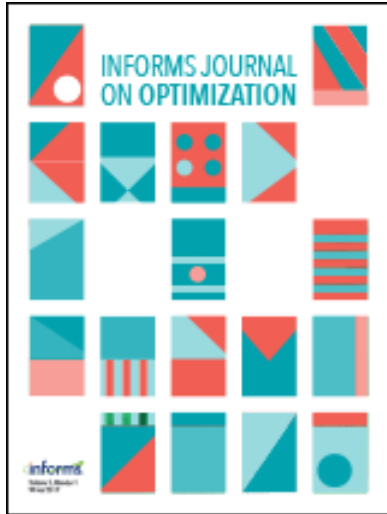




INFORMS Journal on Optimization

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

Novel Target Discovery of Existing Therapies: Path to Personalized Cancer Therapy

Dimitris Bertsimas, Ying Daisy Zhuo

To cite this article:

Dimitris Bertsimas, Ying Daisy Zhuo (2020) Novel Target Discovery of Existing Therapies: Path to Personalized Cancer Therapy. INFORMS Journal on Optimization 2(1):1-13. <https://doi.org/10.1287/ijoo.2019.0019>

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 © 2020, 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>

Novel Target Discovery of Existing Therapies: Path to Personalized Cancer Therapy

Dimitris Bertsimas,^a Ying Daisy Zhuo^a

^a Sloan School of Management and Operations Research Center, Massachusetts Institute of Technology, Cambridge, Massachusetts 02139

Contact: dbertsim@mit.edu,  <http://orcid.org/0000-0002-1985-1003> (DB); zhuo@mit.edu,  <http://orcid.org/0000-0002-8579-932X> (YDZ)

Received: August 12, 2018

Revised: December 29, 2018

Accepted: February 24, 2019

Published Online in Articles in Advance:
January 27, 2020

<https://doi.org/10.1287/ijoo.2019.0019>

Copyright: © 2020 INFORMS

Abstract. Discovering new drugs involves tremendous effort and financial resources, often at a significant risk of failed trials. Identifying new targets of existing drugs provides a promising direction, especially for molecular targeted cancer therapies. This paper presents a novel, machine learning, and optimization-based method that identifies potential targets of existing drugs to expand the treatable patient population. The method has the following advantages: (1) It is based on clinical and genomic data from a large national cancer hospital; (2) it incorporates state-of-the-art knowledge of cancer molecular biology and signaling pathways; and (3) it models patient heterogeneity explicitly outside genomics. The output is an ordered list of therapy–target pairs that our algorithm identifies as highly promising to be further tested. The results are highly accurate when validated against known mechanisms of action for existing drugs, where relationships such as pertuzumab–ERBB2, cetuximab–EGFR, and erlotinib–EGFR were independently identified. We found similar results in the external The Cancer Genome Atlas data set. The findings suggest that a data-driven optimization approach to precision cancer medicine may lead to breakthroughs in the drug-discovery process and recommend effective personalized cancer treatments given patient-specific genomic and phenotypic information.

History: Brian Denton served as editor-in-chief on this paper.

Keywords: mixed integer optimization • targeted therapy • cancer genomics

1. Introduction

Drug discovery is a costly and slow process. The average cost of developing a new pharmaceutical drug has risen to \$2.5 billion in the United States (DiMasi et al. 2016), and the average time from discovery to market is estimated to be 13.5 years (Paul et al. 2010). The high cost does not necessarily translate to the clinical efficacy of drug candidates identified; in fact, only 11% of drugs in clinical trials are eventually approved.

This increasing burden, risk, and low return-on-investment from de novo drug discovery have forced drug developers to look at alternatives. Repurposing existing drugs for novel uses is an attractive option to improve the research-and-development productivity, so long as the drugs turn out to be clinically effective on the novel uses identified. As the toxicity profile and pharmacokinetics of existing drugs has been extensively tested in prior trials, effective drug repurposing can significantly reduce the cost and regulatory approval time frame.

Cancer therapeutic development presents a particular opportunity for alternatives to de novo drug discovery. The advances in understanding of cancer biology at the molecular level and the wide availability of genetic biomarker testing expanded cancer drugs from traditional cytotoxic chemotherapies to more focused targeted therapies in the past decade (Hanahan and Weinberg 2011). These drugs interfere with specific molecular targets that are involved in the growth, progression, and spread of cancer and have demonstrated promising efficacy and tolerability for cancer patients in clinical settings.

Yet, as cancer therapies are becoming more precision-based and individualized, the size of patient cohorts with the specific molecular biomarker for a new drug continues to decline. In 2016, more than 40% of newly approved drugs had the orphan designation (indicated for fewer than 200,000 Americans), among which half are cancer drugs (U.S. Food and Drug Administration 2016). Between 2009 and 2015, the U.S. Food and Drug Administration approved 12 orphan cancer drugs indicated for biomarker-defined subsets (Kesselheim et al. 2017). If additional targets are identified for these highly specialized drugs, the increased market size can justify the high research-and-development cost.

Drug repurposing in cancer has seen some success stories. A notable example is crizotinib, first developed for anaplastic large-cell lymphoma as a MET inhibitor. Its ALK-inhibiting properties were later discovered and therefore repositioned to treat a subpopulation of non-small-cell lung cancer (NSCLC). The approval process

for NSCLC took only 4 years, significantly shorter than the average time (Shaw et al. 2011, Li and Jones 2012). Imatinib, originally approved in 2001 for chronic myelogenous leukemia, now has many indications, including for patients with KIT-positive gastrointestinal stromal tumors (Druker 2004, Novartis 2017). Many such hidden interactions remain to be discovered; it is estimated from known databases that, on average, a targeted cancer drug interacts with six molecular targets (Mestres et al. 2009).

Identifying such novel drug–target relationships accurately is, unfortunately, a difficult problem. The current approach to find new targets of existing drugs in practice typically relies on pure serendipity or anecdotal evidence. The combinatorial nature of the drug–target network prohibits the conventional trial-and-error process to be effective, as the number of possibilities explodes with respect to number of potential targets. With next-generation genome sequencing becoming more widely available, additional potential targets are identified, creating further opportunities as well as challenges for drug repurposing.

As a result, systematic, statistical, and computational approaches are needed for such discovery. In the past decade, progress has been made leveraging various data sources to suggest new targets and drug indications (Li et al. 2015, Chen and Butte 2016). In one strategy, the relationship is inferred based on the structural and chemical similarity between drugs (Li and Lu 2012) and/or based on existing indications and documented side effects (Campillos et al. 2008, Chiang and Butte 2009, Dudley et al. 2011). These studies incorporate the rich information from expert curated biological knowledge databases, which represent the current understanding, but may often be incomplete and inaccurate. In addition, the validation for these studies is typically based on internal data only. As a result, most of the identified drug-repositioning opportunities have not been successfully translated into clinical practice.

Recent advances in genomic technology opened up opportunities for large-scale association studies on drugs and biomarkers. Garnett et al. (2012) screened several hundred cancer cell lines under 130 drugs and identified some mutated genes that were associated with sensitivity to certain drugs. The Connectivity Map (Lamb et al. 2006) offers an excellent data source on the effect of compounds on the gene expression level of cancer cell lines and has been used by researchers for drug repositioning (Dudley et al. 2011, Zerbini et al. 2014). However, the data are from a few isolated cell lines in vitro, which may not reflect the true effect in vivo. Consequently, most studies typically focus on the gene signature changes rather than long-term clinical outcomes. Further, the effects of phenotypic information are often ignored in these studies (Hanahan and Weinberg 2011, Rubio-Perez et al. 2015, Spainhour and Qiu 2016).

The end goal for the drug-repositioning research would be to inform the selection of personalized drug combinations given a patient's genomic and phenotypic data, completing the bench-to-bedside translation. In clinical settings, treatment selection is still based on one or two signatures. To extend to multiple targets and combination treatments, some network-based approaches were taken to infer combinatorial drug effects (Tang et al. 2013). Alternatively, gene-expression-based methods were also considered—for instance, Pritchard et al. (2013) used a signature-based approach and identified cytotoxic chemotherapy combination mechanisms; Zhao et al. (2014) further incorporated drug efficacy and side effects to derive optimal combinations. On the phenotypic side, some recent work by Bertsimas et al. (2016) built a large database of clinical trials and used machine learning/optimization to design combination chemotherapy regimens for cancer. It is an important step toward personalized cancer therapy, although the nature of the aggregate data prevents it from fully capturing individual heterogeneity.

This study offers a data-driven approach to repurpose existing drugs by identifying novel drug–target relations. By integrating clinical and genomic data from thousands of patients at a large national cancer hospital with knowledge from biological processes, we are able to infer the interactions between drugs and targets from the binary coefficients of a nonlinear regression analysis, learned via mixed-integer optimization. Our model outputs a priority list of candidate targets for each drug to be further investigated. The inferred relations are validated against some of the known mechanisms of certain drugs.

Without randomized clinical trials, it is impossible to establish causal relationships from observational data only. We make a number of assumptions for our inferred drug and target interactions to be valid. We assumed a Bayesian network model structure with hidden variables to be the underlying ground-truth causal model. The model is assumed to have enough capacity to represent the relationships (linear effects of the demographic and medical information, for example). We also assume that the data are sufficiently large for the inference. Under these assumptions, the mixed-integer-optimization approach can then correctly recover with high precision the ground-truth drug–target relationships. We show this in a synthetic experiment where the ground truth is known.

Because one cannot verify the validity of the assumptions, given the observational nature of the data, the causal effects we inferred will still need to be confirmed by direct randomized experimentation. As a proxy, we test whether our assumed network model represents ground-truth causal models closer than simple ones

by evaluating the generalization performance against simpler linear models in a hold-out test set and an external data set using The Cancer Genome Atlas.

The findings provide a natural pathway toward personalized, mechanism-guided cancer therapy. Quantifying the impact of each agent in the signaling network allows for more precise treatment selection that targets the relevant pathways therapeutically, preventing the relapses often seen in targeted therapies. Given a patient with demographic, medical, and genomic information, the model allows practitioners to select combination therapies that holistically maximize the long-term effectiveness by targeting the right processes while maintaining low levels of off-target toxicity.

2. Data and Methods

2.1. Data

We retrospectively obtained the electronic health records data for patients at the Dana–Farber/Brigham and Women’s Cancer Center (DFCI) from 2004–2014. There are

- $i = 1, \dots, n$ patients,
- $j = 1, \dots, m$ treatments T_j ,
- $k = 1, \dots, K$ gene mutations G_k , and
- $\ell = 1, \dots, L$ sets S_ℓ that consists of genes (please see the significance in Section 2.2).

For each patient i , we have the following information:

$$G_{ik} = \begin{cases} 1, & \text{if gene mutation } G_k \text{ is present in patient } i, \\ 0, & \text{otherwise,} \end{cases}$$

and

$$T_{ij} = \begin{cases} 1, & \text{if patient } i \text{ received treatment } T_j, \\ 0, & \text{otherwise.} \end{cases}$$

In addition, $\mathbf{x}_i$ is the D -dimensional vector of phenotype covariates for patient i , including demographics, cancer diagnoses, comorbidities, prior treatments, resource utilization, and vital signs/laboratory tests results. y_i is the observed survival time in days for patient i from initiation of anticancer regimen, including chemotherapy, immunotherapy, and targeted therapy.

To be eligible for the study, patients must have initiated at least one anticancer regimen at the cancer center. Furthermore, patients were required to have measurements from *OncoPanel/OncoMap*, a sequencing platform that detects genetic mutations and other cancer-related DNA alterations.

We further obtained auxiliary data from curated biological and pharmacological databases. The sets $S_\ell, \ell = 1, \dots, L$ were obtained from Molecular Signatures Database, a curated database representing canonical pathways (Kanehisa and Goto 2000, Subramanian et al. 2005), which is widely used as a reference knowledge base for the interpretation and analysis of data sets generated from genome sequencing.

Missing values were imputed by using unconditional mean; the results based on other imputation algorithms such as knn-impute (Troyanskaya et al. 2001) and optimal-impute (Bertsimas et al. 2017) are included in sensitivity analyses. We used regimens initiated from 2004–2011 as the training set and regimens initiated from 2012–2014 as the validation set. The institutional review boards of the associated institutions (anonymized here) approved this study.

2.2. Model

It is widely believed that mutations in tumor-suppressor genes or oncogenes are linked to the development of cancer. Such genes are involved in critical signaling pathways that regulate cell cycles. When a mutation is present, the protein expression for that gene can be abnormal and adversely affect cell functions controlled by the pathway, subsequently affecting patient survival. In this section, we build the mathematical model of cancer treatments based on this biological mechanism, the structure of which is based on well-known literature in cancer, including Hanahan and Weinberg (2011), Weinberg (2013), and Sawyers (2004).

Targeted therapies, as an effective treatment of cancer, are known to interfere with the action of some proteins expressed by the mutated genes (G_k) by inhibiting their functions. However, it is unknown whether a treatment T_j targets a specific mutation G_k . Correspondingly, our key decision variable is to infer from data:

$$w_{jk} = \begin{cases} 1, & \text{if treatment } T_j \text{ targets gene mutation } G_k, \\ 0, & \text{otherwise.} \end{cases} \quad (1)$$

We postulate that time of survival y_i of patient i is affected by the vector of phenotype covariates $\mathbf{X}_i$, and specifically decreases when there are abnormal gene expressions that are not targeted by appropriate therapies. In particular, groups of gene expressions S_ℓ function together to affect survival. Mathematically, we assume the following form (where we denote observed data y_i, x_i with lowercase letters, but unobserved variables with capital letters $Y_i, \mathbf{X}_i$):

$$Y_i = \exp \left(\beta_0 + \mathbf{X}_i^T \boldsymbol{\beta} - \sum_{\ell=1}^L r_\ell s_{i\ell} \right), \quad (2)$$

where β_0 is the bias term, $\boldsymbol{\beta}$ is a D -dimensional vector to be estimated that models the influence of the corresponding phenotype $\mathbf{X}_i$, the variable r_ℓ to be estimated is the influence of having an abnormal set S_ℓ , and the decision variable $s_{i\ell}$ is defined as follows:

$$s_{i\ell} = \begin{cases} 1, & \text{if some gene expression for patient } i \text{ in the set } S_\ell \text{ is abnormal,} \\ 0, & \text{otherwise.} \end{cases} \quad (3)$$

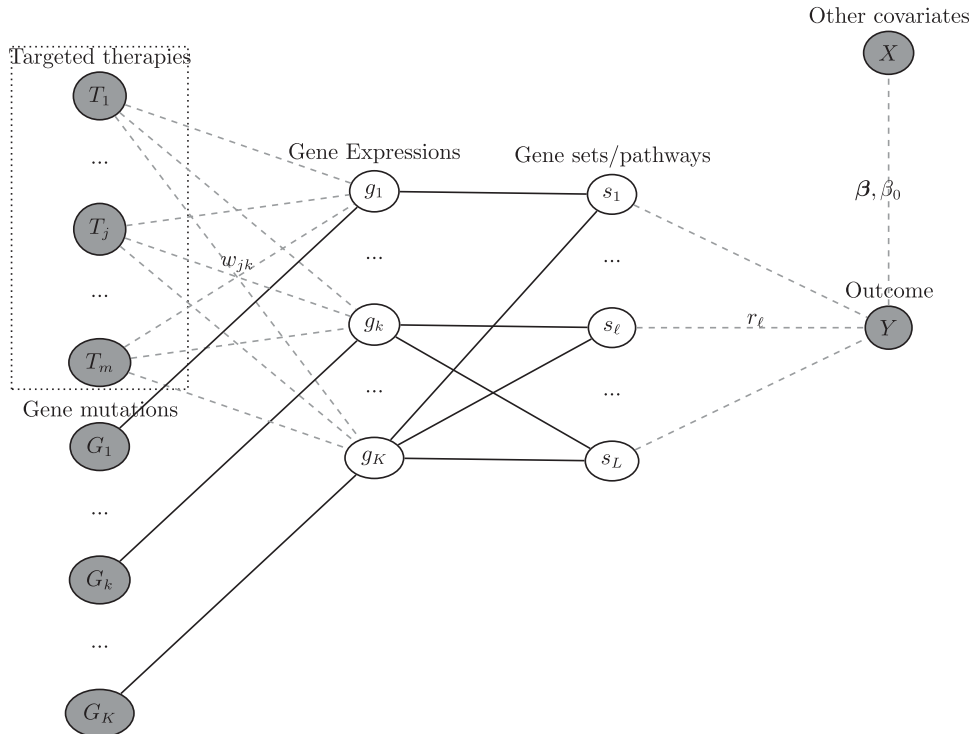
We obtain from Equation (2) that

$$\log(Y_i) = \beta_0 + \mathbf{X}_i^T \boldsymbol{\beta} - \sum_{\ell=1}^L r_\ell s_{i\ell}, \quad (4)$$

that is, presence of an abnormal set S_ℓ of gene expressions decreases life span by r_ℓ log days.

In summary, therapies target gene expressions. A gene expression is abnormal if the corresponding mutation is present and a blocking therapy is not applied to the patient. If a gene expression is abnormal, then all sets S_ℓ that contain it are also abnormal, and, thus, survival time is negatively affected from Equation (4). Our

Figure 1. Model Network Incorporating the Treatments, Gene Mutations, Gene Expressions, Signaling Pathway Gene Sets, Phenotype Covariates, and the Survival Outcomes



Notes. The dashed edges w_{jk} and the coefficients r_ℓ and $\boldsymbol{\beta}$ need to be determined by the mixed integer optimization problem. The solid edges and shaded variables are known and represent inputs to our model. The unshaded variables are derived.

objective is to infer the principal variable w_{jk} , defined in Equation (1), the vector β , and the coefficients r_ℓ . In the process of inferring the above principal variables, we also infer auxiliary variables $s_{i\ell}$ defined in Equation (3), gene expression variable

$$g_{ik} = \begin{cases} 1, & \text{if gene mutation } G_k \text{ is present in patient } i \text{ but not targeted,} \\ 0, & \text{otherwise;} \end{cases} \quad (5)$$

as well as the product between r_ℓ and $s_{i\ell}$, denoted by:

$$v_{i\ell} = r_\ell s_{i\ell} = \begin{cases} r_\ell, & \text{if } s_{i\ell} = 1, \\ 0, & \text{otherwise.} \end{cases} \quad (6)$$

Figure 1 depicts the inference problem we are solving, with the data and variables summarized in Table 1.

Because of the logical relations among the variables, the model can be naturally formulated by using mixed-integer optimization (MIO) (Bertsimas and Weismantel 2005). MIO is a particularly flexible tool for mathematical modeling, especially for incorporating logical constraints and relationships among variables and data. In recent years, advances in algorithms and computational hardware allowed a speedup of over several trillion times for solving MIO models, making even complex, large-scale problems tractable. We therefore selected it as our modeling technique.

The full MIO model formulation is as follows, with $i \in [n]$ and so on as shorthand notation for the enumeration $i \in 1, \dots, n$ and so on:

$$\min_{\mathbf{w}, \mathbf{g}, \beta, \beta_0, \mathbf{s}, \mathbf{r}} \sum_{i=1}^n \left| \log(y_i) - \left(\beta_0 + \mathbf{x}_i^T \beta - \sum_{\ell=1}^L v_{i\ell} \right) \right|, \quad (7a)$$

$$+ \lambda \sum_{j,k} w_{jk} + \gamma \|\beta\|_1, \quad (7b)$$

$$\text{s.t. } g_{ik} \leq G_{ik}, \quad \forall i \in [n], k \in [K], \quad (7c)$$

$$\sum_{j=1}^m T_{ij} w_{jk} \leq \left(\sum_{j=1}^m T_{ij} \right) (1 - g_{ik}), \quad \forall i \in [n], k \in [K], \quad (7d)$$

$$G_{ik} - \sum_{j=1}^m T_{ij} w_{jk} \leq g_{ik}, \quad \forall i \in [n], k \in [K], \quad (7e)$$

$$g_{ik} \leq s_{i\ell}, \quad \forall i \in [n], \ell \in [L], k \in S_\ell, \quad (7f)$$

$$\sum_{k \in S_\ell} g_{ik} \geq s_{i\ell}, \quad \forall i \in [n], \ell \in [L], \quad (7g)$$

$$v_{i\ell} \leq M s_{i\ell}, \quad \forall i \in [n], \ell \in [L], \quad (7h)$$

$$v_{i\ell} \leq r_\ell, \quad \forall i \in [n], \ell \in [L], \quad (7i)$$

$$v_{i\ell} \geq r_\ell - M(1 - s_{i\ell}), \quad \forall i \in [n], \ell \in [L], \quad (7j)$$

Table 1. Data and Variable Descriptions for the Network Model

Data/variable	Definition
T_{ij}	Patient i received treatment T_j
G_{ik}	Gene mutation G_k is present in patient i
$\mathbf{x}_i$	Phenotype covariates of patient i
s_ℓ	Gene set ℓ
y_i	Survival in days for patient i
w_{jk}	Treatment T_j targets gene mutation G_k
g_{ik}	Gene mutation G_k is present in patient i but not targeted
β	Linear effect of covariate $\mathbf{x}$
β_0	Bias term
$s_{i\ell}$	Whether some gene expression for patient i in set S_ℓ is abnormal
r_ℓ	Effect of having an abnormal set S_ℓ

$$v_{i\ell} \geq 0, \quad \forall i \in [n], \ell \in [L], \quad (7k)$$

$$r_\ell \geq 0, \quad \forall \ell \in [L], \quad (7l)$$

$$s_{i\ell} \in \{0, 1\}, \quad \forall i \in [n], \ell \in [L], \quad (7m)$$

$$w_{jk} \in \{0, 1\}, \quad \forall j \in [m], k \in [K], \quad (7n)$$

$$g_{ik} \in \{0, 1\}, \quad \forall i \in [n], k \in [K], \quad (7o)$$

$$\beta \in \mathcal{R}^D, \beta_0 \in \mathcal{R}. \quad (7p)$$

The objective function Equation (7b) finds parameters to achieve high goodness-of-fit on empirical data (first term), as well as giving preference to simpler and sparser models (second and third terms). The penalty λ encourages fewer drug–target interactions identified, and γ is the coefficient for standard lasso-type ℓ_1 regularization (Tibshirani 1996). The values are selected via cross-validation, a procedure where an unseen data set is used to evaluate the best parameter choice for out-of-sample performance.

The constraints describe the biological mechanism and logical relationships in our model. Equations (7c), (7d), and (7e) jointly impose the logical definition in Equation (5). Specifically, from Equation (7c), if $G_{ik} = 0$, then $g_{ik} = 0$ —that is, if mutation G_k is not present in patient i , then gene expression g_k is normal for this patient. From Equation (7d), if there is a therapy T_j that patient i received and T_j targets mutation G_k , then $g_{ik} = 0$. From Equation (7e), if $G_{ik} = 1$ and $\sum_{j=1}^m T_{ij}w_{jk} = 0$ —that is, mutation is present but no targeting therapy is given—then $g_{ik} = 1$.

Equations (7f) and (7g) jointly impose the logical defined in Equation (3). From Equation (7f), if for some k in S_ℓ we have $g_{ik} = 1$, then $s_{i\ell} = 1$ —that is, if a gene expression g_k in patient i is abnormal, then all sets S_ℓ containing g_k are abnormal. From Equation (7g), if $\sum_{k \in S_\ell} g_{ik} = 0$, then $s_{i\ell} = 0$; that is, if all of the gene expressions in S_ℓ are normal, then the set S_ℓ is normal.

Finally, Equations (7h), (7i), and (7j) jointly define the nonlinear relationship between the s_ℓ and the coefficients r_ℓ as defined in Equation (6).

The key output of the MIO (7) are the binary variables w_{jk} that represent our prediction of whether treatment T_j targets mutation G_k . We also produce confidence estimates for w_{jk} by using the bootstrap method via sampling the data with replacement (Efron 1979), solving the MIO (7) multiple times, and aggregating the results. To ensure that the model generalizes well to different populations, in addition to bootstrapping with random resampling, we also include subsamples based on some simple propensity model on the patient demographic information.

2.3. Evaluations

2.3.1. Model Fit Comparisons. To validate that the MIO (7) model fits and generalizes better to real clinical outcomes than existing methods, we compare it against predictions made from linear models. We consider the following three models, each fitted on $\log(y)$ to minimize the total sum of absolute error, to be consistent with the MIO objective for a fair comparison:

- LM1: with phenotype covariates alone;
- LM2: with phenotype covariates plus treatments and gene mutations;
- LM3: with phenotype covariates, treatments, gene mutations, and the interactions between treatments and gene mutations. The interaction terms measure the incremental effects from a given treatment when a patient has a specific mutation. It is a common method for estimating individual treatment effects (Imai et al. 2013).

Regularization is used when the variables are in high dimensions. Each model is cross-validated to select the best regularization-parameter choice. Similar to bootstrapping with a special subpopulation, the cross-validation procedure also includes subpopulation in addition to random sampled population splits. The in-sample and out-of-sample performance measured as mean absolute error (MAE) of each model is reported.

2.3.2. Inferred Therapy–Target Accuracy. The key model output is the inferred values for variable w_{jk} on whether therapy T_j targets mutation G_k . It is, however, unknown what the ground truth is. Therefore, evaluating the accuracy of our inferred therapy–target pairs is not feasible.

To gain a rough sense on the quality of w_{jk} predictions, we use the current biomedical knowledge on existing drugs as a proxy for the ground truth. To do so, we obtain data from the Drugbank database that has current knowledge on drug data and target information (Law et al. 2013). The current understanding of therapy–target relations from this database is considered as ground truth, where our findings are evaluated against it. The accuracy, true-positive, and false-positive rates are reported. Note that because the current knowledge has

likely not discovered some true relationships, the estimated true-positive and false-positive rates from our model are quite conservative.

2.3.3. Further Validations. We further validate the model externally on an additional data set from The Cancer Genome Atlas (TCGA). TCGA is a program launched by the National Cancer Institute and National Human Genome Research Institute, with a goal to accelerate cancer research with comprehensive data on tumor samples. We extracted a similar set of genomic and clinical information from breast cancer, lung cancer, and skin cancer patients and ran the MIO (7) model to estimate the w_{jk} . The same set of model evaluations with model fit and inferred therapy–target accuracy were performed.

Finally, we performed an experiment on synthetically generated data where the ground truth of the therapy–target relations variable w_{jk} was known. The goal was to validate that solving MIO (7) can recover the relationships, assuming the model is an accurate description of the biological process. We first generated values for the principal variables w_{jk} , β , and r_{ℓ} ; next, we simulated the patient phenotype covariates x_i , gene mutation G_{ik} , and treatments T_{ij} . Finally, we calculated the survival outcome y_i based on the model structure described in Section 2.2 with varying levels of additive noise. Because the ground truth for w_{jk} is known, we can evaluate the accuracy, true-positive, and false-positive rates of the inferred values.

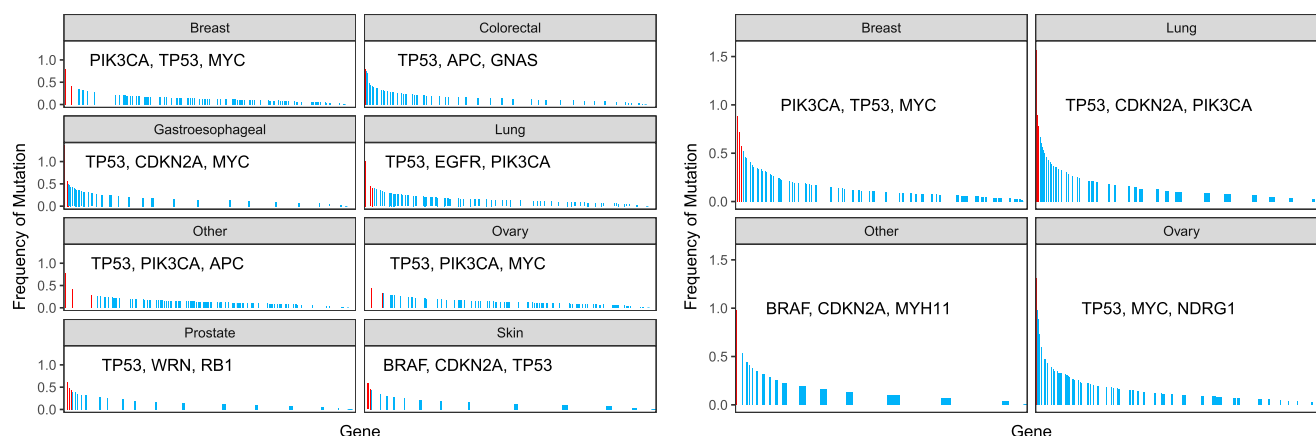
The software used included R (version 3.3.3), Julia (version 0.6.0), and Gurobi Optimizer (version 7.0).

3. Results

From the Dana–Farber Cancer Institute data set, we obtained the genetic and clinical data of 3,118 eligible patients. Common cancer sites included breast, ovary, lung, and colon/rectum. In total, these patients initiated $n = 6,928$ anticancer regimens. The median survival from the initiation of therapy was 550 days. Phenotype covariates (133 variables) included race, age, gender, tumor site, stage, metastatic status, comorbidities, laboratory test results, and major medications, with age at treatment (correlation coefficient of -0.1790 ; same measures below), age at diagnosis (-0.1626), and clinical stage IV (-0.1299) most negatively correlated with survival. The most frequently mutated genes were TP53, PIK3CA, and MYC (Figure 2, top) among $K = 176$ total measured gene mutations. The presence of most mutations was negatively correlated with the survival outcome, with TCF7L2 (-0.0615), CDKN2B (-0.0606), and PTEN (-0.0578) being the most negatively correlated. Among all prescribed chemotherapy drugs, the most commonly prescribed ones include carboplatin, paclitaxel, and cisplatin. Targeted therapies are less frequently prescribed, but we included the following $m = 10$ therapies that have sufficient observations: bevacizumab, trastuzumab, cetuximab, rituximab, ipilimumab, pertuzumab, pembrolizumab, panitumumab, erlotinib, and bortezomib. Among these, the use of trastuzumab (0.1351) and rituximab (0.0811) is most positively correlated with survival.

From The Cancer Genome Atlas data set, we obtained the genetic and clinical data of 1,098 eligible patients from the following projects: breast invasive carcinoma, ovarian serous cystadenocarcinoma, lung squamous cell carcinoma, lung adenocarcinoma, and skin cutaneous melanoma. In total, these patients initiated 6,929 regimens. The median survival from the initiation of therapy was 529 days. Phenotype covariates used include

Figure 2. (Color online) Most Frequently Mutated Genes Among All Observations, Stratified by the Tumor Site for DFCI on the Left and TCGA on the Right



Note. The top three mutations for each site is highlighted and the gene is labeled.

Table 2. Model Performance Comparisons for DFCI and TCGA Data

Model	In-sample MAE	Out-of-sample MAE
DFCI		
LM1	0.65	0.73
LM2	0.64	0.71
LM3	0.62	0.71
MIO	0.66	0.69
TCGA		
LM1	0.89	0.92
LM2	0.77	0.97
LM3	0.73	0.99
MIO	0.88	0.91

Note. The in-sample and out-of-sample mean absolute errors are reported for each of the four models being compared.

race, gender, tumor site, age at diagnosis, history of neoadjuvant treatment, history of radiation therapy, and some breast-cancer-specific ones (menopause status, estrogen-receptor status, etc.). The most frequently mutated genes were TP53, PIK3CA, MYC, NDRG1, EXT1, and RAD21 (Figure 2, bottom) out of a total of 299 common gene mutations under consideration. Among all prescribed chemotherapy drugs, the most commonly prescribed ones include paclitaxel, carboplatin, cyclophosphamide, doxorubicin, docetaxel, cisplatin, and tamoxifen. Targeted therapies are less frequently prescribed, but we included the following that have sufficient observations: bevacizumab, trastuzumab, erlotinib, ipilimumab, vemurafenib, gefitinib, denosumab, aldesleukin, and dabrafenib.

3.1. Model Fit Comparisons

Comparing the proposed mixed-integer optimization (7) model against linear-regression models, we achieved better model generalization performance, as reflected in the lower mean absolute error (Table 2) for both DFCI and TCGA data. Adding to LM1 (phenotypic information only) genomic data and treatment data as additional terms (LM2 and LM3) typically does not improve the generalization performance, even with the appropriate level of regularization. In contrast, because the MIO (7) model directly optimizes the sum of absolute error while maintaining sparsity, it improves upon any linear models significantly under MAE as the goodness-of-fit metric.

Between the two data sources, the DFCI one has very rich phenotype data that are likely to be predictive of the survival outcomes (Bertsimas et al. 2018) and therefore was able to achieve lower MAE both in sample and out of sample. The TCGA data have abundant information on genomic data; however, even with regularization, there is still some discrepancy between the in-sample and out-of-sample performance metrics. We attribute this discrepancy to the inherent high-dimensional nature of the data.

3.2. Inferred Therapy–Target Accuracy

The model outputs therapy–target pairs from solving MIO (7) with the DFCI data. Because the interactions are probabilities averaged from bootstrapped samples, at any chosen probability threshold cutoff, we can obtain a set of inferred interactions. We choose the threshold δ such that if the calculated confidence probability is greater than δ for a particular therapy–target combination, then we include this combination in the output we report, in a way to limit the number of distinct therapy–target combinations. When validating against known

Table 3. Confusion Table for Estimated w_{jk} Based on DFCI and TCGA Data, Comparing Predicted Therapy–Target Relationships Against Ones Known from Literature

	Actual 0	Actual 1
DFCI		
Predicted 0	1,684	3
Predicted 1	70	3
TCGA		
Predicted 0	3,524	11
Predicted 1	49	4

Table 4. Potential Drug–Target Relationships Identified by Model with DFCI Data, Ordered by Descending Probabilities

Drug	Target	Probability
Pembrolizumab	TP53	0.92
Trastuzumab	AKT2	0.74
Pertuzumab	PTK2	0.74
Rituximab	SRC	0.74
Pembrolizumab	CDK4	0.70
Bevacizumab	MAPK1	0.70
Pembrolizumab	MAPK1	0.70
Panitumumab	SRC	0.70
Panitumumab	IGF1R	0.68
Rituximab	PIK3CA	0.68
Erlotinib	TP53	0.68
Pembrolizumab	MTOR	0.66
Rituximab	TP53	0.66
Cetuximab	AKT3	0.64
Panitumumab	PTK2	0.64
Pertuzumab	MAP3K1	0.62
Bevacizumab	PIK3CA	0.62
Panitumumab	TP53	0.62
Trastuzumab	AKT3	0.60
Trastuzumab	RAF1	0.60
Pembrolizumab	EGFR	0.58
Rituximab	PTEN	0.58
Erlotinib	EGFR	0.56
Pembrolizumab	FGFR3	0.56
Bevacizumab	AKT1	0.54
Bevacizumab	BRAF	0.54
Trastuzumab	MAPK1	0.54
Bevacizumab	SRC	0.54
Bevacizumab	TP53	0.54
Panitumumab	MAPK1	0.52
Pertuzumab	PIK3CA	0.52
Cetuximab	RAF1	0.52
Cetuximab	SRC	0.52
Pembrolizumab	BRAF	0.50
Pembrolizumab	JAK2	0.50
Panitumumab	BCL2	0.48
Cetuximab	PIK3CA	0.48
Panitumumab	FGFR1	0.46
Pembrolizumab	MET	0.46
Cetuximab	MTOR	0.46
Ipilimumab	FLT4	0.44
Cetuximab	PTK2	0.44
Cetuximab	BRAF	0.42
Panitumumab	JAK2	0.42
Panitumumab	NTRK3	0.42
Ipilimumab	RAF1	0.42
Rituximab	AKT1	0.40
Bevacizumab	AKT3	0.40
Pertuzumab	ERBB2	0.40
Cetuximab	IGF1R	0.40
Erlotinib	MET	0.40
Pertuzumab	PRKDC	0.40
Trastuzumab	SRC	0.40
Pertuzumab	SRC	0.40
Trastuzumab	SYK	0.40
Pertuzumab	SYK	0.40
Trastuzumab	AKT1	0.38
Pertuzumab	AKT1	0.38
Ipilimumab	EGFR	0.38
Bevacizumab	IGF1R	0.38
Pertuzumab	TP53	0.38

Table 4. (Continued)

Drug	Target	Probability
Cetuximab	AKT2	0.36
Panitumumab	AURKA	0.36
Panitumumab	BCL2L1	0.36
Panitumumab	PTEN	0.36
Pertuzumab	PTPN11	0.36
Cetuximab	EGFR	0.34
Ipilimumab	MAP3K1	0.34
Trastuzumab	PIK3CA	0.34
Bevacizumab	PRKDC	0.34

Note. The relationships known in literature are in bold.

relationships, we aim at keeping around 60 predicted pairs, which correspond to a threshold of 0.34. The prediction achieved an accuracy of 95.8%, a true-positive rate of 50.0%, and a false-positive rate of 4.0% (Table 3, top section).

Sorting the predicted probabilities of therapy–target interactions gives a list to be further investigated in experiments (Table 4). We correctly recovered the well-known therapy–target pairs including erlotinib on EGFR, pertuzumab on ERBB2, and cetuximab on EGFR at this threshold. At a lower threshold of 0.20, we also discovered temsirolimus on MTOR and panitumumab on EGFR. The relationship trastuzumab on ERBB2 was not identified, likely because data do not suggest pronounced incremental treatment effects from trastuzumab on patients with an ERBB2 mutation.

We conducted similar analysis for the TCGA data. To obtain about 60 therapy–target pairs, we selected the threshold of 0.38. The prediction achieved an accuracy of 98.3%, a true-positive rate of 26.7%, and a false-positive rate of 1.4% (Table 3, bottom section). The prediction provides the list of the most promising therapy–target to be further investigated in experiments (Table 5). We correctly recovered the well-known relationships of dabrafenib on BRAF, vemurafenib on BRAF, gefitinib on EGFR, and lapatinib on ERBB2.

While interpreting the results, note that, again, because the known pairs used to validate our model against are not exactly the ground truth, as many pairs remain to be discovered, the reported model performance metrics of true-positive and false-positive rates may be overly conservative.

3.3. Synthetic Data Validation

In the validation studies with synthetic data, we confirmed that when the causal assumptions are met, the inferred values of w_{jk} are accurate. We ran the MIO (7) to assess the quality of recovery of w_{jk} based on synthetically generated data where the underlying parameters are known and the data follow the model structure. Figure 3 presents the model performance at varying levels of noise: The accuracy, true positive, and false positive rates are plotted against noise. At low levels of noise, the model achieves almost perfect recovery of w_{jk} ; as noise increases, the performance generally drops, but still remains reasonably high given the noise level.

4. Discussion

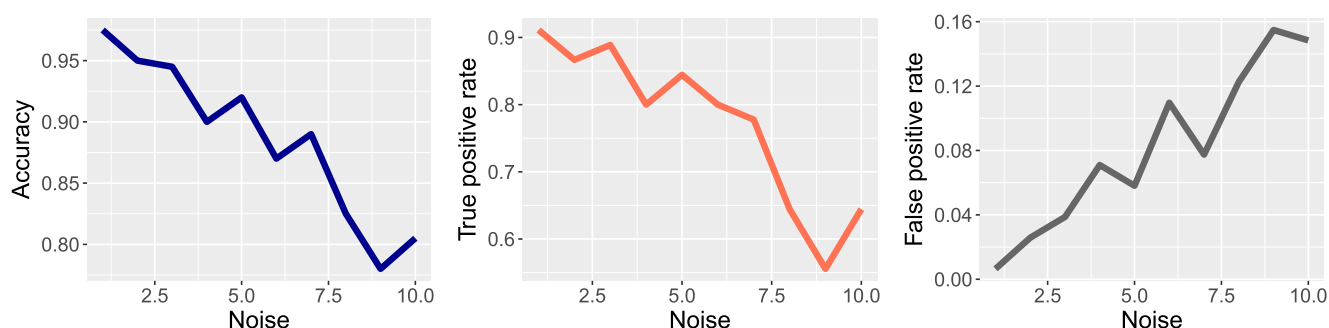
To our knowledge, this is the first study to leverage rich genomic and clinical information from observational in vivo data to inform drug discovery. It uses modern mathematical modeling tools to integrate data and biological knowledge base to identify the most likely therapy–target interactions. The network-based, MIO model is able to capture the nonlinear interaction between therapies and targets. Compared with black-box machine learning models, it outputs interpretable chemical and biological processes while still allowing for the flexibility to represent the complex interactions across the patient inputs.

This novel approach to drug discovery has the following advantages: (1) It is based on large-scale, real-world evidence to infer potential biochemical mechanisms, overcoming the limitation of many in vitro experiments; (2) it incorporates the state-of-the-art understanding of cancer molecular biology and signaling pathways organically into a data-driven method; and (3) it models patient heterogeneity explicitly outside genomic information, controlling for and capturing the rich information from patients' electronic medical records. As a result, the output of potential therapy–target interactions from this model was able to achieve high accuracy in external validation compared with other approaches.

Table 5. Potential Drug–Target Relationships Identified by Model with TCGA Data, Ordered by Descending Probabilities

Drug	Target	Probability
Dabrafenib	BRAF	0.98
Denosumab	IKBKB	0.98
Vemurafenib	BRAF	0.90
Ipilimumab	CACNA1D	0.86
Trastuzumab	MAPK1	0.86
Cetuximab	PRKACA	0.86
Dabrafenib	FGFR2	0.84
Denosumab	PIK3CA	0.84
Cetuximab	TP53	0.84
Erlotinib	NTRK1	0.82
Cetuximab	PIK3CA	0.80
Ipilimumab	AKT2	0.78
Pembrolizumab	BRAF	0.78
Gefitinib	TP53	0.76
Vemurafenib	AKT2	0.74
Dabrafenib	BCR	0.74
Cetuximab	MET	0.74
Dabrafenib	NFKB2	0.74
Erlotinib	TP53	0.74
Cetuximab	BRAF	0.72
Dabrafenib	PTEN	0.72
Ipilimumab	TP53	0.72
Dabrafenib	PDGFRA	0.70
Vemurafenib	PTEN	0.70
Trastuzumab	PIK3CA	0.68
Denosumab	FGFR1	0.66
Bevacizumab	GRIN2A	0.66
Cetuximab	JAK1	0.66
Ipilimumab	ATP1A1	0.64
Ipilimumab	BRAF	0.64
Bevacizumab	PIK3CA	0.64
Pembrolizumab	LCK	0.62
Vemurafenib	MET	0.62
Trastuzumab	TP53	0.62
Pembrolizumab	GRIN2A	0.60
Trastuzumab	IKBKB	0.58
Lapatinib	MAP2K2	0.58
Bevacizumab	IKBKB	0.56
Pembrolizumab	MLH1	0.56
Lapatinib	PIK3CA	0.56
Pembrolizumab	PMS2	0.56
Erlotinib	TSHR	0.56
Gefitinib	EGFR	0.54
Lapatinib	ERBB2	0.54
Trastuzumab	JAK1	0.54
Erlotinib	MDM2	0.54
Erlotinib	PDGFRA	0.54
Aldesleukin	AKT2	0.52
Aldesleukin	BRAF	0.52
Bevacizumab	CACNA1D	0.52
Dabrafenib	KDR	0.52
Cetuximab	MAP2K4	0.52
Lapatinib	TP53	0.52
Lapatinib	BCR	0.50
Erlotinib	IKBKB	0.50
Lapatinib	MALT1	0.50
Erlotinib	PIK3CA	0.50
Bevacizumab	POLE	0.50
Trastuzumab	PRKACA	0.50
Lapatinib	ATM	0.48

Note. The relationships known in literature are in bold.

Figure 3. (Color online) Model Performance on Synthetic Data for Recovering the Therapy–Target Relationship w_{jk} 

Note. The accuracy, true positive, and false positive rates against noise levels are presented.

The study has the following limitations: It is based on observational data, which unavoidably contain confounding factors. As a result, the learned relationships cannot be proven causal unless the assumptions we made hold true. The set of gene mutations measured is relatively small (296 measured mutations) in the DFCI data set; in particular, this data set did not have information on immune cells, therefore restricting our inference on the potential targets for immunotherapies. In addition, the available human knowledge of gene sets is not complete. For future work, drug-dosing information could be incorporated. As some of the target relationships were validated against findings in previous literature in this study, further validation should be conducted in prospective in vitro experiments, animal models, and clinical trials.

Extensions to the model are straightforward. Because the MIO has a very flexible structure, it allows researchers to readily incorporate novel biological discoveries and updated data. In closer collaboration with cancer biologists, geneticists, and pharmacologists, we can further fine-tune the network model to reflect more comprehensive and new knowledge in the targeted therapy mechanism of action. The scalability of the model encourages further computational discoveries using data from other institutions and potentially data consortium from multiple institutions. As the data availability and quality improve, we believe the model will recommend new targets with even higher confidence and accuracy.

Matching patients to the right set of therapies, especially when the genetic signatures and clinical presentations are complex, is the ultimate goal for personalized cancer therapy. In this work, we present a general method with validated findings on novel therapy–target interactions, which provide a firm foundation for an evidence-based design of targeted treatment recommendations. Looking ahead, this approach will present not only significant opportunities for discoveries and breakthroughs in cancer medicine, but also higher-quality, more individualized care for patients.

References

- Bertsimas D, Weismantel R (2005) *Optimization Over Integers* (Dynamic Ideas, Belmont, MA).
- Bertsimas D, Pawlowski C, Zhuo YD (2017) From predictive methods to missing data imputation: An optimization approach. *J. Machine Learn. Res.* 18(1):7133–7171.
- 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, Dunn J, Pawlowski C, Silberholz J, Weinstein A, Zhuo YD, Chen E, Elfiky AA (2018) Applied informatics decision support tool for mortality predictions in patients with cancer. *JCO Clinical Cancer Informatics* 2(2):1–11.
- Campillos M, Kuhn M, Gavin AC, Jensen LJ, Bork P (2008) Drug target identification using side-effect similarity. *Science* 321(5886):263–266.
- Chen B, Butte A (2016) Leveraging big data to transform target selection and drug discovery. *Clinical Pharmacology Therapeutics* 99(3):285–297.
- Chiang AP, Butte AJ (2009) Systematic evaluation of drug–disease relationships to identify leads for novel drug uses. *Clinical Pharmacology Therapeutics* 86(5):507–510.
- DiMasi JA, Grabowski HG, Hansen RW (2016) Innovation in the pharmaceutical industry: New estimates of R&D costs. *J. Health Econom.* 47: 20–33.
- Druker BJ (2004) Imatinib as a paradigm of targeted therapies. *Adv. Cancer Res.* 91:1–30.
- Dudley JT, Deshpande T, Butte AJ (2011) Exploiting drug–disease relationships for computational drug repositioning. *Briefings Bioinformatics* 12(4):303–311.
- Efron B (1979) Bootstrap methods: Another look at the jackknife. *Ann. Statist.* 7(1):1–26.
- Garnett MJ, Edelman EJ, Heidorn SJ, Greenman CD, Dastur A, Lau KW, Greninger P, et al. (2012) Systematic identification of genomic markers of drug sensitivity in cancer cells. *Nature* 483(7391):570–575.
- Hanahan D, Weinberg RA (2011) Hallmarks of cancer: the next generation. *Cell* 144(5):646–674.
- Imai K, Ratkovic M (2013) Estimating treatment effect heterogeneity in randomized program evaluation. *Ann. Appl. Statist.* 7(1):443–470.
- Kanehisa M, Goto S (2000) Kegg: Kyoto Encyclopedia of Genes and Genomes. *Nucleic Acids Res.* 28(1):27–30.

- Kesselheim AS, Treasure CL, Joffe S (2017) Biomarker-defined subsets of common diseases: Policy and economic implications of Orphan Drug Act coverage. *PLoS Medicine* 14(1):e1002190.
- Lamb J, Crawford ED, Peck D, Modell JW, Blat IC, Wrobel MJ, Lerner J, et al. (2006) The connectivity map: Using gene-expression signatures to connect small molecules, genes, and disease. *Science* 313(5795):1929–1935.
- Law V, Knox C, Djoumbou Y, Jewison T, Guo AC, Liu Y, Maciejewski A, et al. (2013) Drugbank 4.0: Shedding new light on drug metabolism. *Nucleic Acids Res.* 42(D1):D1091–D1097.
- Li J, Lu Z (2012) A new method for computational drug repositioning using drug pairwise similarity. *2012 IEEE Internat. Conf. Bioinform. Biomed. (BIBM)* (IEEE, Piscataway, NJ), 1–4.
- Li J, Zheng S, Chen B, Butte AJ, Swamidass SJ, Lu Z (2015) A survey of current trends in computational drug repositioning. *Briefings Bioinformatics* 17(1):2–12.
- Li YY, Jones SJ (2012) Drug repositioning for personalized medicine. *Genome Medicine* 4(3):Article 27.
- Mestres J, Gregori-Puigjané E, Valverde S, Solé RV (2009) The topology of drug–target interaction networks: Implicit dependence on drug properties and target families. *Molecular BioSystems* 5(9):1051–1057.
- Novartis (2017) Gleeevec: Prescribing information. Accessed June 2, 2017, https://www.accessdata.fda.gov/drugsatfda_docs/label/2006/021588s009b1.pdf.
- Paul SM, Mytelka DS, Dunwiddie CT, Persinger CC, Munos BH, Lindborg SR, Schacht AL (2010) How to improve R&D productivity: The pharmaceutical industry’s grand challenge. *Nature Rev. Drug Discovery* 9(3):203–214.
- Pritchard JR, Bruno PM, Gilbert LA, Capron KL, Lauffenburger DA, Hemann MT (2013) Defining principles of combination drug mechanisms of action. *Proc. Natl. Acad. Sci. USA* 110(2):E170–E179.
- Rubio-Perez C, Tamborero D, Schroeder MP, Antolín AA, Deu-Pons J, Perez-Llamas C, Mestres J, Gonzalez-Perez A, Lopez-Bigas N (2015) In silico prescription of anticancer drugs to cohorts of 28 tumor types reveals targeting opportunities. *Cancer Cell* 27(3):382–396.
- Sawyers C (2004) Targeted cancer therapy. *Nature* 432(7015):294–297.
- Shaw AT, Yasothan U, Kirkpatrick P (2011) Crizotinib. *Nature Rev. Drug Discovery* 10(12):897–898.
- Spainhour JCG, Qiu P (2016) Identification of gene-drug interactions that impact patient survival in TCGA. *BMC Bioinformatics* 17(1):Article 409.
- Subramanian A, Tamayo P, Mootha VK, Mukherjee S, Ebert BL, Gillette MA, Paulovich A, et al. (2005) Gene set enrichment analysis: A knowledge-based approach for interpreting genome-wide expression profiles. *Proc. Natl. Acad. Sci. USA* 102(43):15545–15550.
- Tang J, Karhinen L, Xu T, Szwajda A, Yadav B, Wennerberg K, Aittokallio T (2013) Target inhibition networks: Predicting selective combinations of druggable targets to block cancer survival pathways. *PLOS Comput. Biol.* 9(9):e1003226.
- Tibshirani R (1996) Regression shrinkage and selection via the lasso. *J. Royal Stat. Soc. B.* 58(1):267–288.
- Troyanskaya O, Cantor M, Sherlock G, Brown P, Hastie T, Tibshirani R, Botstein D, Altman RB (2001) Missing value estimation methods for DNA microarrays. *Bioinformatics* 17(6):520–525.
- U.S. Food and Drug Administration (2016) Drug innovation—novel drugs summary 2016. Accessed June 2, 2017, <https://www.fda.gov/Drugs/DevelopmentApprovalProcess/DrugInnovation/ucm534863.htm>.
- Weinberg R (2013) *The Biology of Cancer* (Garland Science, New York).
- Zerbini LF, Bhasin MK, de Vasconcellos JF, Paccez JD, Gu X, Kung AL, Libermann TA (2014) Computational repositioning and preclinical validation of pentamidine for renal cell cancer. *Molecular Cancer Therapeutics* 13(7):1929–1941.
- Zhao B, Hemann MT, Lauffenburger DA (2014) Intratumor heterogeneity alters most effective drugs in designed combinations. *Proc. Natl. Acad. Sci. USA* 111(29):10773–10778.