.7 shows the distribution of y as longer simulations are run. For most test

cases generated in the numerical experiment, MH-IC converges faster than MH-MHR. 5,000 runs often result in a satisfactory stationarity of simulated y .

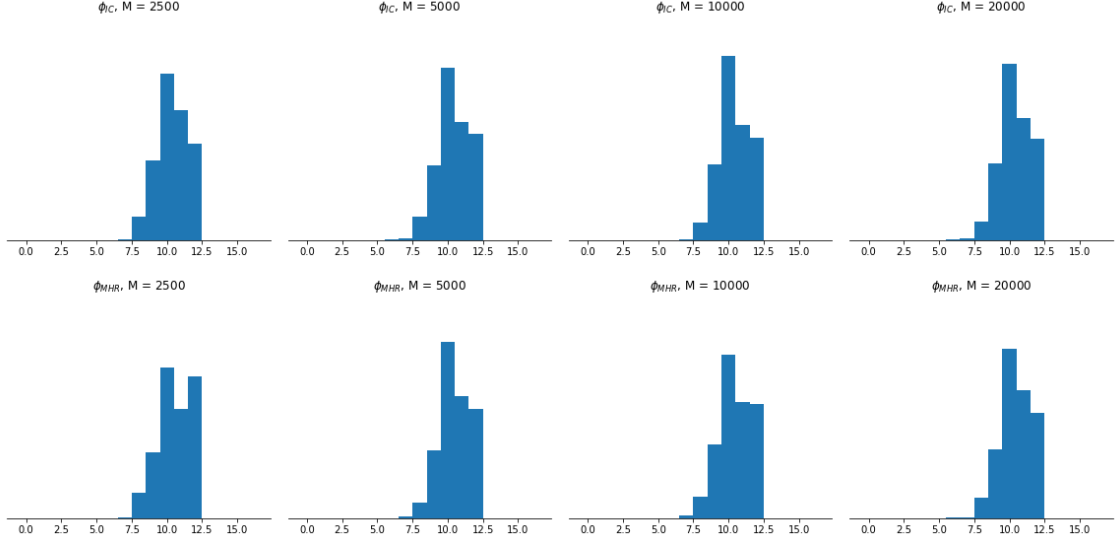


Figure 2.7: Distribution of y^{MH-IC} and y^{MH-MHR} under different number of iterations.

$\bar{w}$	4		6		8	
T	MH-IC	MH-MHR	MH-IC	MH-MHR	MH-IC	MH-MHR
4	0.36	0.56	0.36	0.62	0.31	0.62
6	0.28	0.45	0.26	0.50	0.20	0.50
8	0.23	0.36	0.19	0.41	0.14	0.41
10	0.19	0.28	0.15	0.32	0.10	0.33
12	0.16	0.21	0.12	0.25	0.07	0.25

Table 2.5: Acceptance of an example test case, showing that MH-MHR has a higher acceptance rate than MH-IC.

2.6 Conclusion

Firms often have to make inventory decisions under limited data. Although plenty of recent work has focused on the limiting behavior of non-parametric models, the

performance is often not satisfactory when demand history is short and desired service level is high.

In this paper, we model demand as a compound random variable instead of a single random variable in order to extract more information from the limited demand history. The data on customer arrivals enables us to generate all feasible combinations within prior bound information. Based on this, we propose two approaches, MLE and Metropolis-Hastings sampler, that take advantage of such information and set inventory targets. Compared parametric and non-parametric approaches that based on periodic demand only, MLE and Metropolis-Hastings sampler have superior performance regarding the closeness to the true optimal. Metropolis-Hastings sampler even achieves lower optimality cost gap when compared against the case where the decision maker observes order details and uses it to construct empirical distributions. The numerical investigation also shows that the Metropolis-Hastings sampler considered is more sensitive to the input boundary information, which can be improved by applying a self-regulating bound.

To sum up, setting inventory targets is challenging when available data is limited, demand is low and desired service target is high. In those cases, slight over-estimation usually result in lower cost. However, the magnitude of overestimation should be controlled. The two proposed approaches can improve inventory performance by including observations on customer arrivals, which can be collected at a fairly low cost. We hope to apply those approaches in other settings in future work.

Chapter 3

Set Inventory Targets for Low-Demand SKUs with Integer Pattern Sampling Heuristic

3.1 Problem Description

In this study, we collaborate with a major biotechnology company, Company M, which owns more than one hundred brands and manages a global supply chain network. We study the operation at one of its regional distribution centers, the W facility, for delivering reagents products (one of Company M’s major product families), and help set inventory targets for low-demand Stock Keeping Units (SKUs), based on very limited data. All names and numbers are disguised for company confidentiality

Company M’s supply chain for delivering reagents can be modeled as a three-echelon network: manufacturing sites receive raw materials and convert them into finished goods (FG), which are labelled as “bulk SKUs” and stored at central warehouses. There are two types of bulk SKUs: regular bulk SKUs and specialty bulk SKUs. Regular bulk SKUs are stored in larger containers, e.g. 60-gallon barrels. These regular bulk SKUs are then shipped to regional distribution centers where they are repacked to become catalog SKU. Catalog SKUs are the finished good products in sizes, e.g. 5 ml containers, for fulfilling customer orders. Such repacking

at downstream locations is common in practice for the purpose of reducing shipping cost. Special bulk SKUs are finished goods that are already in the size of catalog SKUs, and can fulfill customer orders directly. At the W facility, the only processing activities consist of receiving, repacking and outbound shipping.

Among all SKUs stocked at the W facility, approximately 75% of the saleable SKUs are repacked from regular bulk SKUs, with the remained being special bulk SKUs. The W facility attempts to manage all catalog SKUs as make to stock with sufficient inventory to immediately satisfy customer orders from inventory. If there is not sufficient on-hand inventory for a catalog SKU order, it triggers a work order to repack corresponding regular bulk SKU into the required sizes. The actual repacking task requires a worker to locate the regular bulk SKUs on the rack and manually repack the catalog SKU at a work station. Depending on the actual category of SKU (hazardous vs. non-hazardous), scheduling, and work station availability, a repack order often requires a lead time between one to two days. When the required bulk SKU is not available, the inventory manager receives a backlog notification indicating that a procurement order should be placed on upstream suppliers, the lead time of which varies by source. Figure 3.1 shows the repacking process in practice. At the W facility, the prefix “bulk” is a relative term and does not necessarily indicate a large container. A regular bulk SKU can be stored in a bottle no larger than 300 ml if the corresponding catalog SKUs are packed in 5 or 10 ml containers. It may also be stored in 60-gallon barrels if catalog sizes are measured in quarts or gallons. In general, most SKUs at the W facility do not occupy a large space, yet the facility does hold a large number of distinct SKUs.

The W facility strives to reach a guaranteed same day shipping promise, with

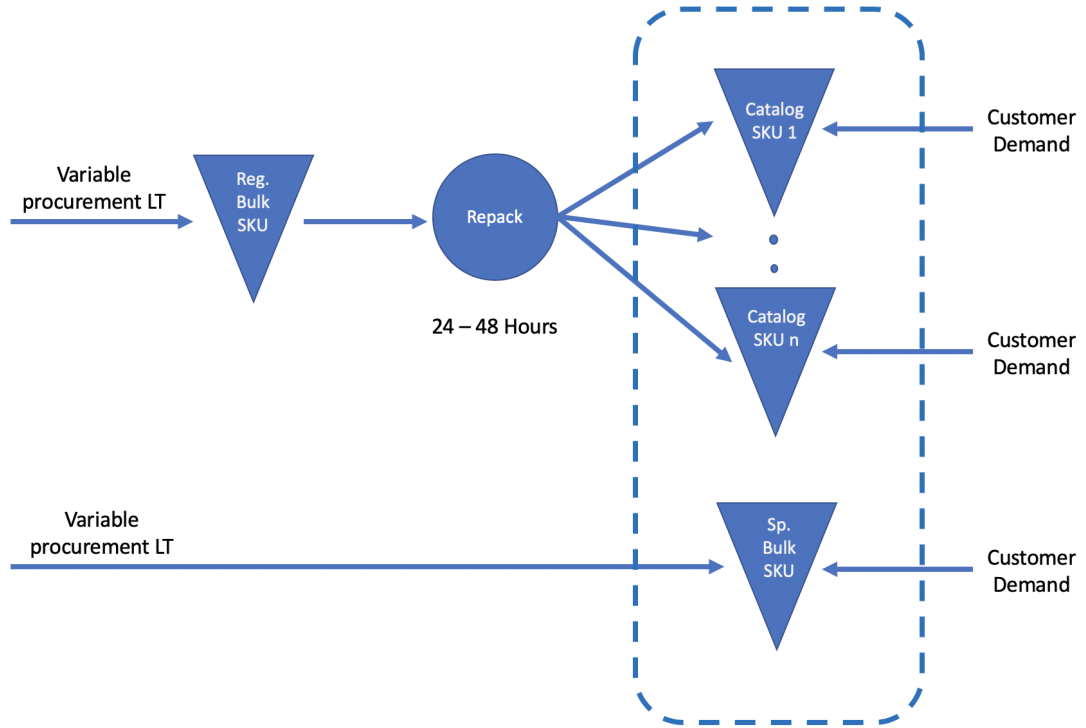


Figure 3·1: Repacking process at Company M's W facility. Repacking lead time for varies between 24-48 hours. Procurement lead times are longer depending on specific sources.

service level measured by the portion of customer orders shipped on the same day as orders arrive. Taking into consideration that lead time for repacking can take up to 48 hours, the service requirement is unlikely to be met when there is not enough catalog SKU inventory. Moreover, when the W facility does not hold enough bulk SKU inventory, the fulfillment time can be drastically extended due to the wait for replenishment. The target service level was above 95% at the W facility, and room for improvement exists.

On a regular work day, the inventory manager at the W facility goes through the backlog list of bulk and catalog SKUs, and manually places both repacking and procurement orders. Given bandwidth of the inventory manager, she can process

several hundred orders on a single day. During busy seasons, this seldom clears backlog list entirely. Prior to this project, the quantity of work and procurement orders were based on a mixture of inventory manager’s intuition, subjective demand forecasts, and practical constraints like minimum order quantity and shipping sizes for bulk SKUs. Inventory manager lacked any formal guidelines regarding how many units to repack or procure proactively to keep inventory at a reasonable level to meet future demand.

The existing repacking process also makes it difficult to efficiently schedule. During busy periods, there are not enough labor nor work stations to repack and ship customer orders; yet those resources are under-utilized when demand is low. Such imbalance adds to the waiting time during high-demand periods, and squanders hours available at work stations during low-demand periods. The inventory manager believes that, under an ideal inventory policy, customer orders are fulfilled from on-hand inventory, and inventory levels are maintained in a proactive manner. For most of special bulk and catalog SKUs, the procurement and repacking should happen at the monthly level, while fast-moving SKUs should be managed more frequently. However, since there was no systematic rules for setting inventory levels at the W facility, such proactive management was not possible at the W facility.

3.1.1 Challenges at the W Facility

Despite the fact that Company M has widely implemented a Reorder Point (ROP) model at its facilities by fitting a normal distribution using a SKU’s trailing 52 weeks of demand data, there are two major challenges at the W facility that prevent a direct usage of the *Normal Approach*. First, the data collected at the W facility was very limited: they consist of the previous six months of sales at the monthly

level, and the total number of orders for that SKU over the same horizon. The first piece of data has 6 observations for each saleable SKU, while the second piece has only 1 observation. This creates the challenge that, with such limited data, no distribution can fit with confidence, and the *Normal Approach* only works well when the true underlying distribution has a shape similar to a normal curve. Second, most catalog SKUs are slow-moving products, which also makes it hard to fit a continuous distribution.

The goal of this study is to set reference inventory targets, for saleable SKUs and regular bulk SKUs, so that whenever the on-hand inventory level of a SKU is less than its targets, a reminder can be sent to the inventory manager so that work and procurement orders can be placed in a proactive manner when idle work station and labor is available. With better knowledge of how many units should be held in stock, the inventory manager can smooth labor scheduling, increase work stations utilization, and achieve a higher service level.

The limited data availability and high service requirement create several challenges regarding whether a parametric or a non-parametric model should be used. Parametric models for discrete time intervals assume that periodic demand follows a specific type of distribution, e.g. normal, Poisson, or belongs to some family of distribution, e.g. exponential, or lognormal. Parameters in those models are either estimated statistically or in a Bayesian fashion (Scarf, 1957, 1959; Azoury, 1985; Akcay et al., 2011). When a demand process consists of both order arrivals and order quantity, it can be modeled as a compound random process. Feeney and Sherbrooke (1966) consider the steady state probability for the number of units in resupply under a compound Poisson process; Sherbrooke (1968) considers a base-depot system, also

using a compound Poisson process, with its mean estimated by a Bayesian procedure. Nevertheless, at the W facility, the available data is not enough to support the compound Poisson assumption, which prevents a direct implementation of those results. Even if the assumption holds, the inventory manager still lacks knowledge of the order arrival and order quantity distribution to estimate required parameters.

Meanwhile, there is increasing interest in models that only require partial knowledge of demand or only estimate an empirical distribution from observed data. For instance, Bookbinder and Lordahl (1989) apply the bootstrap approach to estimate quantiles of lead time demand and variance of estimation, and show that such approach is preferred when the true underlying distribution is not standard. Godfrey and Powell (2001) develop a technique called Concave, Adaptive Value Estimation (CAVE) that approximates the value function by piece-wise functions. Kleywegt et al. (2002) proposed the framework of Sample Average Approximation, which solves the sample path counterpart of a stochastic optimization problem. Levi et al. (2007, 2015) show the bounds on the number of required samples for an arbitrarily close solution to the true optimalities in a repeated newsvendor setting. For the decision problem at the W facility, two practical reasons prevent us from directly implementing one of the existing non-parametric models. First, they usually require a large amount of data when the desired service level is high. The desired service level in our problem is above 95% while the available dataset is very small. The other reason is that, to the author’s best knowledge, almost all non-parametric models assume demand to be a single random variable. This does not raise any theoretical issue since it can depict the outcome of any demand distribution as long as there is enough data. But it also neglects potential information from order arrivals since it is not part of the distribution. In this case, it means that one of the seven

data points is discarded, and the information contained is lost. In this paper, we model demand as a compound process and propose a new solution called *Integer Pattern Sampling (IPS)* that sets inventory targets for the slow-moving SKUs at the W facility, with no parametric assumption on either the order arrivals or the order quantity.

3.2 Integer Pattern Sampling (IPS) Solution for Setting Inventory Targets

The problem at the W facility can be abstracted as follows: for a single SKU, the inventory manager only observes T periods of realized demand $\mathbf{d}_T = (d_1, \dots, d_T)$, as well as the total number of orders $\tilde{z}$ over the same time interval. In each period t , z_t orders arrive with the k -th order requesting $w_{t,k}$ units of the product. However, neither z_t nor $w_{t,k}$ are known to the inventory manager. The goal is to come up with an inventory target based only on available information that covers a high φ ($\varphi \geq 95\%$) quantile of demand uncertainty (Type I service level). We also assume that the inventory targets are not capacitated since space is not a binding constraint at the W facility.

Despite the seemingly limited amount of data available, it is worth noting that each order consists of strictly positive integer units, which reveals possible order arrivals and quantities. Consider 4 period dataset where $\mathbf{d}_4 = (0, 1, 2, 3)$, $\tilde{z} = 4$, and there is no limit on how many units can there be in a single order. Constrained by the fact that each recorded order has to contain at least 1 unit, it can be derived immediately that there is only 1 order in period 2, since only one unit is requested. For the remaining two periods, there are two possibilities: if period 3 only has one

order, then the remaining two orders are in period 4, with order quantity 1 and 2, respectively. If there are two orders in period 3, then each of them must only contains 1 unit, and period 4 has one order of 3 units. This example shows that, combining periodic demand and order arrival information allows us to infer all possibilities from the limited demand data.

In this section, we develop a solution named *Integer Pattern Sampling (IPS)* to take advantage of the discrete nature of the demand process and sets inventory targets for high service requirements under limited observations. To denote those integer patterns, we introduce the notation $\theta = (\mathbf{z}^\theta, \mathbf{w}^\theta)$ to represent a feasible combination of order arrivals and quantities, where $\mathbf{z}^\theta = (z_1^\theta, \dots, z_T^\theta)$ denotes number of orders in each period, and $\mathbf{w}^\theta = (w_{1,1}^\theta, \dots, w_{T,z_T}^\theta)$ denotes all order quantities. In the example above, we have three different θ :

θ	$\mathbf{z}^\theta$	$\mathbf{w}^\theta$
θ^1	(0, 1, 2, 1)	(1, 1, 1, 3)
θ^2	(0, 1, 1, 2)	(1, 2, 1, 2)
θ^3	(0, 1, 1, 2)	(1, 2, 2, 1)

Let $\Theta = \{\theta | \theta = (\mathbf{z}^\theta, \mathbf{w}^\theta), \mathbf{1}^T \mathbf{z}^\theta = \tilde{z}, \sum_{k=1}^{z_t} w_{t,k}^\theta = d_t, \forall t = 1, 2, \dots, T, z_t^\theta \in \mathbb{Z}^*, w_{t,k}^\theta \in \mathbb{Z}^+\}$ be the set that contains all θ given $\mathbf{d}_T$ and $\tilde{z}$. Although θ^2 and θ^3 only differ in the the ordering of $w_{t,k}$, we find it useful to differentiate them.

Boundary Information Although granular information such as order quantities are not observed in this formulation, the inventory manager can still impose prior knowledge as a filter on Θ . Examples of such prior knowledge might be prior distribution, mean, boundary information, or constraint on tail moments (Perakis and Roels, 2008; Levi et al., 2015; Saghafian and Tomlin, 2016). In this study, we

consider the most practical prior knowledge—boundary information (BI), that consists of lower bounds and upper bounds on z_t and $w_{t,k}$. If the inventory manager knows, even by experience, that $\underline{z} \leq z_t \leq \bar{z}$ and $\underline{w} \leq w_{t,k} \leq \bar{w}$ almost surely happens, then we can prune Θ by removing those θ that contains any z_t or $w_{t,k}$ outside the interval. The prior knowledge must be consistent with observations, i.e. for any valid bounds, $\underline{z}T \leq \tilde{z} \leq \bar{z}T$ and $\underline{z}\underline{w} \leq d_t \leq \bar{z}\bar{w}, \forall 1 \leq t \leq T$. Let $BI_0 = \{\underline{z} = 0, \bar{z} = \infty, \underline{w} = 0, \bar{w} = \infty\}$ be the non-informative BI , and Θ_{BI} be the set of combination possibilities under any given BI . In the following sections, we describe how Θ_{BI} can be sampled uniformly.

3.2.1 Uniform Sampler on Θ_{BI_0}

Without any limitation on z_t and $w_{t,k}$, we can generate all θ for any given $\mathbf{d}$ and $\tilde{z}$ as long as both $\sum_{t=1}^T z_t = \tilde{z}$ and $\sum_{k=1}^{z_t} w_{t,k} = d_t$ are satisfied. This can be done by a “wall-insertion method.” Imagine a sequence of $\sum_{t=1}^T d_t$ dots, where each dot represents one unit of demand. Between any adjacent dots, there is one possible location for a wall, and all the dots between two walls denote units requested by the same order. Meanwhile, since no order splits across two or more periods, each period also creates a wall, taking the position between the $\sum_{t=1}^s d_t$ -th unit and the $\sum_{t=1}^s d_t + 1$ -th unit. The first position prior to the first unit is also taken since no order happens before the epoch, and the last position is taken by the last period with positive demand. Thus, the generation of θ is equivalent to picking the locations to insert walls. Figure 3.2 show the example of generating θ for the example in section 3.2.

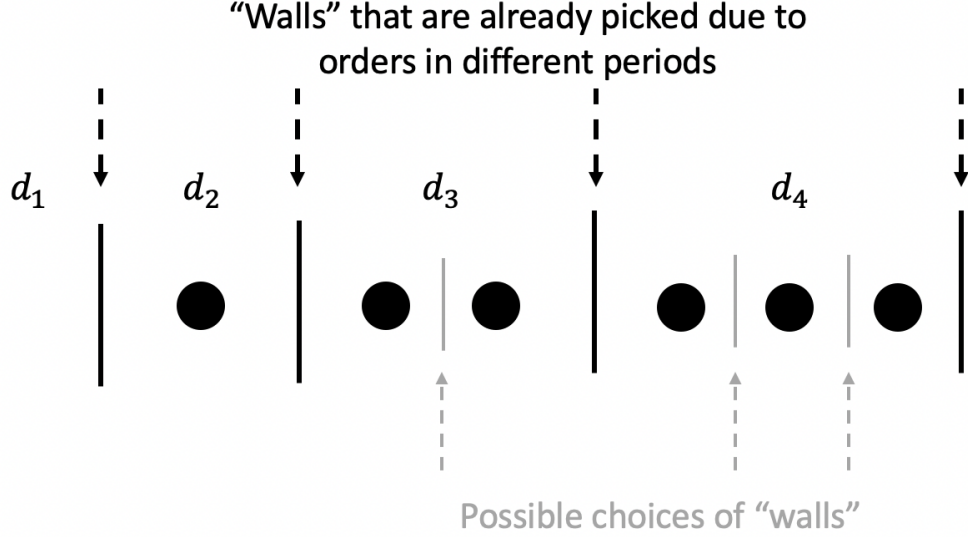


Figure 3.2: Illustration of generating Θ_{BI_0} for $\mathbf{d} = (0, 1, 2, 3)$ and $\tilde{z} = 4$, with wall insertion method.

There are $\sum_{t=1}^T d_t + 1$ locations for walls in total, including ones on both ends. Let $\eta = \sum_{t=1}^T \mathbb{I}_{[d_t \geq 1]}$ denote the number of periods in the observation with non-zero demand, the total number of feasible locations is $\sum_{t=1}^T d_t - \eta$. Since each period with non-zero d_t generates at least one order by default, the number of walls to insert is $\tilde{z} - \eta$. Each θ can be mapped to a unique selection of walls and vice versa. Therefore, we have $|\Theta_{BI_0}| = \binom{\sum_{t=1}^T d_t - \eta}{\tilde{z} - \eta}$. For small Θ_{BI_0} , uniformly sampling Θ_{BI_0} in expectation is equivalent to generating all θ and using each sample once; for larger Θ_{BI_0} , uniformly sampling from Θ_{BI_0} is equivalent to randomly picking $\tilde{z} - \eta$ identical items from a pool of $\sum_{t=1}^T d_t - \eta$ with no replacement, which can be achieved by random number generation.

3.2.2 Uniform Sampler on Θ_{BI}

To sample Θ_{BI} under a general BI , an intuitive approach is to first sample Θ_{BI_0} and then prune the set by removing θ that are not consistent with BI , or only

accept-reject random samples until the number of random samples reaches a preset threshold. Nevertheless, it can be computationally wasteful when a significant portion of random draws falls outside the bounds established by BI . Consider the case where all $\mathbf{d}_8 = (8, 8, 8, 8, 8, 8, 8, 8)$, $\tilde{z} = 12$, and $\bar{w} = 4$. With the wall insertion method under non-informative boundary information, there are $\binom{58}{6} = 5,245,786$ combinations in Θ_{BI_0} . However, there is only one feasible combination given the boundary information, which consists of two orders per period and 4 units in each order. The example is constructed on purpose to demonstrate an extreme case that the acceptance rate can be very close to zero under certain cases while the accept-reject approach can take almost infinite time to generate a desired number of random samples. To resolve such computational burden, we construct an algorithm that samples uniformly from Θ with general BI by applying a hierarchical sampling.

The algorithm first generates all feasible combinations for a given period and a fixed z_t in order to decide the probability of sampling $\mathbf{z}^\theta$. Then it iterates through all periods and samples a combination for each.

First, consider a single period in which d_t , and BI are known, and z_t is fixed. Let $\delta_t^* = (w_{t,z}^*, \dots, w_{t,z_t}^*)$ denote a feasible order quantity combination for period t , such that $w_{t,1}^* \leq \dots \leq w_{t,z_t}^*$, $\sum_{k=1}^{z_t^*} w_{t,k}^* = d_t^*$, and $\underline{w} \leq w_{t,k} \leq \bar{w}$. The superscript in δ_t^* represents that there can be multiple possible combinations given d_t , z_t , and BI . Note that δ_t^* is an ordered vector. Let $\tau(\delta_t^*)$ be the number of unordered permutations of δ_t^* . Let $\Delta_t(d_t, z_t, BI)$ denote the set consisting all possible δ_t^* and their unordered permutations, then $|\Delta_t(d_t, z_t, BI)| = \sum_{\delta_t^* \in \Delta_t(d_t, z_t, BI)} \tau(\delta_t^*)$. Generation of δ_t^* and $\tau(\delta_t^*)$ are described in Appendix F and G, respectively. In order to simulate θ uniformly, the algorithm first generates all possible $\mathbf{z} = (z_1, \dots, z_T)$, in

which each z_t should satisfy $\max \left\{ \underline{z}, \left\lceil \frac{d_t}{\underline{w}} \right\rceil \right\} \leq z_t \leq \min \left\{ \overline{z}, \left\lfloor \frac{d_t}{\underline{w}} \right\rfloor \right\}$. The generation of all $\mathbf{z}$ can be achieved by a simple branch and bound enumeration, which is described in Appendix H. The uniform sampling of Θ_{BI} can be done by the following two steps:

Step 1 Sample $\mathbf{z} = (z_t, \dots, z_T)$ with probability $P(\mathbf{z}) = \frac{\prod_{t=1}^T |\Delta_t(d_t, z_t, BI)|}{\sum_{\mathbf{z}} (\prod_{t=1}^T |\Delta_t(d_t, z_t, BI)|)}$.

Step 2 Initiate θ as an empty vector. For each z_t in the sampled $\mathbf{z}$, sample an ordered δ_t^* with probability $P(\delta_t^*) = \frac{\tau(\delta_t^*)}{|\Delta_t(d_t, z_t, BI)|}$. Add δ_t^* to θ .

By iterating the two steps, the algorithm guarantees that each θ are sampled uniformly. Note that each sampled δ_t^* is ordered, with $w_{t,1} \leq \dots \leq w_{t,z_t}$. In the next section, we show that this is equivalent to uniformly sampling unordered combinations from Θ_{BI} when estimating inventory targets. During implementation, the tuple $(d_t, z_t, \underline{w}, \overline{w}, \delta_t^*, \tau(\delta_t^*))$ is independent of observations and can be stored in a Hash table to reduce computation time.

3.2.3 Estimating Inventory Targets with IPS Heuristic

For any given $\mathbf{d}_T$ and $\tilde{z}$, the inventory manager can set a computational budget K on $|\Theta_{BI}|$. For a general BI , $|\Theta_{BI}|$ can be calculated using the algorithm in Section 3.2.2. If $|\Theta_{BI}| \leq K$, all θ can be generated iteratively. Otherwise, the inventory manager can specify the desired number of samples M , and random draws from Θ_{BI} . For each θ , the e.d.f. of order arrivals is calculated using sampled $\mathbf{z}_T$ and the e.d.f. of order quantity is calculated using sampled $w_{t,1}, \dots, w_{T,z_T}$, both of which are obtained from θ . Then, a numerical convolution is performed to obtain $\hat{F}_D$ for the compound process and the φ quantile is also computed numerically. For each θ , there is an

estimate $\hat{y}^\theta = \hat{F}_D^{-1}(\varphi; \theta)$, and their average is the solution output. Figure 3-3 shows the flow diagram of estimating the φ quantile of D_t using the IPS solution.

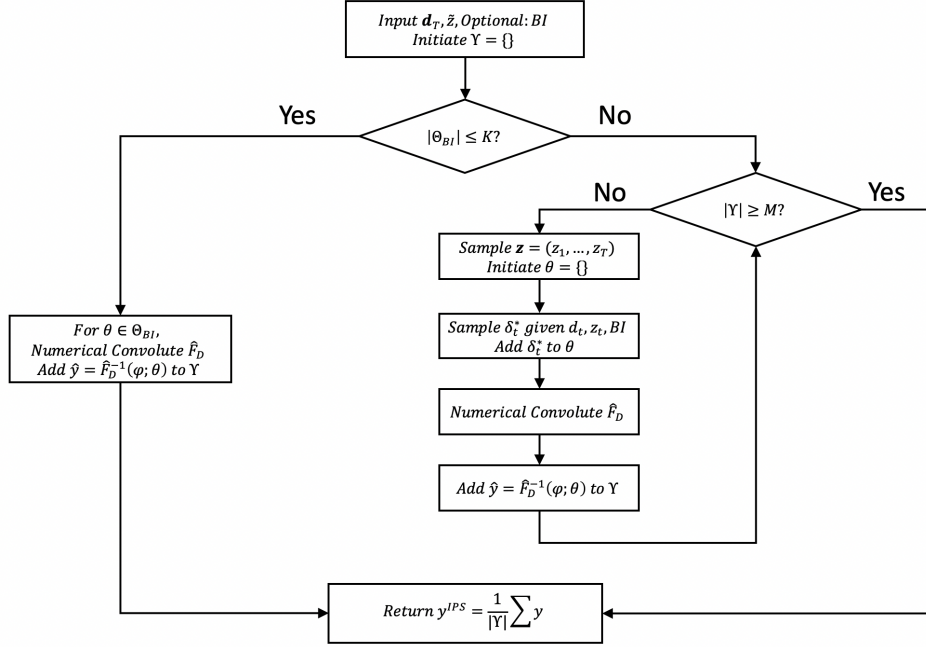


Figure 3-3: Flow diagram of how φ -quantile is estimated by IPS solution.

3.3 Numerical Experiment

We first show the performance of the IPS solution in section 3.2.3 through a synthetic dataset. We randomly generate 1,000 test cases that mimic the demand stream and observations at the W facility. Arrival process Z_t and order quantity $W_{t,k}$ are two streams of independent random variables. Each Z_t are independent and identically distributed and independent of $W_{t,k}$. Each $W_{t,k}$ is also independent of each other. The true distribution of Z_t has support 0 and 4, and true distribution of $W_{t,k}$ has support 1 and 4. In each cases, the true probability mass function of Z_t and $W_{t,k}$ are random draws from a Dirichlet distribution with parameter $\mathbf{1}$ of appropriate dimensions. For the true distributions, mean ranges from 0.5 to 11.7, coefficient of variation ranges

from 0.3 to 3.5, skewness ranges from -2.0 to 3.9, and kurtosis ranges from 1.2 to 18.1. Figure 3.4 shows the distribution of the statistics of the true distributions in the synthetic case.

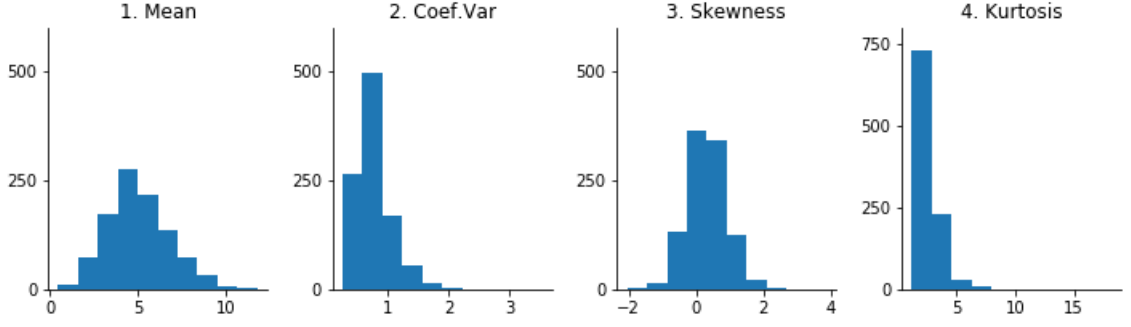


Figure 3.4: Distribution of mean, coefficient of variation, skewness, and kurtosis of the true distribution of 1,000 synthetic test cases

To emulate the inventory manager's task at the W facility, every test case generates a sample path of 6 periods, by first randomly picking z_t and then z_t random draws of $w_{t,k}$. The decision maker observes $\mathbf{d}_6$ and $\tilde{z}$ and sets a inventory target y . To reflect the asymmetric cost structure of overstocking and understocking, we calculate the expected newsvendor cost of a y from each solution under the true distribution. Unit overage cost is normalized to 1, and unit underage cost is $\varphi/(1 - \varphi)$. Target service requirement φ are chosen from $\{0.9, 0.95, 0.98, 0.99\}$ to reflect the W facility's high service promise.

We consider two benchmark solutions. *Normal Approach* assumes that D_t is normally distributed and calculates $y^{Normal} = \hat{\mu} + \beta_\varphi \hat{\sigma}$, where $\hat{\mu}$ and $\hat{\sigma}$ are the sample mean and sample standard deviation from $\mathbf{d}_T$ while β_φ is the factor corresponding to φ quantile. The normal approach represents the widely used parametric paradigm as well as the logic that Company M applies at other facilities. *Max Approach* sets $y^{Max} = \max\{d_t\}$. It is a truthful codification of the inventory manager's intuition,

which applies the idea of one-month-forward-coverage, and uses the maximum of historical observations as an estimate for the upper bound of demand. The max approach also coincides with Sample Average Approximation from Kleywegt et al. (2002) for most of the time since service requirements are high and all targets are rounded to the closest integer. For IPS solution, we consider three types of boundary information:

- BI_0 : $\underline{z} = 0, \bar{z} = \infty, \underline{w} = 0, \bar{w} = \infty$
- BI_{SR} : Self regulating bounds $\underline{z} = 0, \bar{z} = \lceil \gamma \frac{\bar{z}}{T} \rceil, \underline{w} = 1, \bar{w} = \lceil \gamma \frac{\sum dt}{\bar{z}} \rceil$
- BI_{Exact} : $\underline{z} = 0, \bar{z} = 4, \underline{w} = 1, \bar{w} = 4$

where BI_0 is the non-informative bounds; BI_{SR} is a self-regulating bound that ties $\bar{z}$ and $\bar{w}$ to their average values from observations. γ is the relaxing factor indicating how far the upperbound is away from its mean; BI_{Exact} is the exact bounds assuming that the inventory manager has perfect information on the maximum order quantity. Table 3.1 shows the mean and standard deviation of optimality cost gap of the 1,000 synthetic test cases. When service requirement is great than or equal to 95%, IPS under all boundary information consistently outperforms the two benchmark solutions in terms of average optimality cost gap. The difference between solutions is less obvious when service requirement is 90%, but increases dramatically as the service requirement approaches 100%. The improvement from the best performing benchmark to proposed solutions range between 1.3% to 13.2% when service requirement is 95%, 9% to 22.9% when service requirement is 98%, and 23.8% to 42.4% when service requirement is 99%. As more accurate information on the bounds are used as inputs, the IPS algorithm significantly improves optimality cost gap.

Service	Normal	Max	IPS		
Req	Approach	Approach	BI_0	$BI_{SR}, \gamma = 1.5$	BI_{Exact}
90%	21.3%(.35)	17.4%(.30)	21.1%(.29)	18.9%(.29)	14.7%(.27)
95%	34.6%(.63)	40.1%(.69)	33.3%(.47)	28.7%(.48)	21.4%(.45)
98%	63.1%(1.27)	123.5%(1.86)	54.1%(.87)	49.4%(1.01)	40.2%(1.01)
99%	100.4%(2.22)	263.8%(3.71)	76.6%(1.59)	76.1%(1.85)	68.0%(1.90)

Table 3.1: Mean and standard deviation (in paranthesis) of Optimality Cost Gap of solutions under high service requirements.

The root cause of the difference in optimality cost gap is the effect of underestimation and overestimation. The Max Approach, which tends to overestimate the inventory target for the given critical fractile in the long run, is “limited” by the size of observations. When the dataset is small, e.g. less than 10 observations, the probability of observing a high demand is fairly low. As a result, the Max Approach almost always underestimates when service requirement is high, which is an undesired outcome since the unit cost for underestimation and overestimation are strongly asymmetric. The normal approach uses an one-size-fit-all logic, whose performance depends on the similarity between the true distribution and a normal curve. In comparison, IPS takes a different path towards the problem. It constructs individual orders by based on $\mathbf{d}_T$ and $\tilde{z}$. Every generated sample is consistent with $\sum_{t=1}^T z_t = \tilde{z}$ and $\sum_{k=1}^{z_t} w_{t,k} = d_t$. During the process, some extraordinarily high order quantities might be generated, which might not match the stream of orders that actually realized. However, they also self-regulate. If a high order quantity is sampled, then low order quantities must be sampled for the same period due to constraint of d_t limits. The key input is the boundary information, since it also limits how “extraordinary” the sampled orders can be. A tighter input bound is makes the solution less likely to overestimate. Figure 3-5 shows the percentage of underestimation, overestimation, and cases that each solution obtains the optimal inventory target when service

requirement is 98%.

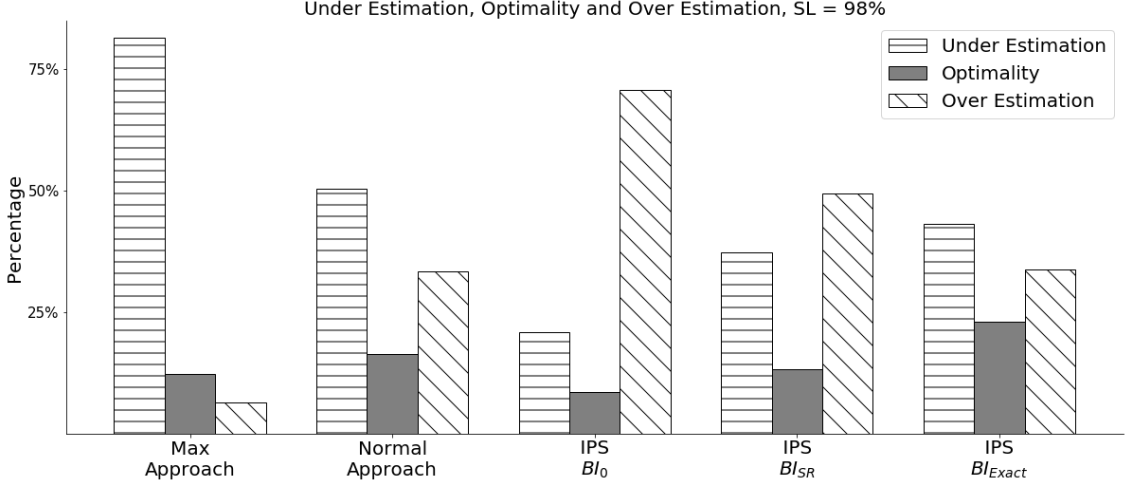


Figure 3.5: Under estimation, over estimation, and optimality of each heuristic when target service level is 98%.

3.4 Implementation at the W Facility

The implementation of IPS at the W facility is split in two sections: inventory targets for regular demand SKUs (average monthly demand greater than a preset threshold) and for low-demand SKUs (average monthly demand less than or equal to the threshold). We apply the normal approach to the regular demand SKUs since the number of feasible combinations becomes intractable for high $\mathbf{d}_T$ and $\tilde{z}$. The inventory targets for those SKUs are set with IPS for $\gamma = 2$ in order to be extra conservative. The targets are then calculated offline and sent back to the inventory manager at the W facility for feasibility validation. For regular bulk SKUs, since the size of its catalog variants should be considered, we estimate the mean and standard deviation from $\hat{F}_D$ and add up the demand from each associated catalog SKU as if they are independent random variables.

3.5 Conclusion

Motivated by the problem faced by the W facility at Company M, we introduce a new solution named Integer Pattern Sampling that takes advantage of the information on number of orders over a short horizon. The solution models demand as a compound process and exploits the integer possibilities for low-demand products. By uniformly sampling feasible combinations that are consistent with the observed data and prior boundary information, IPS consistently outperforms benchmark approaches in synthetic test cases. We believe this demonstrates the importance of characterizing tail behavior of demand distribution with limited observations.

Chapter 4

Identifying Inventory Reduction and Process Improvement Opportunities in Supply Chain

4.1 Introduction

Enterprise Resources Planning (ERP) has been installed to service companies (especially manufacturers) since the 1980s. Extending functions in *Material Requirements Planning* (MRP) and *Manufacturing Resources Planning (MRP II)*, ERP can incorporate modules for almost every business functions, from accounting and financials to human resources. During early 2000s, ERP became passe as a term, and was replaced by *Supply Chain Management (SCM)* (Hopp and Spearman, 2004). Despite the change in acronym, quantitative production and inventory control remains the core of both systems, which takes into consideration a wide range of factors, including but not limited to, bill of materials (BOM), network topology, capacity, required delivery time, inventory status, and the demand forecast. Although an integrated ERP system makes it possible for the company to manage a wide range of products, the system's rigidity can pose a barrier to companies that seek continuous improvement.

Company M (name disguised for company confidentiality) is a leading biotechnology product company that supplies life scientists with an extensive range of products

and services, from next-generation instruments to everyday lab essentials. Company M's global supply chain spans across the globe, and always views product quality and delivery promise as its strategic imperatives. Currently, Company M's supply chain is managed using Oracle's JD Edwards EnterpriseOne Requirement Planning (E1), with built-in the forecast engine—Demantra. Over the past few years, Company M has been seeking opportunities to streamline its supply chain by reducing excessive stock at various levels: raw materials, work-in-process (WIP), and finished goods (FG). We worked closely with the Supply Chain Data Analytics (SCDA) team at Company M on a project to develop a filter that identifies excessive inventory for a given single Stock Keeping Unit-Location (SKU at a location, referred as SKU-L) based on its past demand history. In addition to calculating financial benefit from inventory reduction, an equally important goal of the project is to deep dive into the root causes of overstocks. Similar to other central analytics team, the SCDA team is confronted with the following challenges:

Limitation to Parameter Adjustments Given the fact that existing operation of Company M's supply chain is built on E1, it is not realistic for the SCDA team to propose any changes that are incompatible with E1. Instead, the SCDA team focuses on fine tuning the parameters involved in the decision-making process. In this project, the SCDA team focused on the safety stock level input to E1. The narrower scope of the analysis, nevertheless, is a double-edged sword. On one hand, it simplifies the analysis for the SCDA team. On the other hand, it drives the attention away from the scrutiny of actual inventory trajectory to the metrics that surface service level impact.

Gap Between Recommendation and Implementation For operational

decisions associated with complicated production processes that necessitate domain expertise, the role of a central analytics team (similar to the SCDA in this project) is to suggest recommendations where potential improvement can be made. Likewise, all proposed actions by the SCDA need to be validated by inventory planners in charge before actual changes can be made on the factory floor. However, at a company that competes on service promises, the incentives for a central analytics team and inventory planners are usually different—the performance of a cost reduction project by central analytics team being measured by total amount of achieved savings, while performance of inventory planners being measured by percentage of on-time fulfillment. The key is to come up with the right metric for the communication with planners, which should capture not only the magnitude of inventory reduction, but also the impact on service level.

Preference for Non-Parametric Models with Lead Times Standard distributions like the Normal, and Poisson are commonly used by companies to characterize demand distributions, due to the ease in parameter estimation. However, most SKUs in the scope of this project have demand history that rejects the Shapiro test for normality and the Kolmogorov-Smirnov test for Poisson. This motivated the SCDA team to apply a distribution-free approach. Nevertheless, existing literature on non-parametric inventory models (Levi et al., 2007, 2015; Huh and Rusmevichientong, 2009; Huh et al., 2011) adopts a repeated newsvendor framework, which assumes zero replenishment lead time. Other literature, e.g. Bookbinder and Lordahl (1989) assume that independent lead time demand observations are available. At Company M, demand profile is analyzed at the weekly level corresponding to a weekly review period, while lead time can be as long as several weeks. One possible approach is to transform Company M’s problem into the format

of Bookbinder and Lordahl (1989), by first assuming i.i.d. weekly demand and then aggregate successive periods to obtain non-overlapping lead time demand. The downside of this aggregation approach is that it significantly reduces the length of demand history. At Company M, E1 automatically keeps a trailing 52-week demand history for each SKU-L. Thus, for a SKU-L with four-week replenishment lead time, the aggregation approach would reduce the original 52 observations to 13.

In this paper, we propose an approach that addresses these issues and identifies inventory reduction opportunities with almost no service impact. Our approach was first tested on a dataset collected from a central distribution center (DC) of Company M’s Antibodies (ANT) product line for initial validation. It was then developed into a standard tool on E1 and applied to Company M’s Life Science Solution (LSG) supply chain. The toolbox, according to Sean McCrary, Director of Supply Chain Analytics at Company M, “visualizes clear patterns of overstocking behaviors, and initiates conversations with inventory planners to dive deep into the root causes of excess inventory.”

4.2 Company M’s ANT Supply Chain and Current Operations

Antibodies (ANT) is a major product category offered by Invitrogen, a premier brand of Company M under its Life Science Solution Group, and is widely used by labs to identify, locate, measure and purify proteins as well as other biomolecules. ANT customers comprise of pharmaceutical and biotech companies, clinical and academic research institutions, and government departments. ANT supply chain can be modeled as a three-tier network (Figure 4-1): manufacturing sites (MFG) and

OEM, Tier 1 Distribution Centers (Tier 1 DCs), and Tier 2 Distribution Centers (Tier 2 DCs). In terms of service, Company M's standard policy is to complete delivery for stocked catalog products within 48 hours of receiving the order, which necessitates extremely high service levels at all stages of its supply chain. Given the long manufacturing and shipping lead time (compared to its demanding service promise), most ANT SKUs are make-to-stock inventory.

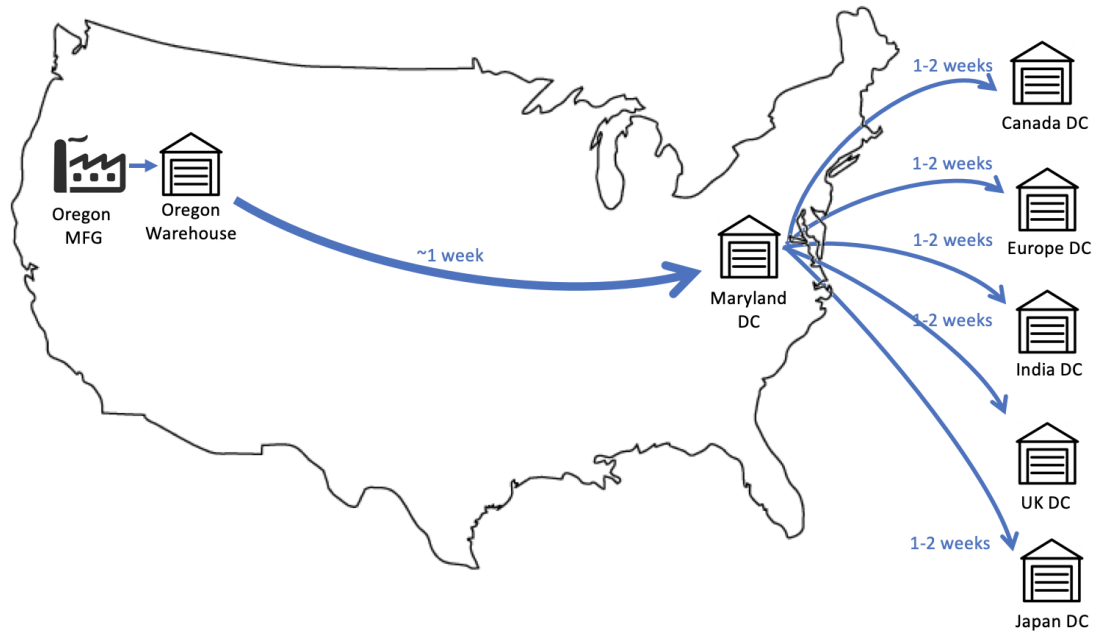


Figure 4·1: Manufacturing and distribution supply chain for ANT product family, which are either made in Company M's facilities in Frederick, Maryland and Eugene, Oregon, or outsourced to OEM.

Manufacturing Sites and OEM: ANT products are either produced at Company M's two factories or outsourced to OEM suppliers. Both facilities receive raw materials from upstream suppliers and transform them into end products (scope of manufacturing of Oregon MFG and Maryland MFG do not overlap). SKUs manufactured at Company M contribute to the majority of inventory in the ANT

FG supply chain, in both inventory quantity and value.

Tier 1 DCs: Both Maryland and Oregon MFG have a warehouse next to them that serves as a Tier 1 DC. ANT SKUs manufactured at Oregon MFG are either stored in Oregon DC, or in Maryland DC. Maryland DC is the major Tier 1 DC, carrying the majority of SKU selection and inventory value in the ANT FG supply chain. As a result, Maryland DC holds inventory in much larger quantities compared to Oregon DC and the regional Tier 2 DCs, and is therefore used as the primary dataset for validating our approach.

Tier 2 DCs: US demand is fulfilled directly from either Oregon DC or Maryland DC, while overseas demand is serviced by regional Tier 2 DCs to achieve shorter lead times. Domestic transportation between Oregon MFG and Maryland DC takes approximately one week; international shipping from Maryland DC to overseas Tier 2 DCs can take between 1 to 2 weeks.

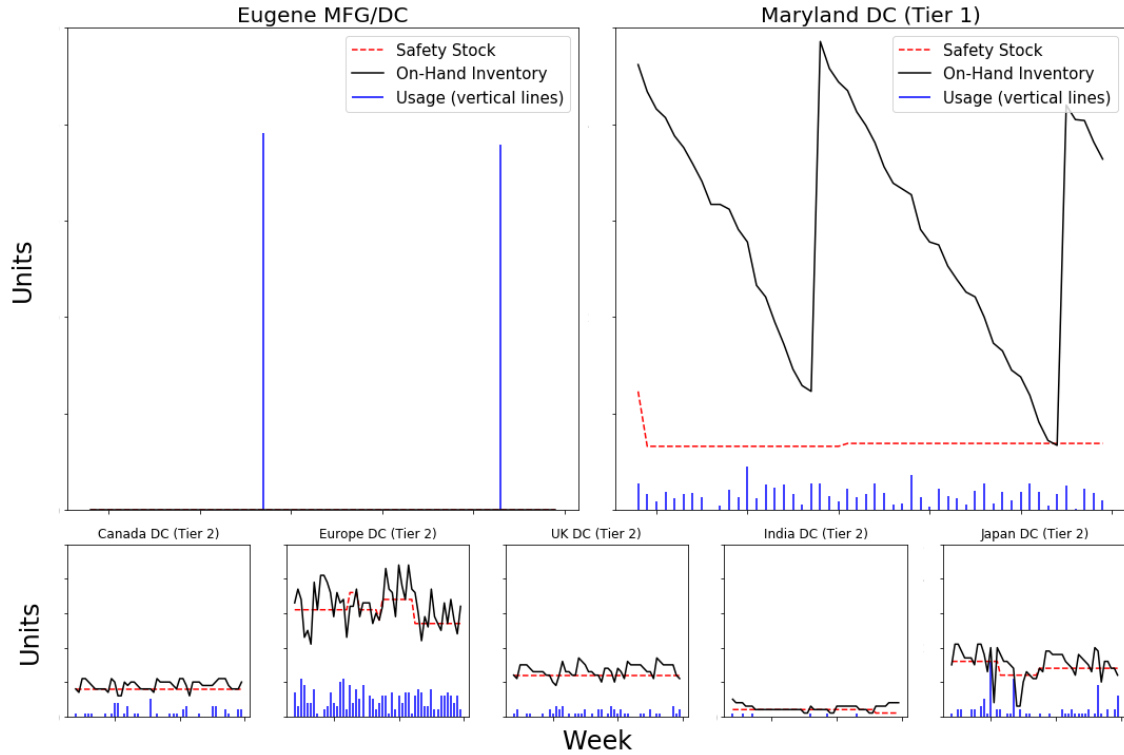


Figure 4.2: Usage profile for a ANT SKU example, which is manufactured at Oregon MFG and stored in Maryland DC, showing that most inventory is stored at Maryland DC.

At Company M, the terminology “usage” represents the number of units shipped out weekly from a facility. Figure 4.2 shows the usage profile of an example ANT SKU, which is manufactured at Oregon MFG and stored in Maryland DC (Tier 1). This SKU is also stored in five other Tier 2 DCs in local markets. The SCDA team approximates actual demand profile for a SKU-L by its usage over a trailing 52-week history. Although “usage” is a censored demand observation (and may include small time lag). It is a close approximation of actual demand given the high service level (above 95%).

Inventory Control Policies: At a monthly cadence, Demantra updates the daily demand forecast at the SKU-L level and covers a 52-week horizon into the

future. At a daily cadence, E1 consumes the demand forecast and compares the total inventory position—sum of on-hand and in-transit inventory, of a SKU-L against a dynamic re-order point (ROP_{SKU-L}). The re-order point consists of two parts, forecasted demand $F(L)_{SKU-L}$ over the SKU's lead time L , and a preset safety stock level, SS_{SKU-L} , which is updated every six month:

$$ROP_{SKU-L} = F(L)_{SKU-L} + SS_{SKU-L} \quad (4.1)$$

If total inventory position of a SKU-L is less than ROP_{SKU-L} , an order is placed against its upstream supplier (an internal transit order at an upstream DC or a work order at MFG).

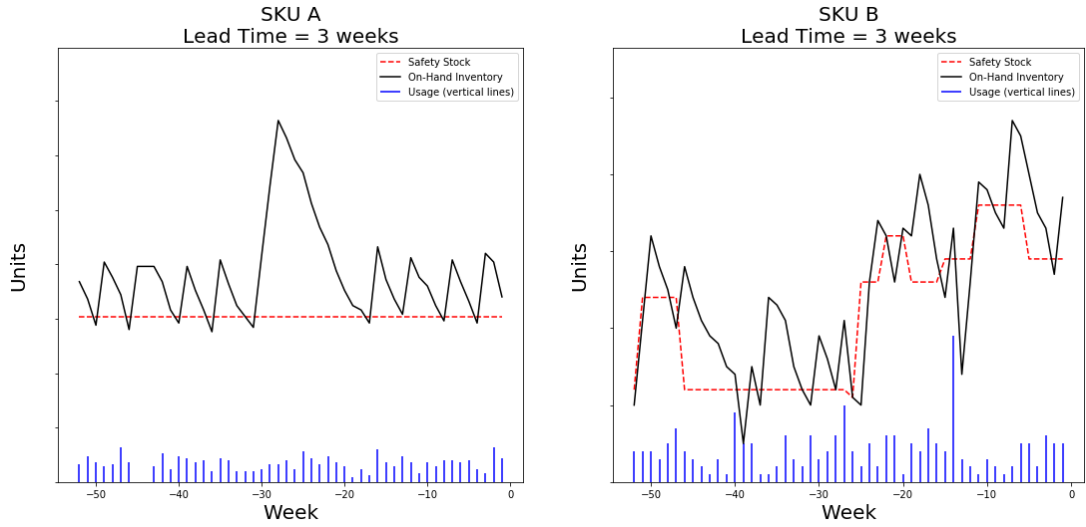


Figure 4.3: Two SKU examples showing that frequencies of safety stock levels adjustment vary between SKUs.

When on-hand inventory is not sufficient to fulfill demand, backlogs are generated and fulfilled at earliest instance possible. At Company M, service is measured by first-order fill rate, which is the percentage of orders filled on the date the customer (both internal and external) first asks for the product. First-order fill

rate conform to a traditional Type I (order fill rate) or Type II (volume fill rate) service level in a textbook. Moreover, due to the large amount of SKUs, the SCDA team performs all the data analytics at the weekly level, which makes it hard to calculate a counter-factual first-order fill rate metric. During the early stage of the project, we had reached a consensus that we use in-stock rate as the service metric.

There are two types of staff for inventory planning at Company M: “Production Planners” take charge for production scheduling and manage the required materials (including raw materials and WIP); “Distribution Requirements Planners” move FG inventory across the network to meet global demand. By default, both types of planners are recommended to set SS_{SKU-L} to be one of the three levels: none, low, high, each corresponds to certain weeks of supply, where the average weekly demand μ_{SKU-L} is updated quarterly using the trailing 52-week of demand history. Furthermore, both types of planners can manually override the safety stock level at any time based on their experience and information on upcoming events (for instance, anticipating that demand for a SKU-L will increase in near future). Collected data also shows that SKUs may receive different “attention” from inventory planners. Figure 4.3 shows two actual SKU examples: for the same 52-week horizon, safety stock for SKU A (left) remains constant, while safety stock for SKU B (right) was changed more frequently. Table 4.1 summarizes the number of SS changes for ANT FG inventory at the five major DCs in the ANT product supply chain over a 52-week horizon.

4.2.1 Problem with Only Considering Safety Stock

As a common treatment in practice, safety stock level at Company M is a function of “weeks of supply.” We first examine the existing practice of safety stock level in E1.

DC	Mean of #Changes in <i>SS</i> Parameters	Std. of #Changes in <i>SS</i> Parameters
Europe DC	1.4	1.6
Maryland DC	1.2	1.0
Oregon DC	0.2	0.5
Japan DC	1.8	1.4
UK DC	0.1	0.3

Table 4.1: Statistics of number of changes in safety stock parameters over 52-week horizon at 5 major DCs in ANT supply chain.

Figure 4.4 (left) shows the histogram of weeks of supply of all ANT SKU-L in the latest snapshot in the dataset. Only 28.5% of the ANT SKU-Ls have safety stock level within the recommended level, and 20.1% of the SKU-Ls have safety stock level parameters higher than 10 weeks of supply, which shows the existence of excessive inventory even by SCDA's standard. Widely used in practice, weeks of supply can be an effective metric in measuring how much inventory is being held compared to average consumption. However, it is not the best metric to adhere when considering the level of safety stock, which is tied to the variance of demand rather than the mean.

We also compare the safety stock level to the case where i.i.d. normal demand is assumed. Figure 4.4 (right) shows a histogram of the service level equivalence (z^*) for each SKU-L under a normal assumption, where $z^* = \Phi^{-1}(SS/\hat{\sigma}\sqrt{L})$, SS being the current safety level, $\hat{\sigma}$ being estimated standard deviation, and Φ^{-1} being the inverse c.d.f. of a standard normal distribution. Although we have shown that most ANT SKU-L do not satisfy a normal distribution, we demonstrate the service level equivalence to show the existence of excess stock.

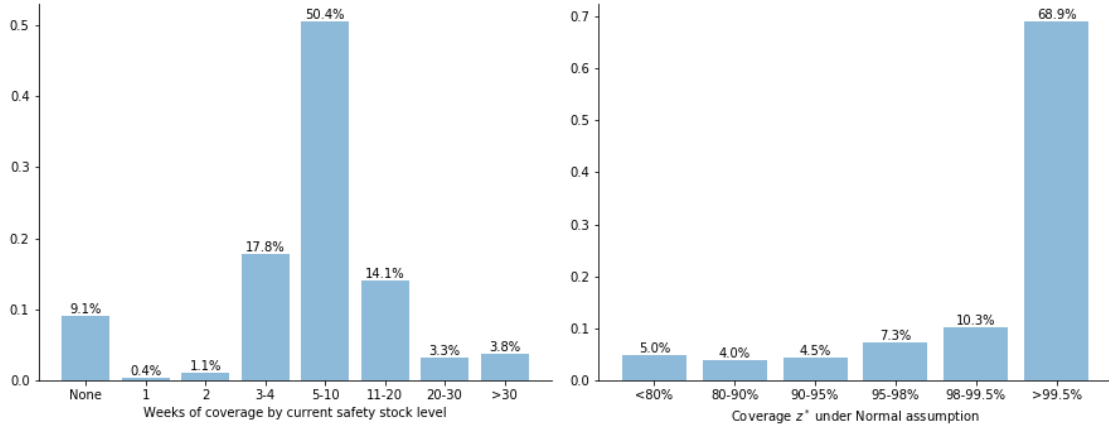


Figure 4.4: Left: distribution of weeks of supply of safety stock parameter snapshot for all ANT SKU-L; right: service level equivalence for all ANT SKU-L under Normal assumption.

Scrutinizing the data reveals inconsistent patterns around safety stock, which partially explains the high weeks of coverage in Figure 4.4. Figure 4.5 shows three examples where (a) depicts a scenario where safety stock is treated as inventory that was never touched within the 52-week horizon, which is commonly seen at central DCs (Oregon DC and Maryland DC for the ANT products); (b) depicts a scenario where safety stock is viewed similarly to the order-up-to level in a base-stock policy; (c) depicts a scenario where safety stock seems to be a general guideline of how the average inventory level is maintained. Although category (b) and (c) constitute the majority of SKU-L at Tier 2 DCs, we did not find a clear mapping between categories and DCs or between categories and planners.

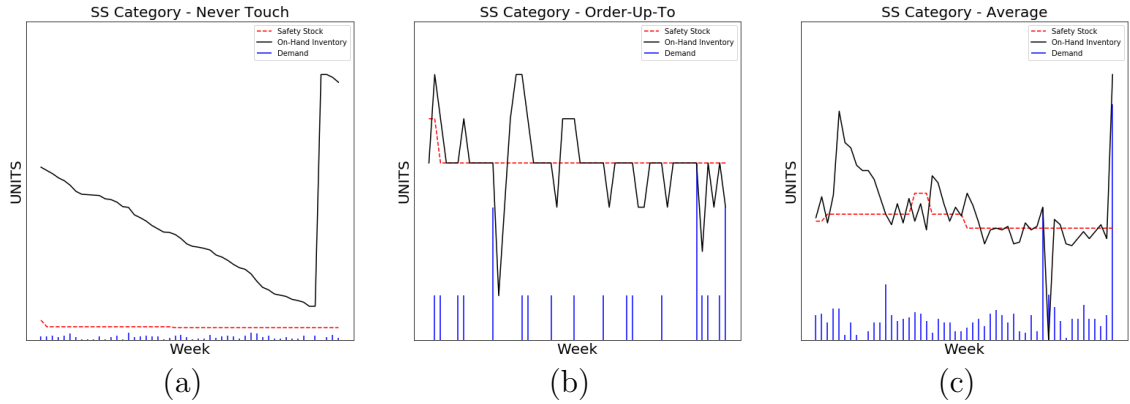


Figure 4-5: Examples of three scenarios indicating that the metric “safety stock” is used differently for different SKUs.

A SKU-L is classified into “Never Touch” category if on-hand inventory is constantly above the safety stock parameter throughout the 52-week period, into “Order-Up-To” category if on-hand inventory is less than the latest safety stock level snapshot for at least 75% of the time, and into “Average” category if for at least 25% of the time, on-hand inventory is higher than and lower than the latest safety stock level snapshot. Since the same parameter may represent different reference point, recommendations that only consider safety stock level can be misleading. Reducing safety stock for SKU-L that never touch the safety stock level has no service impact, it can leads to lower service level for SKU-L that use safety stock parameter as an order-up-to quantity. We also limit Figure 4-4 to Maryland DC (Figure 4-6), which possesses the greatest potential for savings.

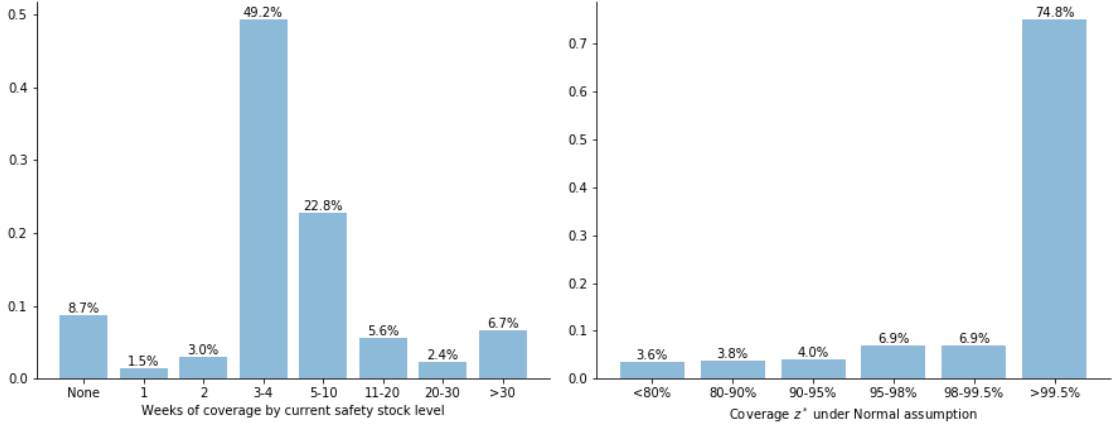


Figure 4-6: Left: distribution of weeks of supply of latest safety stock level snapshot for all ANT SKU-L; right: service level equivalent for all ANT SKU-L under Normal assumption.

4.3 Maximum Forward Coverage Heuristic

As discussed previously, the goal is to develop a tool within E1 that can perform a real-time calculation for any set of SKU-Ls, rather than proposing a brand new ordering policy. More specifically, the tool should be able to quickly identify abnormal inventory behaviors that surface issues underneath current operations, such as unnecessary trigger of procurement or work orders, unreasonable parameters, or even miscounted inventory.

We propose a simple non-parametric Maximum Forward Coverage (MFC) heuristic. Different from existing literature on data-driven (non-parametric) inventory models (e.g. Scarf, 1958; Kleywegt et al. 2001; Levi et al. 2007, 2015; Perakis and Roel 2008), we take the lead time (a SKU-L dependent fixed parameter) into consideration and focus on identifying SKUs with extraordinarily high inventory levels compared to their demand profile. Adopting the terminology used at Company M, let $\mathbf{x} = (x_{-52}, x_{-51}, \dots, x_{-1})$ denote the weekly usage data of a certain SKU-L

collected over a trailing 52-week interval, where x_{-t} denotes the weekly usage within the t -th week in the past. Our approach first obtains a derived lead time demand $y_{-t} = \sum_{s=0}^{L-1} x_{-t+s}$ from existing data by adding demand from successive weeks. Then the maximum lead time demand, y^{MFC} , is compared to the minimum historical inventory level, $\min I_{-t}$, and the cost-weighted gap is identified by δ^{MFC} (detailed derivation described in Appendix J).

$$\delta_{SKU-L}^{MFC} = C_{SKU} \max(\min I_{-t} - y^{MFC}, 0) \quad (4.2)$$

Prior to our participation in the project, the SCDA team developed an approach (referred as the “M Approach”) to calculate a new safety stock level for any given SKU-L. The M Approach calculates a new safety stock level based on the latest weeks of supply and a trimmed mean, in which a parameter $0 < p < 1$ controls how many “extraordinarily high” demand observations are dropped from history to calculate the trimmed mean. A higher p leads to fewer dropped observations and leads to a more conservative reduction (detailed description of the M approach described in Appendix I). Despite the risk of using the safety stock level as the target parameter (which can be used differently across the network), The M approach represents an archetypal logic in inventory practice, which relies on mean and weeks of coverage for measuring amount of buffer provided by safety stock. According to the SCDA team, the average usage numbers overestimated the larger population of customer orders. Replacing the mean demand with a trimmed mean indicates the intention of using safety stock to protect demand variability from “normal” weeks, and use extraordinary measures to cover weeks with high demand. The M Approach tries to justify a lower safety stock by not using it to cover “extraordinarily high demand,” which is exactly the intention of safety stock. Another case in which the

M Approach may fail is SKUs with intermittent demand. Such case can result in zero inventory when applying the M Approach, which can significantly impact service.

When lead time distribution is stationary, regardless of correlations between successive periods, y^{MFC} is a consistent estimator of the right support of the lead time distribution. In other word, δ^{MFC} captures the gap between historical sample path of inventory and an extremely conservative inventory target, estimated by y^{MFC} . To validate the approach, we compare the M Approach, and MFC in a simulation-based test using actual dataset from Maryland DC. As a benchmark, we also include the commonly used parametric method assuming that the weekly demand has an i.i.d. normal distribution, which will be referred as the “Normal Approach.” Prior to the numerical test, SKUs that shows strong correlations (absolute correlation coefficients within 4 period lags greater than 0.3) and non-stationarity (fail to reject the null-hypothesis of Dickey-Fuller Test) are dropped from the dataset. For each approach, we report three metrics: (i) Service Level (in-stock rate) under Average Case (SL-AC), Service Level Lower Bound (SL-LB), and Reported Inventory Reduction (RIR). Since we tend not to include Demantra’s forecast as part of the model (nor does the SCDA team), SL-AC measures the service level when Demantra’s forecast is proxied by the mean lead time demand. SL-LB measures the service level lowerbound by replacing Demantra’s forecast with zero, in the case of which, all lead time demand will be satisfied by safety stock. And RIR captures the inventory reduction opportunities identifies by each approach. Appendix K describes the setting of simulation-based test. Table 4.2 summarizes the test results. Numbers are disguised to protect company confidentiality

The M approach, despite the lack of direct linkage between parameter p and

Approach	M App		Normal	MFC
p	0.95	0.9	0.95	na
SL-AC	95.4%	94.4%	87.3%	99.9%
SL-LB	93.5%	92.2%	66.6%	99.8%
RIR	\$41.5	\$79.2	\$192.0	\$191.0

Table 4.2: Disguised results of the simulation based test at the Maryland facility. All numbers have been disguised to protect company confidentiality.

service level, shows an average service level between 94.4% and 95.4%, and a lower-bound between 92.2% and 93.5%, indicating a potential service impact compared to the existing service level. The widely used Normal assumption performs poorly, indicating that empirical distribution tend to have a heavier tail than normal distributions. MFC, by design, successfully drives attention to the SKUs where significant inventory reduction can be done with almost guaranteed service.

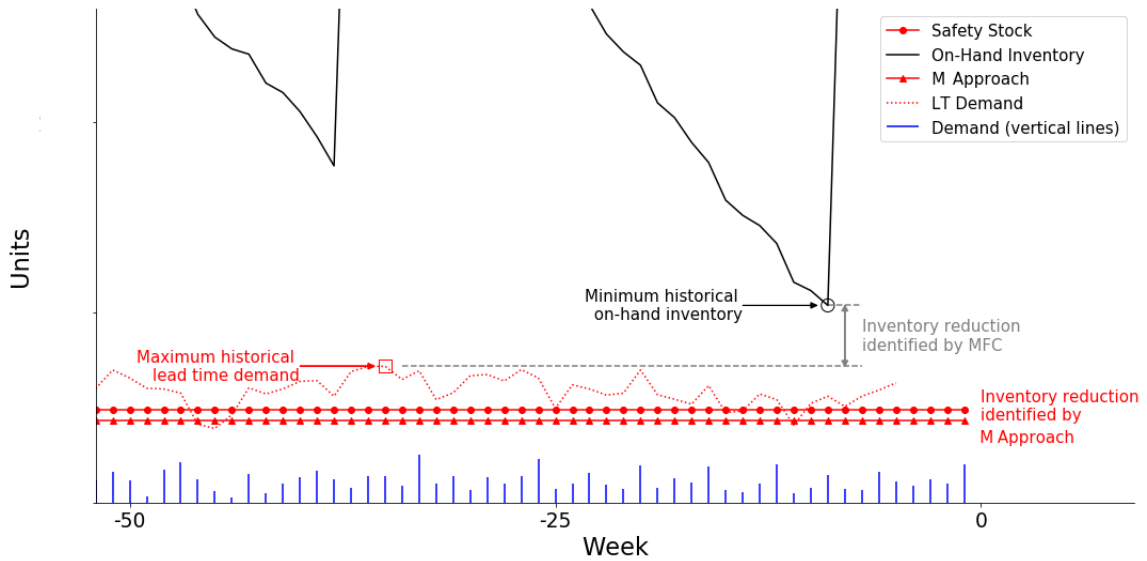


Figure 4.7: Different portions of inventory are identified as reduction opportunities by the M Approach and MFC.

Unlike the M Approach and Normal approach, which report inventory reduction

between the current and newly prescribed safety stock level, MFC reports saving as the gap between the lowest historical inventory level. Figure 4.7 demonstrates that different portions of inventory are identified as reduction opportunities. The more interesting question to ask is which can help the company to identify areas with the most potential for improvement. MFC, does this by comparing the actual realization of a potentially complex inventory control policy with the most conservative cases, and identifies the gap. We also compare the list of the top 20 SKUs identified by each approach, and find that while the M Approach and Normal point to basically the same set of SKUs, MFC shares only 1-2 SKUs with the rest. Given that MFC identifies inventory reduction which is unlikely to cause a service impact, we can conclude that the M Approach and Normal Approach have their limitation by only looking at a single parameter.

4.4 Full Implementation and Discussion

We present the results of implementing MFC on Maryland DC to the SCDA team and the Senior Director of Supply Planning at Company M. Top 20 SKUs with the most inventory reduction at Maryland DC were brought to corresponding planners to validate the reduction opportunities and investigate the root causes. Feedback from planners confirmed the identified opportunities. The SCDA team also found out that MFC is not only robust for different types of SKUs (the non-parametric feature), but it is also easier to visualize MFC results for planners than trying to argue that safety stock should be reduced without showing that service level will be preserved at the same time. MFC was later developed into a toolbox that is available for quickly filtering SKUs for any given facility. In addition, Company M imposes another 20% buffer on y^{MFC} , resulting an even more conservative metric.

When using the tool, a planner specifies a site and the scope of SKUs. MFC then generates a list consisting of the top SKUs with the most reduction opportunities. Planners can then deep dive into the root cause for causing a certain SKU to have excess inventory and summarize the findings in a linked spreadsheet for further analysis.

According to the Director of Supply Chain Analytics team, “despite its simplicity, the approach is really powerful in framing a narrative. When the chart shows that a layer of inventory can clearly be removed, it directly results to the action of finding the root cause for such phenomenon.” During the first phase of implementation in early 2017, the tool was distributed to all Company M’s Tier 1 DCs. Each inventory planner in the supply chain was asked to comment on the root cause of the the top five SKUs identified by MFC. The first phase identifies multi-million dollars of saving opportunities in all stages of Company M’s major product families. Long-term impact of the SKUs that received comments are shown in Figure 4-8.

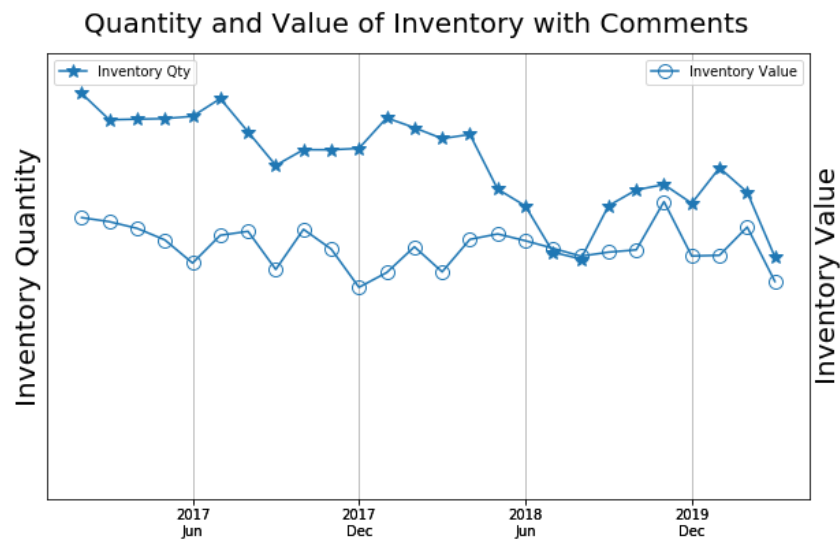


Figure 4-8: Longitudinal impact on inventory identified by MFC at Company M.

Rank	Cause	Percentage
1	Planning Parameter Driven	19.0%
2	Business Requirement	13.1%
3	Item Discontinuation	9.2%
4	Forecasts	8.1%
5	Production Strategy	6.4%

Table 4.3: Top 5 category of root causes of excessive inventory identified by MFC.

After the first phase of the implementation came out as a success, Company M has made the SKU ranking with MFC a monthly check procedure, in which each inventory planner comments on potential root causes and corrective actions for the top SKUs identified by the MFC approach. Table 4.3 lists the top 5 categories that contribute to excessive inventory identified by MFC. The analysis is based a random selection of recent comments from inventory planners. Major corrective actions involve adjusting parameters, rebalancing inventory, and investigation in potential scrap.

Chapter 5

Conclusions

5.1 Summary of the thesis

This thesis considers inventory problems under limited data, and high customer service level requirements. This chapter summarizes results in each chapter and discuss future research direction.

Setting Inventory Targets with Compound Demand in Data-Scarce Environments

Chapter 2 and Chapter 3 consider a relatively similar problem, in which the decision maker observes periodic demand with some prior knowledge on the boundary condition of the order quantity distribution. Chapter 2 formulates a theoretical framework based on a repeated newsvendor model with a compound demand process. The decision maker observes not only periodic demand but also the number of customers in each period. Two proposed approaches, Maximum Likelihood Estimation and Metropolis-Hastings sampler, take advantage of the integer nature of the demand process and construct estimations for order quantity distribution. Both algorithm outperform benchmark approaches that only consider periodic demand, and show different sensitivity to input boundary information.

Chapter 3 considers a more practical problem faced by the inventory manager

at a distribution center of a major biotechnology firm, which actually motivates the theoretical work in Chapter 2. However, even less information was provided during the project and the customer service level requirement was very high. We propose the Integer Pattern Sampling approach in order to sample feasible combinations that are consistent to the observations.

One interesting future research area is to consider the number of periods required to provide a specific probabilistic guarantee. This has been shown for the Sample Average Approximation in a repeated newsvendor framework by Levi et al. (2007) and Levi et al. (2015). Since this thesis models demand as a compound process, there lacks an analytical handle to capture quantile behavior of the outcome random variable after the convolution. However, it would be worthwhile to investigate the probability guarantee when arrival process takes a more tractable form, such as Poisson. It is also interesting to find a more efficient way of numerically convolving two random variables, which takes the most of computation time at this point.

Identify Safe Inventory Reduction Opportunities

Chapter 4 takes a fresh look at the objective function of a typical inventory improvement project by an analytics team at a firm. Since plenty of firms have automated their inventory management using third-party ERP software, the flexibility to change can be very limited. Thus, a typical solution for such project starts with a scrutiny of existing parameters. We show that sometimes such a route can be misleading, especially when terminologies are defined differently. Instead, we propose a heuristic that directly compares on-hand inventory and a conservative inventory target derived from past demand data. The heuristic turns out to be very useful and has been developed into a monthly routine at the firm that we collaborated with. A potential

future research direction is to carefully describe a firm's existing reordering policy and develop a more rigorous data-driven method that can prescribe inventory targets with a probabilistic guarantee.

Appendix A

Recursion for Calculating $\tau(\delta)$

Algorithm 3 Recursion for generating number of permutations for $\delta \in \Delta$

```
1: procedure  $\tau(\delta)$ 
2:    $\tau \leftarrow 1, n = \dim(\delta)$ 
3:   for  $w \leftarrow \min(\delta) : \max(\delta)$  do
4:      $n_w \leftarrow \sum_{k=1}^{z_t} \mathbb{I}[w_k = w]$ 
5:      $\tau \leftarrow \tau \times \binom{n}{n_w}$ 
6:      $n \leftarrow n - n_w$ 
7:   end for
8:   Return  $\tau(\delta) \leftarrow \tau.$ 
9: end procedure
```

Appendix B

Proof Lemma 1

PROOF. OF LEMMA 1. (i) Since ϕ_{IC} assigns equal probability to all points on the probability simplex, the proof is trivial. To show that ϕ_{MHR} is a transition kernel for an irreducible Markov chain, it is equivalent to show that it can reach any point $\mathbf{q}^*$ in the probability simplex from any initial point $\mathbf{q}^{(0)}$. In fact, we can show that there is a finite path from $\mathbf{q}^{(0)}$ to $\mathbf{q}^*$. Consider a path $\mathbf{q}^{(0)}, \mathbf{q}^{(1)}, \dots$, in an n -dimension probability simplex, in which $\mathbf{q}^{(m)} = (q_1^{(m)}, \dots, q_n^{(m)})$, with subscript i referring to the dimension index rather than $P\{W_{t,k} = i\}$. Assume that for some point $\mathbf{q}^{(m)}$, we have $q_i^{(m)} = q_i^*, \forall i = 1, \dots, k$ and $q_{k+1}^{(m)} \neq q_{k+1}^*$, we can show that, within finite steps, we can generate a state $\mathbf{q}^{(m+m')}$, such that $q_i^{(m+m')} = q_i^*, \forall i = 1, \dots, k+1$. For $q_{k+1}^{(m)}$ and q_{k+1}^* , we have two cases: (I) $q_{k+1}^{(m)} > q_{k+1}^*$, and (II) $q_{k+1}^{(m)} < q_{k+1}^*$. For (I), we can write $q_{k+1}^{(m)} = q_{k+1}^* + \epsilon, \epsilon > 0$. By ϕ_{MHR} , we can redistribute $q_{k+1}^{(m)}$ and $q_{k+2}^{(m)}$ such that $q_{k+1}^{(m+1)} = q_{k+1}^*$ and $q_{k+2}^{(m+1)} = q_{k+1}^{(m)} - \epsilon$, which gets to the desired $\mathbf{q}^{(m+m')}$ in one step. For (II), assume $q_{k+1}^* = q_{k+1}^{(m)} + \epsilon, \epsilon > 0$. Since

$$\begin{aligned}
 \sum_{i=k+2}^n q_i^{(m)} - \epsilon &= \left(1 - \sum_{i=1}^k q_i^{(m)} - q_{k+1}^{(m)} \right) - \epsilon \\
 &= \left(1 - \sum_{i=1}^k q_i^* - (q_{k+1}^* - \epsilon) \right) - \epsilon \\
 &= \sum_{i=k+2}^n q_i^* \geq 0
 \end{aligned} \tag{B.1}$$

we can have $q_{k+1}^{(m+m')} = q_{k+1}^*$ by iteratively adding the dimensions $k+3, k+3, \dots$ to the $k+1$ -th dimension, which can be obtained by $m' \leq n-k-1$ intermediate states. Repeating this process from $k=0, \dots, n$, we can construct a path from any given $\mathbf{q}^{(0)}$ to $\mathbf{q}^*$ within length $n(n-1)/2$. On the other hand, there is a positive chance that $\mathbf{q}^{(m)} = \mathbf{q}^{(m+1)}$ with ϕ_{MHR} , thus the Markov chain with transition kernel ϕ_{MHR} is aperiodic on the probability simplex.

(ii) Part (i) of Lemma 1 shows that for any two given point $\mathbf{q}^{(0)}$ and $\mathbf{q}^*$ in the n -th dimension probability simplex, there is a path with non-zero probability that starts from $\mathbf{q}^{(0)}$ and ends at $\mathbf{q}^*$ under either ϕ_{IC} or ϕ_{MHR} . We need to show that the transition kernel of the MH sampler, $K_S(\cdot, \cdot)$, is π_T -irreducible, which is expressed as:

$$K_S(\mathbf{x}^{(i-1)}, \mathbf{x}^{cand}) = \begin{cases} \phi_S(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) \alpha_T(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) & , \text{ if } \mathbf{x}^{cand} \neq \mathbf{x}^{(i-1)} \\ 1 - \int \phi_S(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) \alpha_T(\mathbf{q}^{(i-1)}, \mathbf{q}^{cand}) & , \text{ if } \mathbf{x}^{cand} = \mathbf{x}^{(i-1)} \end{cases} \quad (\text{B.2})$$

For ϕ_{IC} , since it is always able to move from one point to another in $\{\mathbf{q}; \pi_T(\mathbf{q}) > 0\}$, it is π_T -irreducible. For ϕ_{MHR} , consider a path between $\mathbf{q}^{(0)}$ and $\mathbf{q}^*$. Part (i) shows that there is a path $\mathbf{q}^{(0)}, \dots, \mathbf{q}^{(n)}$ with $\prod_{k=0}^{n-1} \phi_S(\mathbf{q}^{(k)}, \mathbf{q}^{(k+1)}) > 0$ and $\mathbf{q}^{(n)} = \mathbf{q}^*$. Under $K_S(\cdot, \cdot)$, the probability of the path is $\prod_{k=0}^{n-1} \phi_S(\mathbf{q}^{(k)}, \mathbf{q}^{(k+1)}) \alpha_T(\mathbf{q}^{(k)}, \mathbf{q}^{(k+1)})$. In order to show $K_S(\cdot, \cdot)$ is π_T -irreducible, it is equivalent to show that $\alpha_T(\mathbf{q}^{(k)}, \mathbf{q}^{(k+1)})$ is always positive when $\pi_T(\mathbf{q}^{(k)})$ is positive. In fact, if $\pi_T(\mathbf{q}^{(k)}) > 0$, by the construction of $\mathbf{q}^{(k)}$ to $\mathbf{q}^{(k+1)}$, ϕ_{MHR} only redistributes two p_i and p_j without rendering anyone of them to zero. Thus $\pi_T(\mathbf{q}^{(k+1)}) > 0$ and there is always a non-zero probability to accept $\mathbf{q}^{(k+1)}$, which completes the proof.

Appendix C

Proof Theorem 1

PROOF. OF THEOREM 1 (i) Roberts and Smith (1994) shows that, under the general condition, if the transition kernel of the MH sampler is π -irreducible and aperiodic, then the convergence is true. Lemma 1 shows that ϕ_{IC} and ϕ_{MHR} are both π_T -irreducible. Meanwhile, the aperiodicity of ϕ_{IC} and ϕ_{MHR} also guarantees that the corresponding MH sampler is aperiodic (see Roberts and Smith, 1994, Thm 3(i)). Thus the theorem holds for the MH sampler of interest. For a more detailed proof of the general case, see Tierney (1994) and Nummelin (1984). (ii) Mengersen et al. (1996), Thm 2.1 show that an independence chain is uniformly ergodic if there exists $\beta > 0$ such that

$$\frac{\phi_{IC}(\mathbf{q})}{\pi_T(\mathbf{q})} \geq \beta \tag{C.1}$$

then the MH sampler is uniformly ergodic. In this case, ϕ_{IC} is a constant and Equation 2.14 shows that the posterior distribution is bounded from above, which guarantees Equation C.1 holds.

Appendix D

Statistics for Test Cases

Statistics table for the test cases used in the numerical study. Each test case is a compound process from two independent integer random variables. The support of the outcome random variable are 0 and 16, with other statistics summarized in the following table

Test Case	Z_t	$W_{t,k}$	Mean	Coef.Var	Skewness	Kurtosis
1	Uniform	Uniform	4.0	0.9	0.6	2.5
2	Uniform	Decreasing	2.2	1.0	1.0	3.6
3	Uniform	Increasing	5.8	0.8	0.3	2.0
4	Uniform	Normal-like	4.0	0.8	0.4	2.3
5	Uniform	Ushape	4.0	0.9	0.7	2.7
6	Decreasing	Uniform	2.2	1.3	1.3	4.2
7	Decreasing	Decreasing	1.2	1.5	1.7	5.9
8	Decreasing	Increasing	3.2	1.2	1.0	3.3
9	Decreasing	Normal-like	2.2	1.2	1.2	3.7
10	Decreasing	Ushape	2.2	1.4	1.4	4.4
11	Increasing	Uniform	5.8	0.6	0.1	2.4
12	Increasing	Decreasing	3.2	0.8	0.6	3.0
13	Increasing	Increasing	8.4	0.5	-0.3	2.4
14	Increasing	Normal-like	5.8	0.5	-0.1	2.4
15	Increasing	Ushape	5.8	0.7	0.2	2.4
16	Normal-like	Uniform	4.0	0.7	0.5	3.0
17	Normal-like	Decreasing	2.2	0.9	1.0	3.7
18	Normal-like	Increasing	5.8	0.6	0.3	2.8
19	Normal-like	Normal-like	4.0	0.6	0.4	3.0
20	Normal-like	Ushape	4.0	0.8	0.6	2.9
21	Ushape	Uniform	4.0	1.0	0.6	2.1
22	Ushape	Decreasing	2.2	1.2	1.0	3.4
23	Ushape	Increasing	5.8	0.9	0.2	1.6
24	Ushape	Normal-like	4.0	0.9	0.4	1.8
25	Ushape	Ushape	4.0	1.0	0.7	2.4

Table D.1: Statistics of 25 test cases. Each generates 40 random sample paths for the numerical experiment in Section 2.5

Appendix E

Table appendix

SL	T	Parametric		Non-Parametric		MLE			MH								FED
		Poisson	Normal	Max	SAA	4	6	8	IC			Self-Reg	MHR				
									4	6	8		4	6	8		
90	4	45.7%	35.9%	34.9%	36.2%	26.4%	27.0%	27.9%	19.7%	31.3%	51.3%	25.8%	19.6%	31.6%	52.5%	27.5%	
	6	36.2%	23.6%	21.5%	23.8%	18.2%	18.5%	19.3%	14.1%	24.5%	44.2%	18.2%	14.1%	25.1%	45.4%	17.4%	
	8	33.6%	18.6%	17.8%	17.1%	13.6%	13.6%	14.1%	10.8%	19.9%	38.3%	13.4%	10.9%	20.6%	39.9%	12.9%	
	10	31.1%	14.8%	16.0%	13.5%	11.7%	11.9%	12.1%	9.2%	16.9%	33.8%	10.9%	9.4%	17.7%	35.8%	10.7%	
	12	28.3%	11.9%	15.4%	14.6%	8.9%	9.1%	9.5%	7.5%	14.9%	31.1%	9.2%	7.8%	16.0%	33.4%	8.4%	
95	4	82.9%	64.8%	81.8%	82.6%	46.6%	46.4%	47.9%	30.8%	42.0%	64.1%	39.4%	30.9%	42.3%	64.9%	47.0%	
	6	65.3%	41.3%	45.9%	47.8%	28.9%	29.5%	30.7%	20.0%	31.4%	54.1%	26.4%	20.0%	31.9%	55.2%	26.0%	
	8	58.9%	32.3%	32.1%	33.5%	19.9%	20.1%	20.8%	13.5%	24.4%	46.6%	18.2%	13.6%	25.1%	47.9%	18.1%	
	10	54.6%	25.7%	23.4%	23.8%	15.8%	15.6%	16.3%	10.9%	20.7%	41.6%	14.3%	11.2%	21.4%	43.5%	13.4%	
	12	52.4%	20.1%	17.4%	18.7%	12.9%	13.0%	13.3%	8.7%	18.3%	38.7%	12.0%	9.0%	19.4%	40.8%	10.9%	
98	4	164.7%	135.6%	228.6%	228.6%	104.3%	100.3%	102.9%	62.7%	68.5%	91.9%	74.7%	62.8%	69.0%	92.6%	98.8%	
	6	126.6%	80.9%	134.8%	134.9%	60.9%	59.7%	61.3%	36.6%	46.3%	71.5%	46.6%	36.7%	46.8%	72.5%	49.4%	
	8	115.6%	62.1%	95.2%	95.1%	36.5%	35.7%	37.0%	20.6%	32.2%	57.4%	28.6%	20.7%	32.9%	58.5%	31.3%	
	10	105.9%	47.9%	69.5%	69.6%	25.9%	25.8%	26.8%	15.5%	26.1%	50.6%	21.4%	15.6%	26.8%	52.4%	20.4%	
	12	99.1%	36.5%	49.8%	49.9%	18.8%	19.0%	19.6%	10.9%	22.2%	46.6%	16.0%	11.3%	23.4%	48.8%	15.1%	
99	4	267.0%	219.7%	467.1%	467.1%	191.3%	181.8%	181.4%	113.6%	108.6%	132.5%	129.6%	113.3%	108.6%	133.1%	176.9%	
	6	196.1%	133.0%	283.8%	283.8%	105.7%	101.0%	102.9%	62.5%	67.5%	93.9%	75.1%	62.6%	68.1%	94.8%	81.7%	
	8	186.9%	98.6%	204.8%	204.8%	60.2%	58.8%	61.1%	32.0%	41.5%	68.5%	41.9%	32.2%	42.1%	69.6%	49.9%	
	10	164.2%	75.6%	153.1%	153.1%	41.2%	40.8%	42.1%	22.2%	32.1%	58.5%	29.1%	22.6%	32.9%	60.2%	31.4%	
	12	151.1%	56.2%	112.6%	112.6%	28.6%	28.3%	29.2%	14.4%	26.2%	52.8%	21.6%	14.6%	27.4%	55.0%	20.5%	
99.5	4	430.0%	392.5%	928.9%	928.9%	364.9%	346.4%	343.3%	209.5%	182.4%	206.2%	223.8%	209.9%	182.5%	206.9%	327.2%	
	6	322.2%	219.9%	575.3%	575.3%	189.8%	181.6%	182.1%	111.2%	105.1%	132.1%	124.5%	111.1%	105.5%	133.0%	144.0%	
	8	282.7%	160.7%	421.0%	421.0%	103.3%	98.7%	101.3%	54.4%	58.4%	86.8%	66.6%	54.8%	58.5%	87.8%	81.7%	
	10	252.4%	120.1%	320.0%	320.0%	72.1%	69.1%	70.9%	35.7%	42.0%	69.9%	42.8%	36.1%	42.7%	71.3%	50.5%	
	12	228.3%	87.1%	239.6%	239.6%	46.5%	45.1%	46.2%	22.7%	32.4%	60.8%	30.6%	23.0%	33.7%	62.7%	30.6%	

Table E.1: Average OCG for MLE, MH, and all benchmark algorithms over 1,000 synthetic test cases.

SL	T	Parametric		Non-Parametric		MLE			MH								FED
		Poisson	Normal	Max	SAA	4	6	8	IC			MHR					
									4	6	8	Self-Reg	4	6	8		
90	4	0.50	0.56	0.57	0.62	0.57	0.43	0.43	0.43	0.34	0.37	0.44	0.34	0.37	0.45	0.41	
	6	0.30	0.33	0.39	0.49	0.39	0.31	0.30	0.31	0.24	0.28	0.36	0.24	0.28	0.36	0.27	
	8	0.24	0.25	0.31	0.45	0.28	0.23	0.22	0.22	0.18	0.21	0.30	0.17	0.22	0.31	0.20	
	10	0.18	0.20	0.24	0.40	0.21	0.20	0.20	0.20	0.14	0.18	0.26	0.14	0.18	0.28	0.14	
	12	0.15	0.19	0.20	0.36	0.23	0.15	0.15	0.16	0.11	0.15	0.24	0.11	0.16	0.26	0.12	
95	4	0.97	1.20	1.09	1.18	1.21	0.81	0.80	0.79	0.66	0.63	0.65	0.66	0.63	0.65	0.72	
	6	0.58	0.75	0.73	0.93	0.77	0.59	0.58	0.59	0.48	0.48	0.51	0.48	0.48	0.52	0.52	
	8	0.40	0.55	0.59	0.82	0.57	0.37	0.37	0.37	0.26	0.27	0.33	0.26	0.27	0.34	0.32	
	10	0.26	0.38	0.45	0.72	0.40	0.29	0.29	0.29	0.20	0.21	0.29	0.20	0.22	0.29	0.21	
	12	0.21	0.28	0.35	0.65	0.31	0.24	0.24	0.24	0.15	0.18	0.26	0.15	0.18	0.27	0.17	
98	4	2.32	2.98	2.51	2.58	2.98	1.91	1.90	1.88	1.60	1.52	1.49	1.60	1.52	1.49	1.68	
	6	1.37	1.98	1.59	2.00	1.98	1.42	1.40	1.39	1.20	1.16	1.15	1.20	1.16	1.15	1.25	
	8	0.87	1.49	1.25	1.74	1.49	0.84	0.81	0.80	0.54	0.50	0.52	0.54	0.50	0.52	0.66	
	10	0.49	1.09	0.89	1.47	1.09	0.57	0.56	0.56	0.41	0.37	0.41	0.41	0.37	0.41	0.44	
	12	0.36	0.82	0.68	1.33	0.82	0.39	0.38	0.38	0.24	0.22	0.29	0.24	0.23	0.30	0.26	
99	4	4.43	5.76	4.44	4.55	5.76	3.66	3.64	3.59	3.13	3.00	2.95	3.13	3.00	2.95	3.24	
	6	2.60	3.89	2.90	3.37	3.89	2.71	2.66	2.63	2.35	2.28	2.25	2.35	2.28	2.25	2.42	
	8	1.55	2.95	2.23	3.05	2.95	1.50	1.48	1.49	1.02	0.91	0.90	1.02	0.92	0.91	1.18	
	10	0.90	2.21	1.54	2.44	2.21	1.00	0.99	0.99	0.76	0.65	0.66	0.76	0.66	0.67	0.73	
	12	0.55	1.68	1.13	2.17	1.68	0.66	0.65	0.65	0.41	0.32	0.36	0.41	0.33	0.37	0.43	
99.5	4	8.50	11.10	8.49	8.31	11.10	6.98	6.97	6.88	5.98	5.74	5.69	5.98	5.74	5.69	6.04	
	6	5.07	7.60	5.31	6.28	7.60	5.27	5.22	5.14	4.67	4.53	4.50	4.67	4.53	4.50	4.72	
	8	2.82	5.77	4.03	5.11	5.77	2.84	2.80	2.80	1.98	1.79	1.76	1.98	1.75	1.76	2.16	
	10	1.68	4.36	2.71	4.07	4.36	1.94	1.92	1.91	1.46	1.26	1.24	1.46	1.26	1.24	1.34	
	12	0.97	3.34	1.93	3.57	3.34	1.25	1.22	1.21	0.80	0.60	0.61	0.80	0.61	0.61	0.75	

Table E.2: Standard deviation of OCG for MLE, MH, and all benchmark algorithms over 1,000 synthetic test cases.

SL	T	Parametric Model		Non-Parametric Model		Full
		Poisson	Normal	MaxApp	SAA	Emp. Dist.
0.90	4	74.4/11.4/14.2	62/14.3/23.7	58.2/14.9/26.9	65.4/15.1/19.5	48.4/18.2/33.4
	6	76.5/11.3/12.2	57.1/17.1/25.8	43.1/17.1/39.8	59.4/17.3/23.3	44.2/21/34.8
	8	78.4/14/7.6	53.2/20.8/26	33.9/17.8/48.3	52/19.5/28.5	42.8/23.3/33.9
	10	80/13.6/6.4	52.4/22.8/24.8	26.7/18.2/55.1	45.9/21.2/32.9	42.2/24.7/33.1
	12	81.9/13.7/4.4	51.4/24.1/24.5	19.5/18.5/62	53.6/22.1/24.3	38.4/28.7/32.9
0.95	4	77/10.6/12.4	63.7/13.2/23.1	75.6/10.9/13.5	77.4/10.5/12.1	51/17.4/31.6
	6	80.2/9.5/10.3	58.7/15.9/25.4	63.2/14.8/22	68.3/13.9/17.8	45.4/20.2/34.4
	8	81.3/11.3/7.4	56.4/16.7/26.9	55.1/16.6/28.3	62.6/15.8/21.6	43/21.9/35.1
	10	82/12.2/5.8	55.1/19.5/25.4	48.2/18.4/33.4	55.1/20/24.9	39.8/25.5/34.7
	12	84.3/12/3.7	54.5/21.3/24.2	41/20.3/38.7	53.9/24.3/21.8	36.1/29.5/34.4
0.98	4	80.3/8.7/11	66.5/10.3/23.2	89.4/6.7/3.9	89.4/6.7/3.9	58.1/15.6/26.3
	6	82/8.4/9.6	61.8/13.4/24.8	82.7/10.2/7.1	82.7/10.3/7	49.7/19.1/31.2
	8	84.7/8.6/6.7	61/12.2/26.8	77.5/12.5/10	77.8/12.3/9.9	46.3/22.3/31.4
	10	85.3/9.7/5	59.7/14.8/25.5	73.2/14.6/12.2	74.5/14.3/11.2	44.5/25.3/30.2
	12	87.1/7.9/5	57.9/17.1/25	68.7/17/14.3	70.4/17/12.6	41.2/29.6/29.2
0.99	4	81.4/7.4/11.2	66.4/10.6/23	95.3/3.1/1.6	95.3/3.1/1.6	60.1/17.4/22.5
	6	83.2/7.7/9.1	63.4/11.4/25.2	91.8/5.1/3.1	91.8/5.1/3.1	52.8/20.4/26.8
	8	86.2/7.1/6.7	61/12.9/26.1	88.5/6.8/4.7	88.5/6.8/4.7	47.8/25.1/27.1
	10	86.2/8.9/4.9	60.2/13.2/26.6	85.5/8.6/5.9	85.5/8.6/5.9	45.5/28/26.5
	12	87.6/7.4/5	59.9/14.7/25.4	83/10.1/6.9	83/10.1/6.9	42.5/31.6/25.9
0.995	4	80.4/7.4/12.2	65.6/9.8/24.6	97.4/1.6/1	97.4/1.6/1	61.2/18.3/20.5
	6	81/8.5/10.5	60.9/11.2/27.9	95.4/2.6/2	95.4/2.6/2	53.1/22.5/24.4
	8	83.6/7.2/9.2	58.9/12.8/28.3	93.2/3.9/2.9	93.2/3.9/2.9	46.5/26.8/26.7
	10	83/8.8/8.2	58.1/12.1/29.8	91.2/5.7/3.1	91.2/5.7/3.1	43.4/30.2/26.4
	12	84.3/8.2/7.5	56.7/14.1/29.2	89.6/6.7/3.7	89.6/6.7/3.7	39.6/32.6/27.8

Table E.3: Percentage of Under-Estimation/Optimality/Over-Estimation of each benchmark algorithm with service level 90%, 95%, 98%, 99%, 99.5%.

Appendix F

Generation of Δ_t

Algorithm 4 Iteration for generating Δ_t under d_t , z_t , and BI

```

1: procedure GENDELTA( $d_t, z_t, \underline{w}, \overline{w}$ )
2:   Initiate  $\Delta_t \leftarrow \{(\underline{w}), (\underline{w} + 1), \dots, (\min\{\overline{w}, d_t\})\}$ 
3:   for  $k \leftarrow 2 : z_t$  do
4:     for  $\delta_t^* \in \Delta_t$  do
5:        $\Delta_t \leftarrow \Delta_t \setminus \{\delta_t^*\}$ 
6:       if  $k < z_t$  then
7:         for  $w_{t,k}^* \leftarrow w_{t,k-1} : \min\{\overline{w}, d_t - \mathbf{1}^T \delta\}$  do
8:           if  $w_{t,k}^* \leq (d_t - \sum_{s=1}^{k-1} w_{t,s}^\delta - w_{t,k}^*) / (z_t - k) \leq \overline{w}$  then
9:              $\Delta_t \leftarrow \Delta_t \cup \{(w_1^\delta, \dots, w_{k-1}^\delta, w_k^*)\}$ 
10:          end if
11:        end for
12:      else if  $k = z_t$  and  $w_{t,z_t-1}^\delta \leq d_t - \mathbf{1}^T \delta \leq \overline{w}$  then
13:         $w_k^* \leftarrow (d_t - \mathbf{1}^T \delta)$ 
14:         $\Delta_t \leftarrow \Delta_t \cup \{(w_1^\delta, \dots, w_{z_t-1}^\delta, w_k^*)\}$ 
15:      end if
16:    end for
17:  end for
18:  Return  $\Delta_t$ .
19: end procedure

```

Appendix G

Calculation of $\tau(\delta_t)$

Algorithm 5 Recursion for generating number of unordered permutations for δ_t

```
1: procedure  $\tau(\delta)$ 
2:    $\tau \leftarrow 1, n = \dim(\delta)$ 
3:   for  $w \leftarrow \min(\delta) : \max(\delta)$  do
4:      $n_w \leftarrow \sum_{k=1}^{z_t} \mathbb{I}[w_k = w]$ 
5:      $\tau \leftarrow \tau \times \binom{n}{n_w}$ 
6:      $n \leftarrow n - n_w$ 
7:   end for
8:   Return  $\tau$ .
9: end procedure
```

Appendix H

Generation of All Feasible $\mathbf{z}_T$

Algorithm 6 Generating all feasible $\mathbf{z}_T$ under $\mathbf{d}_T$, $\tilde{z}$, and BI

```

1: procedure GENZLIST( $\mathbf{z}_T, \tilde{z}, BI$ )
2:   for  $t \leftarrow 1 : T$  do
3:      $z_{t,LB} \leftarrow \max \left\{ \underline{z}, \left\lceil \frac{d_t}{w} \right\rceil \right\}$ 
4:      $z_{t,UB} \leftarrow \min \left\{ \bar{z}, \left\lfloor \frac{d_t}{w} \right\rfloor \right\}$ 
5:   end for
6:   Initiate  $ZList \leftarrow \{(z_{1,LB}), \dots, (z_{1,UB})\}$ 
7:   for  $t \leftarrow 2 : T$  do
8:     for  $\mathbf{z}^* \in ZList$  do
9:        $ZList \leftarrow ZList \setminus \{\mathbf{z}^*\}$ 
10:      if  $t < T$  then
11:        for  $z^* \leftarrow z_{t,LB} : z_{t,UB}$  do
12:           $\xi \leftarrow \tilde{z} - (\mathbf{1}^T \mathbf{z}^* + z^*)$   $\triangleright \xi$  is sum of remaining orders
13:          if  $\xi \geq \sum_{s=t+1}^T z_{s,LB}$  and  $\xi \leq \sum_{s=t+1}^T z_{s,UB}$  then
14:             $ZList \leftarrow ZList \cup \{(\mathbf{z}^*, z^*)\}$ 
15:          end if
16:        end for
17:      else if  $t = T$  and  $z_{T,LB} \leq \tilde{z} - \mathbf{1}^T \mathbf{z}^* \leq z_{T,UB}$  then
18:         $ZList \leftarrow ZList \cup \{(\mathbf{z}^*, \tilde{z} - \mathbf{1}^T \mathbf{z}^*)\}$ 
19:      end if
20:    end for
21:  end for
22:  Return  $ZList$ .
23: end procedure

```

Appendix I

Company M's Approach

In addition to $\mathbf{x}$ and input parameter p , we introduce the following notations to mathematically describe the Company M's Approach (the M App):

- $\mu_0 = \frac{1}{52} \sum_{t=1}^{52} x_{-t}$
- SS_0 : Current safety stock
- WOS_0 : Current weeks of supply
- θ_p : p -quantile of the empirical distribution of x_t

First, $WOS_0 = SS_0/\mu_0$ is calculated by current safety stock level and mean usage over the trailing 52 week history. Then a trimmed mean is calculated by $\mu_p = \frac{1}{[52p]} \sum_{x_{-s} \leq \theta_p} x_{-s}$. Finally, the M Approach prescribes the new safety stock level (corresponding to parameter p) to be μ_p : $SS_p = \lceil WOS_0 \times \mu_p \rceil$.

Appendix J

MFC Approach

For each ANT SKU-L, the M Approach takes an input $0 < p < 1$ and calculates a “trimmed mean” by dropping the top $(1 - p)$ portion of usage history. Then the a new safety stock level (with parameter p) is set to be the product of weeks of supply by current safety stock level (with mean demand approximated by trailing usage history) and the trimmed mean (detailed description in Appendix I). Figure J.1 shows a SKU example of the M Approach at a regional warehouse: 90th-percentile recommendation suggests safety stock reduced from 30 units to 9 units; 75th-percentile suggest a further reduction to 4 units.

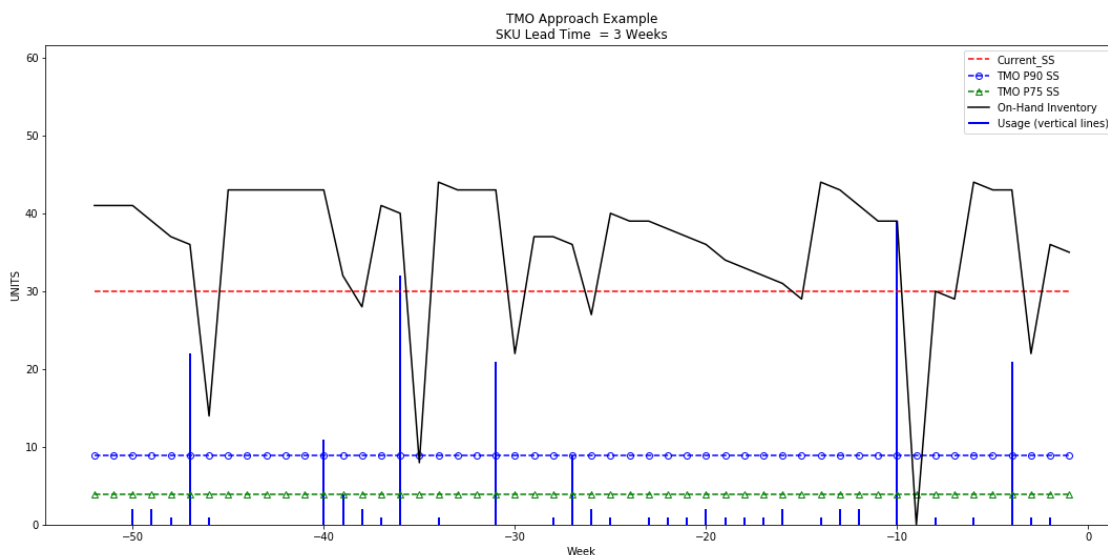


Figure J.1: Example of the M Approach with parameter $p = 0.75$, and $p = 0.9$.

First, we define a “cycle” to be the period between two successive receipts of replenishment. The minimum on-hand inventory within a cycle coincides with the arrival of the second replenishment. Assuming that the second replenishment arrives in period t_{end} , then the ending inventory level is $I(t_{end}) = ROP_{t'-L} - X_{t'-L,t'}$, where $ROP_{t'-L}$ is the reorder point $ROP_{t'-L}$ triggered in period $t' - L$, and $X_{t'-L,t'} = X_{t'-L} + X_{t'-L+1} + \dots + X_{t'-1}$ is the random variable of lead time demand between $t' - L$ to t' .

Substituting $ROP_{t'-L}$ with Equation (4.1), we have $I_{t'} = F_{t'-L,t'} + SS_{t'-L} - X_{t'-L,t'}$, where $F_{t'-L,t'}$ is the lead time demand forecast. Under a guaranteed service model, the ending inventory $I_{t'}$ should always be greater than or equal to zero, then we have $SS_{t'-L} \geq X_{t'-L,t'} - F_{t'-L,t'}$. Let G_{t_1,t_2} be the distribution of demand between t_1 and t_2 with right support $\bar{X}_{t'-L,t'}$. When forecast is unbiased, i.e. $F_{t'-L,t'} = E[X_{t'-L,t'}]$, we should set $SS_{t'-L} = G_{t'-L,t'}^{-1}(\alpha) - E[X_{t'-L,t'}]$. Moreover, let $\bar{X}(L)$ be the right support for any lead time demand between $t - L$ and t . Since we do now know the accuracy of $F_{t'-L,t'}$, and distribution $G_{t'-L,t'}$, we can set $SS_{t'-L} = \bar{X}(L)$ by relaxing the $X_{t'-L,t'}$ to $\bar{X}(L)$ and set $F_{t'-L,t'} = 0$, which guarantees that $I_{t'} \geq 0$.

Let I_{-t} be the inventory in period $-t$, we should have $x_t \geq x_{t'} = SS_{t'-L} - (X_{t'-L,t'} - F_{t'-L,t'})$, where the first inequality represents that on-hand inventory level in period t should be higher than the ending on-hand inventory in the same cycle. For unbiased demand forecast, we should expect on-hand inventory level drop below $SS_{t'-L}$. However, the minimum on-hand inventory should never be higher than $\bar{X}(L)$ since that even the most extreme cases (forecast is 0 and lead time demand realizes to be $\bar{X}(L)$) would be covered by setting $SS_t = \bar{X}(L)$. In other word, as long as the right support of lead time demand remain constant over time (no assumption on

the periodic demand distribution or correlation), the minimum inventory level can be reduced to $\bar{X}(L)$ while guaranteeing 100% service. Let $y_{-t} = \sum_{s=-t}^{-t+L-1} x_s$ be the lead time demand from period $-t$, we estimate $\bar{X}(L)$ by $y^{MFC} = \max\{y_{-t}\}$, and calculate the gap by

$$\delta^{MFC} = \max(\min I_{-t} - y^{MFC}, 0) \quad (\text{J.1})$$

Appendix K

Simulation-based Numerical Test on Frederic Facility

For each SKU, we first construct the empirical distribution using the SKU's usage history over the 52-week history. Then a large number of $M(= 5,000)$ independent lead time demand is simulated, each of which is the sum of L independent draws from the empirical distribution (assuming the underlying weekly demand are i.i.d. distributions after filtering out SKUs with high autocorrelations and non-trivial long-term trend). We compare the reorder point in E1, which is calculated by

$$ROP^{Approach} = F(L) + SS^{Approach}$$

For the “average performance,” we approximate $F(L)$ by average historical demand, μ_0 , and calculate $SL - AC$ of an approach as

$$SL-AC^{Approach} = \frac{1}{M} \sum_{m=1}^M \mathbb{I} \left[L\mu_0 + SS^{Approach} \geq X_1^{(m)} + \dots + X_L^{(m)} \right]$$

For the lowerbound, we replace $F(L)$ with 0 to imitate the case where forecast is inaccurate (to the extreme extent), forcing all lead time demand to be fulfilled by safety stock. Formally, SL-LB of an approach is calculated by

$$SL-LB^{Approach} = \frac{1}{M} \sum_{m=1}^M \mathbb{I} \left[SS^{Approach} \geq X_1^{(m)} + \dots + X_L^{(m)} \right]$$

For Reported Inventory Reduction, the dollar value for each approach is calculated

by

- $RS^{M-p} = C_{SKU} \max\{SS_0 - SS_p^M, 0\}$
- $RS^{Normal} = C_{SKU} \max\{SS_0 - \Phi^{-1}(p; \mu_0, \sigma_0), 0\}$
- $RS^{MFC} = C_{SKU} \delta^{MFC}$

where $SS_p^{MApp} = \lceil WOS_0 \times \mu_p \rceil$, $\Phi^{-1}(\cdot; \mu_0, \sigma_0)$ is the inverse normal function with mean μ_0 and standard deviation σ_0 (direct estimation from the history), δ^{MFC} is calculated by Equation (4.2), and C_{SKU} is the standard cost of the SKU.

Bibliography

- Akcay, A., Biller, B., and Tayur, S. (2011). Improved inventory targets in the presence of limited historical demand data. *Manufacturing & Service Operations Management*, 13(3):297–309.
- Azoury, K. S. (1985). Bayes solution to dynamic inventory models under unknown demand distribution. *Management science*, 31(9):1150–1160.
- Bélisle, C. J., Romeijn, H. E., and Smith, R. L. (1993). Hit-and-run algorithms for generating multivariate distributions. *Mathematics of Operations Research*, 18(2):255–266.
- Bookbinder, J. H. and Lordahl, A. E. (1989). Estimation of Inventory Re-Order Levels Using the Bootstrap Statistical Procedure. *IIE Transactions*, 21(4):302–312.
- Chib, S. and Greenberg, E. (1995). Understanding the Metropolis-Hastings Algorithm. *The American Statistician*, 49(4):327–335.
- Feeney, G. J. and Sherbrooke, C. C. (1966). The (s- 1, s) inventory policy under compound poisson demand. *Management Science*, 12(5):391–411.
- Gallego, G. and Moon, I. (1993). The distribution free newsboy problem: review and extensions. *Journal of the Operational Research Society*, 44(8):825–834.
- Gallego, G., Ryan, J. K., and Simchi-Levi, D. (2001). Minimax analysis for finite-horizon inventory models. *Iie Transactions*, 33(10):861–874.
- Geman, S. and Geman, D. (1987). Stochastic relaxation, gibbs distributions, and the bayesian restoration of images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 6(6):721–741.
- Godfrey, G. A. and Powell, W. B. (2001). An Adaptive, Distribution-Free Algorithm for the Newsvendor Problem with Censored Demands, with Applications to Inventory and Distribution. *Management Science*, 47(8):1101–1112.
- Hastings, W. K. (1970). Monte Carlo Sampling Methods Using Markov Chains and Their Applications. *Biometrika*, 57(1):97–109.
- Hayes, R. H. (1969). Statistical estimation problems in inventory control. *Management Science*, 15(11):686–701.

- Hopp, W. J. and Spearman, M. L. (2004). To pull or not to pull: what is the question? *Manufacturing & service operations management*, 6(2):133–148.
- Huh, W. T., Levi, R., Rusmevichientong, P., and Orlin, J. B. (2011). Adaptive Data-Driven Inventory Control with Censored Demand Based on Kaplan-Meier Estimator. *Operations Research*, 59(4):929–941.
- Huh, W. T. and Rusmevichientong, P. (2009). A Nonparametric Asymptotic Analysis of Inventory Planning with Censored Demand. *Mathematics of Operations Research*, 34(1):103–123.
- Iglehart, D. L. (1964). The dynamic inventory problem with unknown demand distribution. *Management Science*, 10(3):429–440.
- Kleywegt, A. J., Shapiro, A., and Homem-de Mello, T. (2002). The sample average approximation method for stochastic discrete optimization. *SIAM Journal on Optimization*, 12(2):479–502.
- Kottas, J. F. and Lau, H.-S. (1980). The use of versatile distribution families in some stochastic inventory calculations. *Journal of the Operational Research Society*, 31(5):393–403.
- Levi, R., Perakis, G., and Uichanco, J. (2015). The Data-driven Newsvendor Problem: New Bounds and Insights. *Operations Research*, 63(6):1294–1306.
- Levi, R., Roundy, R., and Shmoys, D. B. (2007). Provably Near-Optimal Sampling-Based Policies for Stochastic Inventory Control Models. *Mathematics of Operations Research*, 32(4):821–839.
- Mengersen, K. L., Tweedie, R. L., et al. (1996). Rates of convergence of the hastings and metropolis algorithms. *The annals of Statistics*, 24(1):101–121.
- Metropolis, N., Rosenbluth, A. W., Rosenbluth, M. N., Teller, A. H., and Teller, E. (1953). Equation of State Calculations by Fast Computing Machines. *The Journal of Chemical Physics*, 21(6):1087–1092.
- Moon, I. and Gallego, G. (1994). Distribution free procedures for some inventory models. *Journal of the Operational research Society*, 45(6):651–658.
- Müller, P. (1991). *A generic approach to posterior integration and Gibbs sampling*. Purdue University, Department of Statistics.
- Nummelin, E. (1984). *General Irreducible Markov Chains and Non-Negative Operators*. Cambridge Tracts in Mathematics 83. Cambridge University Press.
- Perakis, G. and Roels, G. (2008). Regret in the newsvendor model with partial information. *Operations Research*, 56(1):188–203.

- Roberts, G. O. and Smith, A. F. (1994). Simple conditions for the convergence of the gibbs sampler and metropolis-hastings algorithms. *Stochastic processes and their applications*, 49(2):207–216.
- Saghafian, S. and Tomlin, B. (2016). The Newsvendor under Demand Ambiguity: Combining Data with Moment and Tail Information. *Operations Research*, 64(1):167–185.
- Scarf, H. (1959). Bayes solutions of the statistical inventory problem. *The annals of mathematical statistics*, 30(2):490–508.
- Scarf, H. E. (1957). A min-max solution of an inventory problem. Technical report, RAND CORP SANTA MONICA CALIF.
- Sherbrooke, C. C. (1968). Metric: A multi-echelon technique for recoverable item control. *Operations research*, 16(1):122–141.
- Tierney, L. (1994). Markov chains for exploring posterior distributions. *the Annals of Statistics*, 22(4):1701–1728.

CURRICULUM VITAE

