



Manufacturing & Service Operations Management

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

Data Analytics in Operations Management: A Review

Velibor V. Mišić, Georgia Perakis

To cite this article:

Velibor V. Mišić, Georgia Perakis (2020) Data Analytics in Operations Management: A Review. Manufacturing & Service Operations Management 22(1):158-169. <https://doi.org/10.1287/msom.2019.0805>

Full terms and conditions of use: <https://pubsonline.informs.org/Publications/Librarians-Portal/PubsOnLine-Terms-and-Conditions>

This article may be used only for the purposes of research, teaching, and/or private study. Commercial use or systematic downloading (by robots or other automatic processes) is prohibited without explicit Publisher approval, unless otherwise noted. For more information, contact permissions@informs.org.

The Publisher does not warrant or guarantee the article's accuracy, completeness, merchantability, fitness for a particular purpose, or non-infringement. Descriptions of, or references to, products or publications, or inclusion of an advertisement in this article, neither constitutes nor implies a guarantee, endorsement, or support of claims made of that product, publication, or service.

Copyright © 2019, INFORMS

Please scroll down for article—it is on subsequent pages



With 12,500 members from nearly 90 countries, INFORMS is the largest international association of operations research (O.R.) and analytics professionals and students. INFORMS provides unique networking and learning opportunities for individual professionals, and organizations of all types and sizes, to better understand and use O.R. and analytics tools and methods to transform strategic visions and achieve better outcomes. For more information on INFORMS, its publications, membership, or meetings visit <http://www.informs.org>

20th Anniversary Invited Article

Data Analytics in Operations Management: A Review

Velibor V. Mišić,^a Georgia Perakis^b

^a Anderson School of Management, University of California, Los Angeles, Los Angeles, California 90095; ^b Sloan School of Management, Massachusetts Institute of Technology, Cambridge, Massachusetts 02139

Contact: velibor.misic@anderson.ucla.edu,  <https://orcid.org/0000-0002-8952-5617> (VVM); georgiap@mit.edu,

 <https://orcid.org/0000-0002-0888-9030> (GP)

Received: March 28, 2019

Revised: March 29, 2019

Accepted: March 29, 2019

Published Online in Articles in Advance:
September 19, 2019

<https://doi.org/10.1287/msom.2019.0805>

Copyright: © 2019 INFORMS

Abstract. Research in operations management has traditionally focused on models for understanding, mostly at a strategic level, how firms should operate. Spurred by the growing availability of data and recent advances in machine learning and optimization methodologies, there has been an increasing application of data analytics to problems in operations management. In this paper, we review recent applications of data analytics to operations management in three major areas—supply chain management, revenue management, and healthcare operations—and highlight some exciting directions for the future.

History: This paper has been accepted for the *Manufacturing & Service Operations Management* 20th Anniversary Special Issue.

Funding: G. Perakis received financial support from the National Science Foundation [Grant CMMI-1563343].

Keywords: data analytics • machine learning • operations management

1. Introduction

Historically, research in operations management (OM) has focused on models. These models, based on micro-economic theory, game theory, optimization, and stochastic models, have been mostly used to generate strategic insights about how firms should operate. One of the reasons for the prevalence of this model-based approach in the past has been the relative scarcity of data, coupled with limitations in computing power.

Today, there has been a shift in research in OM. This shift has been driven primarily by the increasing availability of data. In various domains—such as retail, healthcare, and many more—richer data are becoming available that are more voluminous and more granular than ever before. At the same time, the increasing abundance of data has been accompanied by methodological advances in a number of fields. The field of machine learning, which exists at the intersection of computer science and statistics, has created methods that allow one to obtain high-quality predictive models for high-dimensional data, such as L1 regularized regression [also known as *least absolute shrinkage and selection operator* (LASSO) regression; Tibshirani (1996)] and random forests (Breiman 2001). The field of optimization has similarly advanced. Numerous scientific innovations in optimization modeling, such as robust optimization (Bertsimas et al. 2011), inverse optimization (Ahuja and Orlin 2001), and improved integer-optimization formulation techniques (Vielma 2015), have extended both the scope of what

optimization can be applied to and the scale at which it can be applied. In both machine learning and optimization, researchers have benefited from the availability of high-quality software for estimating machine-learning models and for solving large-scale linear, conic, and mixed-integer optimization problems.

These advances have led to the development of the burgeoning field of *data analytics* (or *analytics* for short). The field of analytics can most concisely be described as using *data* to create *models* leading to *decisions* that create *value*. In this paper, in honor of the 20th anniversary of *Manufacturing & Services Operations Management*, we highlight recent work that applies the analytics approach to problems in OM. We divide our review along three major application areas: supply chain management (Section 2), where we cover location, omnichannel, and inventory decisions; revenue management (Section 3), where we cover choice modeling and assortment optimization, pricing and promotion planning, and personalized revenue management; and healthcare (Section 4), where we cover applications at the policy, hospital, and patient levels. We conclude in Section 5 with a discussion of some future directions—such as causal inference, interpretability, and “small-data” methods—that we believe will be increasingly important in the future of analytics in OM.

Because of the page limitations imposed by this Anniversary Issue, our review is necessarily brief and only covers a small set of representative examples in

each application area. As a result, there are many other excellent papers that apply analytics to OM that we were unable to include in this review; we apologize in advance to those authors whose papers we have not cited.

2. Analytics in Supply Chain Management

Many major problems in supply chain management are being reexamined under the lens of analytics; in this section, we focus on location and omnichannel decisions (Section 2.1), as well as inventory decisions (Section 2.2).

2.1. Location and Omnichannel Operations

In this subsection, we discuss the application of analytics to omnichannel and, more generally, location problems. The term *omnichannel* refers to the integration of an e-commerce channel and a network of brick-and-mortar stores. This integration allows a retailer to do cross-channel fulfillment—that is, fulfill online orders in any store location, leading to interesting location and inventory-management problems. Glaeser et al. (2018) consider a location problem faced by a real “buy online, pickup in store” retailer, which fulfills online orders through delivery trucks parked at easily accessible locations (e.g., at schools or parking lots). The retailer needs to decide at which locations and times to position its trucks in order to maximize profit. To solve this problem, the paper first builds a random forest model (Breiman 2001) to predict demand at a given location at a given time, using a diverse set of independent variables such as demographic attributes of the location (total population, population with a postsecondary degree, median income, and so on), the retailer’s operational attributes (such as whether the retailer offers home delivery in this location), and other location attributes (e.g., number of nearby competing businesses). Then the paper uses fixed-effects regression to account for cannibalization effects. Using the retailer’s data, the paper shows how a heuristic approach based on greedy construction and interchange ideas together with the combined random forest and fixed-effects model, leads to improvements in revenue of up to 36%.

Acimovic and Graves (2014) study how to manage the omnichannel operations of a large online retailer. They address the fundamental questions of what is the best way to fulfill each customer’s order in order to minimize average outbound shipping costs. The paper develops a heuristic that makes fulfillment decisions by minimizing the immediate outbound shipping cost, plus an estimate of future expected outbound shipping costs. These estimates are derived from the dual values of a transportation linear program. Through experiments on industry data, the model captures 36% of the opportunity gap assuming

clairvoyance, leading to reductions in outbound shipping costs on the order of 1%. In the same spirit, Avrahami et al. (2014) discuss their collaboration with the distribution organization within the Yedioth Group in Israel to improve how Yedioth distributes print magazines and newspapers. This paper uses real-time information on the newspaper sales at the different retail outlets to enable pooling of inventory in the distribution network. The methods developed were implemented at Yedioth and generated substantial cost savings as a result of both a reduction in magazine-production levels and a reduction in return levels.

In a different domain, He et al. (2017) consider the problem of how to design service regions for an electric-vehicle-sharing service. A number of car-sharing firms, such as Car2Go, DriveNow, and Autolib, offer electric vehicles to customers. Providing such a service in an urban area requires the firm to decide the regions in which the service will be offered—that is, in which regions of the urban area are customers allowed to pick up and drop off a vehicle; this, in turn, informs major investment decisions (how large the fleet needs to be and where charging stations should be located). The main challenge in this setting comes from the fact that there is uncertainty in customer adoption, which depends on which regions are covered by the service. To solve this problem, the paper formulates an integer programming problem where customer adoption is represented through a utility model; because of the limited data from which such a utility model can be calibrated, the paper uses distributionally robust optimization (see, e.g., Delage and Ye 2010) to account for the uncertainty in customer adoption. Using data from Car2Go, the paper applies the approach to designing the service region in San Diego and shows how the approach leads to higher revenues than several simpler, managerially motivated heuristic approaches to service region design.

2.2. Inventory Management

Ban and Rudin (2019) consider a data-driven approach to inventory management. The context that the paper studies is when one has observations of demand together with features that may be predictive of demand, such as weather forecasts or economic indicators like the consumer price index. To make inventory decisions in this setting, one might consider building a demand distribution that is feature dependent and then finding the optimal order quantity for the distribution corresponding to a given realization of the features. Instead, the paper by Ban and Rudin (2019) proposes two alternate approaches. The first, based on empirical risk minimization, involves finding the order quantity by solving a single problem, where the decision variables is the decision rule

that maps the features to an order quantity, and the objective is to minimize a sample-based estimate of the cost. The second approach involves using kernel regression to model the conditional demand distribution and applying a sorting algorithm to determine the optimal order quantity. The paper derives bounds on the out-of-sample performance of both approaches and applies the approaches to the problem of nurse staffing in a hospital emergency room using data from a large teaching hospital in the United Kingdom. The paper finds that the proposed approaches outperform the best-practice benchmark by up to 24% in out-of-sample cost.

In a similar vein, Bertsimas and Kallus (2019) develop a general framework for solving a collection of sample-average approximation (SAA) problems, where each SAA problem has some contextual information. For example, in an inventory setting, a firm might be selling different products, which are described by vectors of attributes, for which the firm has historical observations of demand; the firm would then like to determine an order quantity for a (potentially new) product. A naive approach to this problem would be to simply pool all the data together and ignore the contextual information. Instead, Bertsimas and Kallus (2019) propose building a machine learning model first to predict the uncertain quantity as a function of contextual information; in our inventory example, this would amount to predicting demand as a function of the product attributes. Then, instead of solving the SAA problem using all the data on the uncertain parameter, one uses the machine-learning model to obtain a context-specific conditional distribution and then solves the SAA problem with respect to this conditional distribution. In the inventory setting, consider a regression tree (Breiman et al. 1984) for predicting demand for a given product attribute vector. If one runs the product's attribute vector down the regression tree, one will obtain a point prediction, but by considering the historical demand observations contained in the corresponding leaf, one also obtains an estimate of the conditional distribution. The paper develops theoretical guarantees for this general procedure for specific types of machine-learning models, proving asymptotic optimality of the procedure as the number of observations grows. The paper also applies the approach to a real distribution problem for a media company that needs to manage inventories of different products (specifically, movie CDs, DVDs, and Blu-ray discs) across different retail locations. The paper builds machine-learning models of demand as a function of the store location, properties of the movie (such as its genre, rating on Rotten Tomatoes, box office revenue, and so on), and other large-scale information (such as localized Google search queries for the movie). The paper shows how the approach is able

to close the cost gap between the naive approach (which does not consider contextual data) and the perfect foresight approach (which knows demands before they are realized) by 88%.

Ban et al. (2018) consider the problem of dynamic procurement. In this problem, one has to decide how much of a product to order over a finite time horizon from different sources that have different lead times and different costs so as to meet uncertain demand over the horizon. The problem setting studied in the paper is motivated by a practical problem faced by Zara, one of the world's largest fast-fashion retailers. The paper proposes an approach that involves using linear and regularized linear regression to predict the demand for a given product using the historical trajectories of other products. Then the paper formulates a multistage stochastic programming problem where the scenario tree is generated by using the previously estimated predictive model. Using real data from Zara, the paper shows how the existing approach, which ignores covariate information, can lead to costs that are 6%–15% higher than the proposed approach.

At the core of many inventory problems—and OM problems in general—lies the problem of demand estimation. The results of this estimation can then be used, among others, to address pricing, distribution, and inventory-management questions. Baardman et al. (2018a) develop a data-driven approach for predicting demand for new products. The paper develops a machine learning algorithm that identifies comparable products in order to predict demand for the new product. The paper demonstrates both analytically and empirically, through a collaboration with Johnson & Johnson, the benefits of this approach.

3. Analytics in Revenue Management

Analytics is having a major impact on how firms make decisions about how to sell their products to customers; in this section, we discuss applications in choice modeling and assortment optimization (Section 3.1), pricing and promotion planning (Section 3.2), and personalized revenue management (Section 3.3).

3.1. Choice Modeling and Assortment Optimization

Assortment optimization refers to the problem of deciding which products to offer to a set of customers so as to maximize the revenue earned when customers make purchase decisions from the offered products. A major complicating factor in assortment optimization is the presence of substitution behavior, where customers may arrive with a particular product in mind but purchase a different product if that first product is not available.

The paper by Farias et al. (2013) proposed a non-parametric model for representing customer choice

behavior by modeling the customer population as a probability distribution over rankings of the products. Such a model is able to represent any choice model based on random utility maximization, which includes the majority of popular choice models (e.g., the multinomial logit [MNL] model, the nested logit model, and the latent-class MNL model). The paper shows, using data from a major automobile manufacturer, how the proposed approach can provide more accurate predictions than traditional parametric models such as the MNL model. In later work, a number of papers have considered how to efficiently find assortments that maximize expected revenue under such a model. From a tractability standpoint, Aouad et al. (2018) showed that the assortment problem is NP-hard even to approximate (APX-hard) and developed an approximation algorithm; Feldman and Topaloglu (2019) later provided an approximation algorithm with an improved guarantee. For specific classes of ranking-based models, Aouad et al. (2015) developed efficient approaches based on dynamic programming for solving the assortment-optimization problem exactly. In a different direction, the paper by Bertsimas and Mišić (2019) proposed a new mixed-integer optimization formulation for the problem as well as a specialized solution approach, based on Benders decomposition, for solving large-scale instances efficiently. Using a product line design problem based on a real conjoint analysis data set, the paper shows how a large-scale product line design problem that required 1 week of computation time to solve in earlier work is readily solved by the proposed approach in 10 minutes.

In addition to the ranking-based model, the paper by Blanchet et al. (2016) recently introduced a Markov chain-based choice model. In this model, substitution from one product to another is modeled as a state transition in a Markov chain. The paper shows that under some reasonable assumptions, the choice probabilities computed by the Markov chain-based model are a good approximation of the true choice probabilities for any random utility-based choice model and are exact if the underlying model is a generalized attraction model, of which the MNL model is a special case. In addition to the theoretical bounds, the authors conduct numerical experiments and observe that the average maximum relative error of the choice probabilities of their model with respect to the true probabilities for any offer set is less than 3%, where the average is taken over different offer sets.

3.2. Pricing and Promotion Planning

An exciting example of how analytics is being applied in pricing is the paper by Ferreira et al. (2015), which studies the pricing of Rue La La's weekly sales. Rue La La is an online fashion retailer that sells

designer apparel items (called *styles*) at significantly discounted prices in weekly limited-time sales (called *events*). The paper develops a strategy for optimally pricing styles in an event that combines machine learning and optimization. The machine-learning component of the approach involves developing a bagged regression tree model (Breiman 1996) for predicting the demand of a style based on the price of the style and the average price of other styles in the event. The optimization component involves formulating an integer optimization problem where the goal is to find the prices of the styles that maximize the total revenue from all styles, which is just the product of the price of the style and the demand predicted by the bagged regression tree model. This optimization problem is, in general, very difficult. However, if one fixes the average price of all styles to a given value, then the problem simplifies, and one can solve a small number of smaller integer-optimization problems to find the optimal collection of prices. The paper reports on a live pilot experiment to test the methodology at Rue La La, which resulted in an increase in revenues of about 10%.

Another example where analytics has made a difference in the field of pricing relates to promotion planning. Cohen et al. (2017a, b), as well as Cohen-Hillel et al. (2019a, b), consider how analytics applies to price-promotion planning first for a single item and, subsequently, for multiple items. Price-promotion planning is concerned with when and how deeply to promote each item over a time horizon. The problem is challenging because sales of an item in a given time period are affected not only by the price of that item in that period but also by what the prices of that item were before then and what the prices of other items are. For example, a promotion on a coffee stock-keeping unit (SKU) today may result in a large increase in sales if it has not been recently promoted but may only result in a modest change if the coffee SKU has been promoted heavily in the preceding weeks (because customers may, for example, have stockpiled the item during earlier promotion periods). The paper by Cohen et al. (2017b) considers the single-item problem and introduces a general mixed-integer nonlinear optimization formulation of the problem, where one first learns a demand model from available data and then uses that demand model in the objective of the optimization model. The objective in this body of work corresponds to profit, and the constraints reflect common business rules imposed by retailers in practice (e.g., limits on the number of successive promotions over any fixed window). Although the optimization problem is in general theoretically intractable, Cohen et al. (2017b) show how an approximate formulation, based on linearizing the nonlinear objective, leads to an efficiently solvable mixed-integer linear problem. The paper shows how the suboptimality gap

of the promotion-price schedule derived from the approximation can be bounded for different classes of demand functions. This approach was also later extended to multi-item price-promotion planning problem in Cohen et al. (2017a).

More recently, Cohen-Hillel et al. (2019a, b) have considered a demand model that only considers a limited set of pricing features; in particular, one considers the price of the item in the current time period and the price in the previous period, as well as the minimum price within a bounded set of past periods. These papers consider the pricing problem under this model, which is referred to as the *bounded memory peak-end demand model*. These papers show how one can exploit the structure arising from these price features to introduce a compact dynamic programming formulation of the pricing problem, allowing the problem to be efficiently solved. In the multi-item version of the problem, the papers show that the problem is NP-hard in the strong sense but establish a polynomial-time approximation scheme. Furthermore, Cohen-Hillel et al. (2019b) show that even if the true underlying demand model does not follow the bounded memory peak-end demand model but rather follows some of the more traditional models in the literature, the bounded memory peak-end demand model still performs very well; the paper provides a theoretical bound for this result, which is tested with data. Using data from a grocery retailer, these papers show how the optimization approaches introduced can lead to profit increases of between 3% and 10%.

The above-mentioned papers deal with price promotions, which constitute one type of promotion. Retailers have access to many other methods of promotion (termed *promotion vehicles*), such as flyers, coupons, and announcements on social media websites (e.g., Facebook). The paper by Baardman et al. (2018c) considers how one should schedule such promotion vehicles so as to maximize profit. The paper shows that the problem is APX-hard and subsequently introduces a nonlinear integer optimization formulation for scheduling promotion vehicles under various business constraints (such as limiting the promotion vehicles to use each time period, among others). The paper considers a greedy approach as well as a linear mixed-integer optimization approximation that is more accurate than the greedy method. Using real data, the paper illustrates the quality of the two approximation methods as well as the potential profit improvement, which is on the order of 2%–9%.

In the omnichannel retail setting, consumers can compare store prices with online prices, giving rise to challenging pricing problems. Harsha et al. (2019) study the price-optimization aspect of omnichannel operations. Because of the presence of cross-channel

interactions in demand and supply, the paper proposes two pricing policies that are based on the idea of *partitions* to the store inventory, which approximate how it will be used by the two different channels. These two policies are, respectively, based on solving either a deterministic or a robust mixed-integer optimization problem that incorporates several business rules as constraints. Through an analysis of a pilot implementation at a large U.S. retailer, the paper estimates that the methodology leads to a 13.7% increase in clearance-period revenue.

3.3. Personalized Revenue Management

Another major area of research is in personalized revenue management. Retailers have become increasingly interested in personalizing their products and services, including promotions. *Personalization* refers to the ability to tailor decisions to the level of individuals. In revenue management, this entails offering different prices and changing the product assortment, as well as other decisions such as promotions and product bundling.

Many papers have recently considered the problem of personalized pricing. Chen et al. (2015) consider a setting in which a firm sells a single product to a customer population, where each customer chooses to buy the product according to a logit model that depends on customer covariates and the price, and the firm has access to historical transaction records, where each record contains the covariates of the customer, the price at which the product was offered, and a binary indicator of whether the customer bought the product or not. The paper proposes a simple approach to personalized pricing: estimate the parameters of the logistic regression model using regularized maximum-likelihood estimation and then set the price for future customers so as to maximize the predicted expected revenue (price multiplied by purchase probability) under the estimated parameters. The paper bounds the suboptimality gap between the expected revenue of this policy and that of an oracle pricing policy that knows the customer choice model perfectly. Using airline sales data, the paper shows how the proposed method provides an improvement of over 3% in revenue over a pricing policy that offers all customers the same price.

Ban and Keskin (2017) develop an algorithm for dynamic pricing of a single product, where arriving customers are described by a feature vector, and the demand in each period is given by a model that depends on the price and only a subset of the features; the demand model allows for feature-dependent price sensitivity. In this setting, the paper proves a lower bound on the regret of any dynamic pricing algorithm in terms of the number of relevant features and develops algorithms with near-optimal regret. The paper

demonstrates, using data from an online automobile loan company, how the proposed algorithms accrue more revenue than the company's current practice, as well as other recently proposed dynamic-pricing algorithms.

Personalization of product offerings has also led to the need to develop new personalized demand models. Unfortunately, in practice (and especially in the retail space), the data available are not necessarily at the level of granularity needed to develop meaningful demand models at the personalized level. A potential path to developing personalized demand models is to use data on customers from social media, which can directly inform a retailer on how one customer's purchases affect the purchases of a different customer. However, for most retailers, acquiring these data is costly and poses challenges in terms of privacy concerns. Baardman et al. (2018b) focus on the problem of detecting customer relationships from transactional data and using them to offer targeted price promotions to the right customers at the right time. The paper develops a novel demand model that distinguishes between a base purchase probability, capturing factors such as price and seasonality, and a customer-trend probability, capturing customer-to-customer trend effects. The estimation procedure is based on regularized bounded-variables least squares combined with the method of instrumental variables, a commonly used method in econometrics/causal inference (Angrist and Pischke 2008). The resulting customer trend estimates subsequently feed into a dynamic promotion targeting optimization problem, formulated as a nonlinear mixed-integer optimization problem. Although the problem is NP-hard, the paper also proposes an adaptive greedy algorithm and proves that the customer-to-customer trend estimates are statistically consistent, but also that the adaptive greedy algorithm is provably good. Using actual fashion data from a client of Oracle, the paper shows that the demand model reduces the prediction error by 11% on average, and the corresponding optimal policy increases profits by 3%–11%.

Besides these papers, there has been other recent research in the area of dynamic pricing with learning. Javanmard and Nazerzadeh (2019) consider a single-product setting where, in each period, there is only one customer that purchases the product if the price is below that customer's valuation and propose a regularized maximum-likelihood policy that achieves near-optimal expected regret. Cohen et al. (2016) consider a firm selling a sequence of arriving products, and each product is described by a feature vector. Each product is sold if its price is lower than the market value of the product, which is assumed to be linear in the features of the product. As more and more products arrive and the firm observes how

the market responds to the product (buy/no buy decisions), the firm can refine a polyhedron representing the uncertainty in the parameters of the market's valuation model. The paper proposes a policy based on approximating this polyhedron with an ellipsoid and setting prices in a way that balances the revenue gained against the reduction in uncertainty and show that this policy achieves a worst-case regret that scales favorably in the number of features and the time horizon.

There has also been much interest in personalized assortment planning. In addition to personalized pricing, Chen et al. (2015) develop similar theoretical guarantees as discussed earlier for the problem of assortment personalization based on a finite sample of data. Golrezaei et al. (2014) consider dynamic-assortment personalization when there is a limited inventory of products. The difficulty in the problem arises from the limited inventory of products—given a customer now, we may choose to offer him or her a particular product as part of an assortment, but if the customer chooses to purchase that product, the inventory of that product will be depleted, and we may run out of that product for later customers. The paper proposes a heuristic, called *inventory balancing*, that attains significantly improved revenues over simpler benchmarks. At the same time, there is also a growing literature that studies assortment personalization in a dynamic setting, where one must learn the underlying choice model while making assortment decisions. For example, Bernstein et al. (2019) consider a bandit approach to assortment personalization based on dynamic clustering. This approach assumes that there exists a clustering of customer *profiles* (vectors of attributes); however, the number of clusters and the clustering (mapping of profiles to clusters) are unknown. The approach involves modeling the distribution over clusterings by using the Dirichlet process mixture model, which is updated in a Bayesian fashion with each transaction.

Personalization has also been applied to other types of revenue-management decisions. Ettl et al. (2019) consider the problem of making personalized price and product-bundle recommendations over a finite time horizon to individual customers so as to maximize revenue with limited inventory. The paper develops a number of approximation algorithms related to the inventory-balancing heuristic (Golrezaei et al. 2014) and Lagrangian relaxation methods for weakly coupled stochastic dynamic programs (Hawkins 2003). Using two different data sets, one in airline ticket sales and one from an online retailer, the paper shows how the proposed approach can improve revenues over pricing and bundle strategies used in practice by 2%–7%.

4. Analytics in Healthcare Operations

Machine-learning and big-data methodologies are beginning to also have a significant impact in healthcare operations. We divide our discussion in this part of the paper to problems at the policy level (Section 4.1), the hospital level (Section 4.2), and the patient level (Section 4.3).

4.1. Policy-Level Problems

A number of papers have considered the application of analytics to healthcare policy problems. The paper by Aswani et al. (2019) studies the Medicare Shared Savings Program (MSSP). The MSSP was introduced by Medicare to combat rising healthcare costs: Medicare providers that enroll in MSSP and that reduce their costs below a certain financial benchmark receive bonus payments from Medicare. However, few Medicare providers have been able to successfully reduce their costs in order to receive the bonus payments because of the investment required. The paper by Aswani et al. (2019) models the interaction between Medicare and each provider as a principal-agent problem, where Medicare sets the financial benchmark for each provider, and each provider decides what investment to make to reduce costs. The authors use data from Medicare, as well as an estimation methodology based on inverse optimization (Ahuja and Orlin 2001), to learn the parameters of their principal-agent model. Using this estimated model, they propose an alternate contract where the payment depends on both reducing cost below a financial benchmark and the amount invested and show that this alternate contract can increase Medicare savings by up to 40%.

Bertsimas et al. (2013) consider the design of a dynamic allocation policy for deceased donor kidneys to waiting-list patients. In the design of an allocation policy, a key objective is efficiency—that is, the benefit garnered from allocating (transplanting) a kidney to a patient. However, another objective is fairness—that is, making sure that patients with certain characteristics are not systematically underserved by the allocation policy. The paper proposes formulating the problem as a linear optimization problem where the objective is to maximize efficiency (the aggregate match quality) subject to a fairness constraint and uses linear regression to derive a scoring rule from the solution of this linear optimization problem to predict the fairness-adjusted quality of a match using attributes of both the patient and the donor kidney. The paper shows that compared with the existing allocation policy at the time, the proposed policy can satisfy the same fairness criteria while resulting in an 8% increase in aggregate quality-adjusted life years.

Gupta et al. (2017) consider the problem of targeting individuals for a treatment or intervention.

Clinical research studies often establish that a treatment has some sort of aggregate benefit in a large population. However, there may be significant heterogeneity in the treatment effect across different subgroups of that population. For example, an individual who is very sick may benefit less from a treatment than a less sick individual. A natural question, then, is how should one allocate the intervention to individuals in a population to maximize the aggregate benefit, subject to a limit on how many individuals may be treated (i.e., it is not possible to simply treat everyone) and a lack of knowledge of individual-level treatment effects? The paper formulates this problem as a type of robust optimization problem where the uncertainty set represents all heterogeneous treatment effects that are consistent with the aggregate characteristics of the research study. The paper uses data on a case management intervention for reducing emergency department utilization by adult Medicaid patients and demonstrates settings under which the proposed robust optimization method provides a benefit over scoring rules typically used by medical practitioners.

In a different direction, the paper by Bertsimas et al. (2016) considers the design of combination chemotherapy clinical trials for gastric cancer. The authors of the paper constructed a large database of gastric cancer clinical trials and used this database to build two predictive models: one model to predict the median overall survival of patients in a trial based on the characteristics of the trial patients and the dosages of different drugs and another model to predict whether the trial will exhibit an unacceptably high fraction of patients with severe toxicities. Using these predictive models, the paper then formulates a mixed-integer optimization model for deciding the next combination of drugs to test for a given cohort of patients so as to maximize the predicted median overall survival subject to limits on the predicted toxicity. The approach is evaluated by using a method based on simulation and one based on matching a proposed clinical trial with the one most similar to it in the data; both evaluation methods suggest that the proposed approach can improve the overall efficacy of regimens that are chosen for testing in phase III clinical trials.

4.2. Hospital-Level Problems

At the hospital level, the paper by Rath et al. (2017) studies the problem of staffing and scheduling surgeries at the Ronald Reagan Medical Center (RRMC) at the University of California, Los Angeles, one of the largest medical centers in the United States. Patients who undergo surgery at the RRMC require three different types of resources: an anesthesiologist, a surgeon, and an operating room. Prior to this research, anesthesiologists at RRMC would manually schedule surgeries, resulting in high operating costs. The paper

presents a methodology, based on a two-stage robust mixed-integer optimization problem, that accounts for the uncertainty in surgery duration. The proposed method was implemented and is currently in use, saving RRMC approximately \$2.2 million per year through reduced anesthesiologist overtime. The follow-up paper by Rath and Rajaram (2018) focuses specifically on staff planning at the RRMC (deciding how many anesthesiologists are needed on regular duty and how many on overtime). Although such decisions involve tangible, explicit costs (e.g., requiring an on-call anesthesiologist to come in requires a payment), they also involve implicit costs (e.g., having an anesthesiologist on call but not calling him or her). The paper proposes an approach that learns such implicit costs from data on prior staffing decisions and then uses these estimates within a two-stage integer stochastic dynamic program to make staffing decisions that lead to improvements in total (implicit plus explicit) cost.

Another example of optimization being applied to hospital operations is the paper by Zenteno et al. (2016). This paper describes the successful implementation by Massachusetts General Hospital (MGH) of a large-scale optimization model in order to reduce the surgical patient flow congestion in the perioperative environment without requiring additional resources. The paper studies an optimization model that rearranges the elective block schedule in order to smooth the average inpatient census by reducing the maximum average occupancy throughout the week. This model was implemented by MGH and resulted in increasing the effective capacity of the surgical units at MGH. Outside of scheduling and staffing, the paper by Bravo et al. (2019) develops a linear optimization-based framework, inspired by network revenue management, for making resource allocation and case mix decisions in a hospital network. Using data from an academic medical center, the paper shows how the approach can lead to better case mix decisions than current practice, which relies on prioritizing services based on traditional cost-accounting methods.

Outside of optimization, the paper by Ang et al. (2015) focuses on predictive modeling, specifically the problem of predicting emergency department (ED) wait times in hospitals. Many approaches have been proposed to predict ED wait times, such as rolling-average estimators, metrics based on fluid models, and quantile regression. The paper proposes a method that combines the LASSO method of statistical learning (Tibshirani 1996) with predictor variables that are derived from a generalized fluid model of the queue. The paper shows, using real hospital data, how the proposed method outperforms other approaches to predicting ED wait times, leading to reductions in mean squared error of up to 30% relative to the

rolling-average method that is commonly used in hospitals today.

4.3. Patient-Level Problems

Finally, the increasing availability of electronic health record data has led to the application of analytics methods to patient-level problems that are of a medical nature. For example, Bastani and Bayati (2015) consider a multiarmed bandit problem where, at each period, the decision maker can choose one of finitely many decisions, and the reward of each decision is a function of a high-dimensional covariate vector. The problem is to find a policy that achieves minimal regret. To do this, the paper proposes a novel algorithm based on the LASSO method. The specific healthcare application studied in the paper is that of learning to dose optimally: a physician encounters patients described by various features (e.g., patient's gender, patient's age, whether the patient has other preexisting conditions, whether the patient is taking other drugs, and so on) and must prescribe a dose to each patient but must learn the effect of different dosage levels on the fly. Using a data set on warfarin dosing, the paper demonstrates how the approach can be used effectively to prescribe initial doses to patients who are starting warfarin therapy.

Another example is the work of Bertsimas et al. (2017). This paper considers the problem of recommending a personalized drug treatment regimen to a patient with diabetes based on his or her medical attributes so as to minimize the patient's glycated hemoglobin A1c (HbA1c) level, a measure of a patient's baseline blood glucose level over the preceding two to three months. The paper proposes an approach that uses the k -nearest-neighbors algorithm to determine, for a given patient, which patients are most similar to that patient in terms of their demographics (such as age and sex), medical history (such as days since first diabetes diagnosis), and treatment history (such as the number of regimens previously tried). By focusing on patients similar to the one at hand, one can obtain an estimate of what the patient's HbA1c level will be under each possible drug regimen and thereby recommend a good drug regimen. Using data from Boston Medical Center, the paper shows how the approach can lead to clinically significant reductions in HbA1c level relative to the current standard of care.

5. Conclusions and Future Directions

In this paper, we have highlighted recent work in OM, leveraging methodological advances in machine learning and optimization that allow large-scale data to be used for complex decision making. We conclude our review by discussing some methodological directions that we expect to grow in importance in the future.

5.1. Causal Inference

Machine learning is usually concerned with predicting a dependent variable using a collection of independent variables. In many OM applications, the decision enters into the machine learning model as an independent variable. For such settings, a more relevant perspective to take is in understanding the *causal* impact of the decision on the dependent variable. This concern for the validity of decisions motivates the study of causal inference methods, which is a major area of focus in statistics and econometrics.

The relevance of causal inference to empirical OM has been highlighted previously (Ho et al. 2017), but less research has highlighted the importance of considering causality in prescriptive modeling. We have already discussed one example, the work of Baardman et al. (2018b), which develops a specialized estimation method for detecting customer-to-customer trends that uses the instrumental variables approach. Another example is the work of Bertsimas and Kallus (2016), which considers the problem of pricing from observational data. This paper shows that a standard approach to data-driven pricing, which involves fitting a regression model to price and then optimizing revenue as price multiplied by predicted demand, can actually be highly suboptimal when there exist confounding factors affecting both historical prices and demand. The paper develops a statistical hypothesis test for determining whether the price prescription from a predictive model is optimal and applies this approach to a car loan data set.

Outside of revenue management, Fisher et al. (2016), Bandi et al. (2018), and Chaurasia et al. (2019) study causal inference and prescriptive modeling questions in supply chain management, particularly relating to the speed of delivery of online orders and subsequent product returns. Fisher et al. (2016) use a quasi-natural experiment to demonstrate the economic value of faster delivery for an online retailer. Chaurasia et al. (2019) take this one step further by analyzing a transaction data set from a large fashion e-retailer and establishing a causal relationship between product delivery gaps in the supply chain and product returns. Furthermore, the paper embeds this causal model within a data-driven optimization framework in order to reduce product returns by optimizing product-delivery services. Using the e-retailer's data, the paper shows that a two-day delivery policy—that is, when all orders are attempted to be delivered within two days of order placement—can lead to cost savings of as much as \$8.9 million every year through reduced product returns. Bandi et al. (2018) also investigate the causes of product returns, with a focus on how dynamic pricing can lead to more product returns in the online retail industry. In particular, using data from the same online retailer, they develop a logistic

regression model that allows them to find that a postpurchase price drop is another cause of returns because it gives rise to opportunistic returns.

Recently, new approaches have also been emerging that combine machine learning with prescriptive analytics to deal with issues of heterogeneity. Alley et al. (2019) provide a framework for estimating price sensitivity when pricing tickets in a secondary market. Because of the heterogeneous nature of tickets, the unique market conditions at the time each ticket is listed, and the sparsity of available tickets, demand estimation needs to be done at the individual-ticket level. The paper introduces a double-orthogonalized machine-learning method for classification that isolates the causal effects of pricing on the outcome by removing the conditional effects of the ticket and market features. Furthermore, the paper shows how this price sensitivity estimation procedure can be embedded in an optimization model for selling tickets in a secondary market. This double-orthogonalization procedure can also be applied to other settings, such as personalized pricing and estimating intervention effects in healthcare.

5.2. Interpretability

Another exciting future direction is in the development of interpretable models in OM. In machine learning, an *interpretable model* is one in which a human being can easily see and understand how the model maps an observation into a prediction (Freitas 2014). Interpretability is a major area of research in machine learning for two reasons. First, interpretable models can provide insights into the prediction problem that black-box models, such as random forests and neural networks, cannot. Second, and more important, machine-learning models often do not affect decisions directly but instead make predictions or recommendations to a human decision maker; in many contexts (e.g., medicine), a decision maker is unlikely to accept a recommendation made by a machine-learning model without the ability to understand how the recommendation was made. In addition, there is also growing legislation, such as the General Data Protection Rules in the European Union (Doshi-Velez and Kim 2017), requiring that algorithms affecting users must be able to provide an explanation of their decisions.

Although interpretable machine learning is still new to OM, some recent papers have considered interpretability in the context of dynamic decision making. Bravo and Shaposhnik (2018) develop an approach called *mining optimal policies* (MinOP), wherein one analyzes exact solutions of stochastic dynamic programs using interpretable machine learning methods; this approach provides a modeler with insight into the structure of the optimal policy for a given family of

problems. Ciocan and Mišić (2018) propose an algorithm for finding policies for optimal stopping problems in the form of a binary tree, similar to trees commonly used in machine learning; the paper shows how the proposed tree policies outperform (non-interpretable) policies derived from approximate dynamic programming in the problem of option pricing while being significantly simpler to understand. With the continuing proliferation of machine learning-based decision support systems in modern business, we expect interpretable decision making to become a major area of research in OM.

5.3. “Small-Data” Methods

One challenge that is being recognized with the increasing availability of data is that these data are often small. This paradoxical statement refers to the fact that, often, although the number of observations one might have access to is large, the number of parameters that one must estimate is also large. For example, imagine an online retailer offering a very large selection of products in a particular category; the retailer might be interested in estimating the demand rate of each product but may only see a handful of purchases for each product in a year. The paper by Gupta and Rusmevichientong (2017) considers how to solve linear optimization problems in this small-data regime and shows that standard approaches in stochastic optimization, such as sample-average approximation, are suboptimal in this setting. The paper develops two different solution approaches, based on empirical Bayes methods and regularization, that enjoy desirable theoretical guarantees as the number of uncertain parameters grows and showcases the benefits of these approaches in an online advertising portfolio application.

A different approach, which specifically considers prediction in the small-data regime, can be found in the paper by Farias and Li (2019), which studies the following problem: a retailer offers a collection of products to a collection of users. For each product–user pair, we see whether the user has purchased that product or not, leading to a 0–1 matrix with users as rows and products as columns. Based on these data, we would like to estimate the probability of any given user purchasing any given product, leading to a matrix of user–product purchase probabilities; such a matrix is useful because it can guide product recommendations, for example. The problem is made challenging by its scale (e.g., Amazon has on the order of 100 million products and 100 million users) and the fact that there is a very small number of transactions per user, which inhibits one’s ability to accurately estimate the user–product purchase–probability matrix. To solve the problem, Farias and Li (2019) leverage

the availability of *side information*. For example, Amazon may see other types of user–product interactions—for example, when users view a product page, add a product to their wish list, add a product to their shopping cart, and so on. The data for each interaction can be represented as a user–product matrix in the same way as the principal interaction (making a purchase); each matrix thus can be viewed as a slice of a three-dimensional tensor. The paper develops a novel algorithm for recovering slices of the data tensor based on singular value decomposition and linear regression; using data from Xiami, an online music streaming service in China, the paper shows how the proposed method can more accurately predict user–song interactions than existing methods that do not use side information.

In an entirely different application, the tensor completion method of Farias and Li (2019) has been used for the problem of designing a blood test for cancer in the paper by Mahmoudi et al. (2017). Specifically, an emerging idea for designing a blood test for cancer is to insert nanoparticles into a blood sample and measure the binding levels of these nanoparticles; this leads to data in the form of a three-dimensional tensor. Although the raw data have too much noise to be directly used in a predictive model, Mahmoudi et al. (2017) leverage the method of Farias and Li (2019) to denoise the data. Using data from a longitudinal study, Mahmoudi et al. (2017) show how the resulting models can predict whether a patient will develop cancer with diagnostically significant accuracy seven to eight years *after* the blood sample was drawn.

5.4. New Approaches to “Predict-Then-Optimize”

Many of the papers described earlier rely on the “predict-then-optimize” paradigm: first, build a predictive model using data, and then embed that model within the objective function of an optimization problem. Machine learning models are typically estimated so as to lead to good out-of-sample predictions; however, such models may not necessarily lead to good out-of-sample decisions. Rather than estimating models that minimize loss functions measuring predictive performance, the paper by Elmachtoub and Grigas (2017) proposes minimizing a loss function that is related to the objective value of the decisions that ensue from the model. The paper shows how this new estimation approach leads to decisions that outperform the traditional approach in fundamental problems, such as the shortest path, assignment, and portfolio selection problems. Given the prescriptive focus of OM, this work highlights the potential opportunities for revisiting how machine-learning models are constructed for prescriptive applications in OM.

References

- Acimovic J, Graves SC (2014) Making better fulfillment decisions on the fly in an online retail environment. *Manufacturing Service Oper. Management* 17(1):34–51.
- Ahuja RK, Orlin JB (2001) Inverse optimization. *Oper. Res.* 49(5): 771–783.
- Alley M, Biggs M, Hariss R, Herman C, Li M, Perakis G (2019) Pricing for heterogeneous products: Analytics for ticket reselling. Working paper, Massachusetts Institute of Technology, Cambridge.
- Ang E, Kwasnick S, Bayati M, Plambeck EL, Aratow M (2015) Accurate emergency department wait time prediction. *Manufacturing Service Oper. Management* 18(1):141–156.
- Angrist JD, Pischke J-S (2008) *Mostly Harmless Econometrics: An Empiricist's Companion* (Princeton University Press, Princeton, NJ).
- Aouad A, Farias VF, Levi R (2015) Assortment optimization under consider-then-choose choice models. Working paper, London Business School, London.
- Aouad A, Farias V, Levi R, Segev D (2018) The approximability of assortment optimization under ranking preferences. *Oper. Res.* 66(6):1661–1669.
- Aswani A, Shen Z-JM, Siddiq A (2019) Data-driven incentive design in the Medicare shared savings program. *Oper. Res.* 67(4): 1002–1026.
- Avrahami A, Herer YT, Levi R (2014) Matching supply and demand: Delayed two-phase distribution at Yedioth Group—models, algorithms, and information technology. *Interfaces* 44(5):445–460.
- Baardman L, Levin I, Perakis G, Singhvi D (2018a) Leveraging comparables for new product sales forecasting. Working paper, Massachusetts Institute of Technology, Cambridge.
- Baardman L, Borjian Boroujeni S, Cohen-Hillel T, Panchamgam K, Perakis G (2018b) Detecting customer trends for optimal promotion targeting. Working paper, Massachusetts Institute of Technology, Cambridge.
- Baardman L, Cohen M, Panchamgam K, Perakis G, Segev D (2018c) Scheduling promotion vehicles to boost profits. *Management Sci.* 65(1):50–70.
- Ban G-Y, Keskin NB (2017) Personalized dynamic pricing with machine learning. Working paper, London Business School, London.
- Ban G-Y, Rudin C (2019) The big data newsvendor: Practical insights from machine learning. *Oper. Res.* 67(1):90–108.
- Ban G-Y, Gallien J, Mersereau AJ (2018) Dynamic procurement of new products with covariate information: The residual tree method. *Manufacturing Service Oper. Management*, ePub ahead of print December 10, <https://doi.org/10.1287/msom.2018.0725>.
- Bandi C, Moreno A, Ngwe D, Xu Z (2018) Opportunistic returns and dynamic pricing: Empirical evidence from online retailing in emerging markets. Harvard Business School Working Paper 19-030, Harvard University, Boston.
- Bastani H, Bayati M (2015) Online decision-making with high-dimensional covariates. Working paper, University of Pennsylvania, Philadelphia.
- Bernstein F, Modaresi S, Sauré D (2019) A dynamic clustering approach to data-driven assortment personalization. *Management Sci.* 65(5):2095–2115.
- Bertsimas D, Kallus N (2016) The power and limits of predictive approaches to observational-data-driven optimization. Working paper, Massachusetts Institute of Technology, Cambridge.
- Bertsimas D, Kallus N (2019) From predictive to prescriptive analytics. *Management Sci.*, ePub ahead of print August 23, <https://doi.org/10.1287/mnsc.2018.3253>.
- Bertsimas D, Mišić VV (2019) Exact first-choice product line optimization. *Oper. Res.* 67(3):651–670.
- Bertsimas D, Brown DB, Caramanis C (2011) Theory and applications of robust optimization. *SIAM Rev.* 53(3):464–501.
- Bertsimas D, Farias VF, Trichakis N (2013) Fairness, efficiency, and flexibility in organ allocation for kidney transplantation. *Oper. Res.* 61(1):73–87.
- Bertsimas D, O'Hair A, Relyea S, Silberholz J (2016) An analytics approach to designing combination chemotherapy regimens for cancer. *Management Sci.* 62(5):1511–1531.
- Bertsimas D, Kallus N, Weinstein AM, Zhuo YD (2017) Personalized diabetes management using electronic medical records. *Diabetes Care* 40(2):210–217.
- Blanchet J, Gallego G, Goyal V (2016) A Markov chain approximation to choice modeling. *Oper. Res.* 64(4):886–905.
- Bravo F, Shaposhnik Y (2018) Mining optimal policies: A pattern recognition approach to model analysis. Working paper, University of California, Los Angeles, Los Angeles.
- Bravo F, Braun M, Farias V, Levi R, Lynch C, Tumolo J, Whyte R (2019) Optimization-driven framework to understand healthcare network costs and resource allocation. Working paper, University of California, Los Angeles, Los Angeles.
- Breiman L (1996) Bagging predictors. *Machine Learn.* 24(2):123–140.
- Breiman L (2001) Random forests. *Machine Learn.* 45(1):5–32.
- Breiman L, Friedman J, Stone CJ, Olshen RA (1984) *Classification and Regression Trees* (CRC Press, Boca Raton, FL).
- Chaurasia M, Pandey S, Perakis G, Rathore HS, Singhvi D, Spanditakis Y (2019) First delivery gaps: A supply chain level to reduce product returns in online retail. Working paper, Massachusetts Institute of Technology, Cambridge.
- Chen X, Owen Z, Pixton C, Simchi-Levi D (2015) A statistical learning approach to personalization in revenue management. Working paper, Massachusetts Institute of Technology, Cambridge.
- Ciocan DF, Mišić VV (2018) Interpretable optimal stopping. Working paper, INSEAD, Fontainebleau, France.
- Cohen M, Lobel I, Paes Leme R (2016) Feature-based dynamic pricing. Working paper, New York University, New York.
- Cohen MC, Kalas J, Perakis G (2017a) Optimizing promotions for multiple items in supermarkets. Working paper, Massachusetts Institute of Technology, Cambridge.
- Cohen MC, Leung N-HZ, Panchamgam K, Perakis G, Smith A (2017b) The impact of linear optimization on promotion planning. *Oper. Res.* 65(2):446–468.
- Cohen-Hillel T, Panchamgam K, Perakis G (2019a) High-low promotion policies for peak-end demand models. Working paper, Massachusetts Institute of Technology, Cambridge.
- Cohen-Hillel T, Panchamgam K, Perakis G (2019b) Bounded memory peak end models can be surprisingly good. Working paper, Massachusetts Institute of Technology, Cambridge.
- Delage E, Ye Y (2010) Distributionally robust optimization under moment uncertainty with application to data-driven problems. *Oper. Res.* 58(3):595–612.
- Doshi-Velez F, Kim B (2017) Toward a rigorous science of interpretable machine learning. Working paper, Harvard University, Cambridge.
- Elmachtoub AN, Grigas P (2017) Smart “predict, then optimize.” Working paper, Columbia University, New York.
- Ettl M, Harsha P, Papush A, Perakis G (2019) A data-driven approach to personalized bundle pricing and recommendation. *Manufacturing Service Oper. Management*, ePub ahead of print July 19, <https://doi.org/10.1287/msom.2018.0756>.
- Farias VF, Li AA (2019) Learning preferences with side information. *Management Sci.* 65(7):3131–3149.
- Farias VF, Jagabathula S, Shah D (2013) A nonparametric approach to modeling choice with limited data. *Management Sci.* 59(2): 305–322.
- Feldman J, Paul A, Topaloglu H (2019) Assortment optimization with small consideration sets. *Oper. Res.* Forthcoming.

- Ferreira KJ, Lee BHA, Simchi-Levi D (2015) Analytics for an online retailer: Demand forecasting and price optimization. *Manufacturing Service Oper. Management* 18(1):69–88.
- Fisher M, Gallino S, Xu J (2016) The value of rapid delivery in online retailing. Working paper, University of Pennsylvania, Philadelphia.
- Freitas AA (2014) Comprehensible classification models: A position paper. *ACM SIGKDD Explorations Newslett.* 15(1):1–10.
- Glaeser CK, Fisher M, Su X (2018) Optimal retail location: Empirical methodology and application to practice. *Manufacturing Service Oper. Management* 21(1):86–102.
- Golrezaei N, Nazerzadeh H, Rusmevichientong P (2014) Real-time optimization of personalized assortments. *Management Sci.* 60(6): 1532–1551.
- Gupta V, Rusmevichientong P (2017) Small-data, large-scale linear optimization with uncertain objectives. Working paper, University of Southern California, Los Angeles.
- Gupta V, Han BR, Kim S-H, Paek H (2017) Maximizing intervention effectiveness. Working paper, University of Southern California, Los Angeles.
- Harsha P, Subramanian S, Uichanco J (2019) Dynamic pricing of omnichannel inventories. *Manufacturing Service Oper. Management* 21(1):47–65.
- Hawkins JT (2003) A Lagrangian decomposition approach to weakly coupled dynamic optimization problems and its applications. Unpublished PhD thesis, Massachusetts Institute of Technology, Cambridge.
- He L, Mak H-Y, Rong Y, Shen Z-JM (2017) Service region design for urban electric vehicle sharing systems. *Manufacturing Service Oper. Management* 19(2):309–327.
- Ho T-H, Lim N, Reza S, Xia X (2017) OM Forum: Causal inference models in operations management. *Manufacturing Service Oper. Management* 19(4):509–525.
- Javanmard A, Nazerzadeh H (2019) Dynamic pricing in high-dimensions. *J. Machine Learn. Res.* 20(9):1–49.
- Mahmoudi M, Caracciolo G, Safavi-Sohi R, Poustchi H, Li AA, Vasighi M, Chiozzi RZ, et al. (2017) Multi-nanoparticle protein corona characterization of human plasma and machine learning enable accurate identification and discrimination of cancers at early stages. Working paper, Harvard Medical School, Cambridge, MA.
- Rath S, Rajaram K (2018) Staff planning for hospitals with cost estimation and optimization. Working paper, University of North Carolina at Chapel Hill, Chapel Hill.
- Rath S, Rajaram K, Mahajan A (2017) Integrated anesthesiologist and room scheduling for surgeries: Methodology and application. *Oper. Res.* 65(6):1460–1478.
- Tibshirani R (1996) Regression shrinkage and selection via the lasso. *J. Roy. Statist. Soc. B.* 58(1):267–288.
- Vielma JP (2015) Mixed integer linear programming formulation techniques. *SIAM Rev.* 57(1):3–57.
- Zenteno AC, Carnes T, Levi R, Daily BJ, Dunn PF, Systematic OR (2016) Block allocation at a large academic medical center. *Ann. Surgery* 264(6):973–981.