

Operator-Induced Structural Variable Selection for Identifying Materials Genes

Shengbin Ye, Thomas P. Senftle & Meng Li

To cite this article: Shengbin Ye, Thomas P. Senftle & Meng Li (2024) Operator-Induced Structural Variable Selection for Identifying Materials Genes, Journal of the American Statistical Association, 119:545, 81-94, DOI: [10.1080/01621459.2023.2294527](https://doi.org/10.1080/01621459.2023.2294527)

To link to this article: <https://doi.org/10.1080/01621459.2023.2294527>

View supplementary material 13'

a

Published online: 12 Feb 2024.

Submit your article to this journal 13'

!

Article views: 336

!J

View related articles 13'

(R}

C. .Muk

View Crossmark data 13'

Operator-Induced Structural Variable Selection for Identifying Materials Genes

Shengbin Ye^a, Thomas P. Senftle^b, and Meng U^a

^aDepartment of Statistics, Rice University, Houston, TX; ^bDepartment of Chemical and Biomolecular Engineering, Rice University, Houston, TX

ABSTRACT

In the emerging field of materials informatics, a fundamental task is to identify physicochemically meaningful descriptors, or materials genes, which are engineered from primary features and a set of elementary algebraic operators through compositions. Standard practice directly analyzes the high-dimensional candidate predictor space in a linear model; statistical analyses are then substantially hampered by the daunting challenge posed by the astronomically large number of correlated predictors with limited sample size. We formulate this problem as variable selection with operator-induced structure (OIS) and propose a new method to achieve unconventional dimension reduction by using the geometry embedded in OIS. Although the model remains linear, we iterate nonparametric variable selection for effective dimension reduction. This enables variable selection based on *ab initio* primary features, leading to a method that is orders of magnitude faster than existing methods, with improved accuracy. To select the nonparametric module, we discuss a desired performance criterion that is uniquely induced by variable selection with OIS; in particular, we propose to employ a Bayesian Additive Regression Trees (BART)-based variable selection method. Numerical studies show superiority of the proposed method, which continues to exhibit robust performance when the input dimension is out of reach of existing methods. Our analysis of single-atom catalysis identifies physical descriptors that explain the binding energy of metal-support pairs with high explanatory power, leading to interpretable insights to guide the prevention of a notorious problem called sintering and aid catalysis design. Supplementary materials for this article are available on line.

ARTICLE HISTORY

Received November 2021
Accepted November 2023

KEYWORDS

BART; Bayesian nonparametrics; Feature engineering; Materials genomes; Nonparametric dimension reduction

1. Introduction

The Materials Genome Initiative set up by the White House is a large-scale effort concerning the utilization of computational tools to accelerate the pace of discovery and deployment of advanced material systems. Since its inception in 2011, there has been a surge of interest in data-driven materials design and understanding (Zhong et al. 2020; Hart et al. 2021; Keith et al. 2021; Lin et al. 2021; Liu et al. 2021). In this nascent area called materials informatics, computational methods that account for physical and chemical mechanisms of a material system play a central role in aiding, augmenting, or even replacing the time-consuming trial and error experimentation. A fundamental task is to identify physicochemically meaningful *descriptors*, or *materials genes* (Ghiringhelli et al. 2015; Poppa et al. 2021). These descriptors, for example, are key to modeling single-atom catalysis and finding or developing more efficient catalytically active materials. In statistical terms, descriptors are high-dimensional predictors but with strong structure in that they are functional transformations of a set of primary features $X = (x_1, \dots, x_p)$. For instance, a simple example of descriptors is $f(X) = \{\exp(x_1) - \exp(x_2)\}^2$, which can be constructed using exponential and squared functions in combination with subtraction.

Suppose the response vector y measures the material property of interest, and the primary features matrix $X =$

$(x_1, \dots, x_p)$ collects physical or chemical properties of the materials such as atomic radii, ionization energies, etc. Then the space of engineered predictors (or descriptors) up to order M is $(\cdot)(M)(X)$, which consists of nonlinear predictors with explicit functional form resulting from M -order compositions of operators $(\cdot)$ on X :

$$o \in (M)(X) = (\cdot)_0 \circ \dots \circ (\cdot)_M(X) = \underbrace{(\cdot)_0 \dots \circ (\cdot)_M}_{M \text{ times}} J(X).$$

For example, some commonly used operators in materials genome are

$$(\cdot) = \{+, -, \times, /, 1 - 1, \exp, \log, |\cdot|, \cdot^{-1}, \sin(\pi \cdot), \cos(\pi \cdot)\}, \quad (1)$$

and the aforementioned descriptor $f(X) = \{\exp(x_1) - \exp(x_2)\}^2$ belongs to $o \in (3)(X)$. We refer to this distinctive geometry encoded in $(\cdot)(M)(X)$ as *operator-induced structure* (OIS). The aforementioned descriptors in materials genome are thus the predictors in a linear model with OIS. Henceforth, we will use *descriptor selection* and variable selection in the presence of OIS interchangeably. Note that the specification of $(\cdot)$ depends on domain knowledge, and we intentionally include the absolute difference operator $| \cdot |$ in $(\cdot)$ because it is directly interpretable in materials science, and it often provides clear intuition on many physical phenomena, such as the metal-oxide binding

energy (Liu et al. 2022). Treating it as a single operator reduces the required number of iterations to generate related descriptors.

A common practice in materials genome (O'Connor et al. 2018; Ouyang et al. 2018; Liu et al. 2020) is to employ modern statistical variable selection developed for linear models. However, the geometry of OIS defined by operators O and high-order compositions induces high correlation and ultra-high dimension to the feature space. As detailed in Section 2, the dimension of $O(M)(X)$ increases *double exponentially* with M and the number of binary operators in O . For example, with $p = 59$ in our real data application, enumerating $O(3)(X)$ gives 1.01×10^{17} predictors while only a handful of them are associated with the response. Moreover, these predictors are highly correlated as a result of iteratively applying unary operators. This along with a small size such as $n = 91$ in our real data application substantially hinders the performance of existing methods that rely on linear variable selection methods. Indeed, materials genomes are an analog concept to genomes, but the dimension of predictors and inherent strong correlation in materials genome-wide association studies, or *materials GWAS*, pose unprecedented challenges to statistical analysis.

In this article, we aim to develop a powerful method for materials GWAS in which we effectively identify materials genes that are associated with the response of interest. To achieve dimension reduction in materials GWAS, we consider an iterative approach by applying a small set of operators and immediately identifying the relevant descriptors, D , before constructing more complex descriptors. This step is iterated in light of the composition structure in OIS, in striking contrast to existing literature that aims to exhaustively generate $O(M)(X)$. In each iteration, $O(D)$ is typically substantially smaller than $O(M)(X)$, and such sparsity achieves dimension reduction and tackles the daunting computational challenges posed by materials GWAS.

Iterative dimension reduction, however, faces two intertwined challenges. First, the constructed descriptors in intermediate steps, unlike the astronomically large $O(M)(X)$, are not necessarily linearly associated with the response. To address this, we propose to use *nonparametric variable selection* for dimension reduction to ensure selection accuracy under the geometry of OIS. That is, while the model is assumed to be linear, we employ nonparametric variable selection to achieve dimension reduction while maintaining high selection accuracy. We refer to this key novelty of our proposed method as "parametrics assisted by nonparametrics": or PAN.

The second challenge pertains to the selection of the nonparametric module. Unlike traditional nonparametric variable selection, OIS variable selection calls for new performance criteria for the nonparametric module to ensure OIS selection accuracy (see Section 2.3 for more details). We introduce a *PAN criterion* to reassess nonparametric selection methods, which elucidates an asymmetric effect between false positives and false negatives and highlights a desired invariance property to unary transformations. In particular, we propose to use a Bayesian additive regression tree variable selection method, BART-G.SE (Bleich et al. 2014), as the nonparametric module, which we show is well suited to satisfy the PAN criterion.

Coupling the PAN strategy with BART-G.SE, together with additional considerations to address the complexities of materials GWAS, leads to a new method for materials GWAS, which we

call iterative BART, or iBART. The iterative framework of iBART reduces the size of the effective descriptor space significantly, mitigating collinearity in the process, and the use of nonparametric variable selection accounts for structural model misspecification in intermediate variable selection steps. Our extensive experiments show that iBART gives excellent performance with accuracy and scalability that are not seen in existing methods. Note that iBART is not a new BART variant for nonparametric regression, but rather an iterative use of BART within the PAN framework specifically tailored for materials GWAS.

The outline of the article is as follows. Section 1.1 reviews related work in materials genome. In Section 2, we introduce the OIS framework, describe the PAN selection procedure, and discuss how to choose the nonparametric module in PAN and some practical considerations regarding PAN. Section 3 contains a simulation study that shows superior performance of iBART relative to existing methods. In Section 4, we apply iBART to a single-atom catalysis dataset and it identifies physical descriptors that explain the binding energy of metal-support pairs with high explanatory power, leading to interpretable insights to guide the prevention of a notorious problem called sintering. We close in Section 5 with a discussion. All proofs, details of variants of iBART, and additional simulation results and discussion are deferred to the supplementary material.

1.1. Related Work

Descriptor selection has attracted growing attention in materials science. Recent methods often build on a *one-shot* descriptor generation and selection scheme followed by modern statistical variable selection approaches (O'Connor et al. 2018; Ouyang et al. 2018; Liu et al. 2020). In particular, they first construct descriptors by applying operators iteratively M times on the primary feature space $X^{(0)} = X = (x_1, \dots, x_p) \in \mathbb{R}^{n \times p}$ to construct an ultra-high dimensional descriptor space $X^{(M)} = O(M)(X)$ of $O(p^M)$ descriptors, assuming binary operators are used in each iteration. Then variants of generic statistical methods are adopted to select variables from $X^{(M)}$. Along this line, the method SISSO (Sure Independence Screening and Sparsifying Operator) proposed by Ouyang et al. (2018) builds on Sure Independence Screening, or SIS (Fan and Lv 2008), O'Connor et al. (2018) uses LASSO (Tibshirani 1996), and Liu et al. (2020) adopts Bayesian variable selection methods.

SISSO is widely perceived as one of the most popular methods for materials genome. It uses SIS to screen out P descriptors, from which the single best descriptor is selected using an l_0 -penalized regression. If a total of k descriptors are desired, this process will be iterated for k times yielding k sets of SIS-selected descriptors, followed by an l_0 -penalized regression to select the best k descriptors from all the SIS-selected descriptors. Note that in each iteration, SIS is employed to screen out P descriptors from the remaining descriptor set with an updated response vector given by its least squares residuals projected onto the space spanned by previously SIS-selected descriptors. Users must define the composition complexity of the descriptors through M , that is, the order of compositions of operators. In a typical application of SISSO, the composition complexity M is no greater than 3, the number of candidate descriptors in each SIS iteration is less than 100, and the number of descriptors k

is no larger than 5 (Ouyang et al. 2018). Note that selecting five descriptors with 100 SIS-selected descriptors in each iteration amounts to fitting at most $\binom{100}{5} \approx 2.6 \times 10^{11}$ different regressions, which is computationally intensive.

A major drawback of these one-shot descriptor construction procedures is the introduction of a highly correlated and ultra-high dimensional descriptor space $\mathbf{x} \prec M$. High correlation often hampers the performance of a variable selection method, and the ultra-high dimensional descriptor space with large p , as common in modern applications, can make it computationally prohibitive for such methods. In practice, these methods often resort to ad hoc adaptation or sacrifice the complexity level of candidate descriptors.

Another thread of work is the well-developed *automatic feature engineering* in machine learning, aiming to generate complex features from given constructor functions adaptively (Markovitch and Rosenstein 2002; Feurer et al. 2015; Khurana, Samulowitz, and Turaga 2018). However, the overwhelming focus of this literature is on increasing the predictive power of the primary features $\mathbf{X}$. We instead focus on discovering the underlying functional relationship between the response and the predictors and revealing data-driven insights into the underlying physics of materials design. In addition, the sample size in materials genome is typically limited, hampering the use of machine learning methods that rely on large training data.

In statistics, transformations have been commonly used to expand the predictor space, including polynomials, logarithmic, power transformations, and the previously noted interactions. The induced feature spaces from these elementary transformations are often overly simple to capture the intricate dynamics of the response in materials genome, particularly compared to high-order compositions of a larger operator set. There has been a rich literature on nonparametric variable selection, but descriptor selection relies on a linear model with feature engineering that favorably points to interpretable insights for domain experts as the functional forms of selected variables are explicitly given and the feature space could be composed using domain-related knowledge. The nonparametric module in PAN only serves as a dimension reduction tool, and the desired performance calls for new investigation under the context of OIS. Overall, materials GWAS may play an analogous role that GWAS have played in motivating new statistical methods and concepts, and to the best of our knowledge, the present article is the first statistical work on this topic.

2. The OIS Framework

2.1. Operator-Induced Structural Model

We begin with a standard ultra-high dimensional linear regression model

$$\mathbf{y} = \beta_0 + \beta_1 \mathbf{x}_1 + \cdots + \beta_{p^*} \mathbf{x}_{p^*} + \boldsymbol{\varepsilon}, \quad (2)$$

where $\boldsymbol{\varepsilon} \sim \mathcal{N}(0, \mathbf{I})$ is a Gaussian noise vector, and the regression coefficients β are sparse. The dimension of predictors p^* is ultra-high, at the materials GWAS scale that typically exceeds

the maximum size of matrices allowed by a modern personal computer, while the sample size n is on the order of tens. In

this article, we assume that this seemingly ultra-dimensional descriptor space obeys an operator-induced structure (OIS). In particular, we assume that the predictors $x_1, \dots, x_{p^*}$ in (2), or *descriptors*, are generated by applying operators in $\mathcal{O}$ iteratively M times on a primary feature space $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_{p^*}) \in \mathbb{R}^{n \times p}$,

$$\begin{aligned} (x_1, \dots, x_{p^*}) &= \mathbf{X}(M) = \mathbf{o}(M)(\mathbf{X}) = \mathbf{O} \circ \mathbf{o}(M-1)(\mathbf{X}) \\ &= \mathbf{O} \underbrace{\circ \cdots \circ}_{M \text{ times}} \mathbf{O}(\mathbf{X}), \end{aligned} \quad (3)$$

where $\mathbf{x} \prec M = \mathbf{o}(M)(\mathbf{X})$ denotes M -composition of $\mathbf{O}$ on $\mathbf{X}$ and can be defined iteratively as above. The operator set $\mathcal{O}$ is user-defined; for concreteness, we focus on the common example given in (1), unless stated otherwise.

We adopt the following convention. Evaluation of operators on vectors is defined to be entry-wise, for example, $\mathbf{X}_i = (x_{i,1}, \dots, x_{i,p^*})^T$ and $\mathbf{X}_1 + \mathbf{X}_2 = (x_{1,1} + x_{2,1}, \dots, x_{n,1} + x_{n,2})^T$. Throughout this article, we assume that all descriptors in $\mathbf{x} \prec M$ are uniquely defined in terms of their numerical values. For instance, only one of the descriptors in $\{x_i, x_j \mid x_i \times x_j\}$ will be kept in $\mathbf{x} \prec M$. This can be easily achieved in practice by identifying and removing perfectly correlated descriptors.

We hereafter refer to the linear regression model in (2) along with the operator-induced structure (OIS) in (3) as the OIS model. To facilitate a precise OIS model definition using predictors with nonzero coefficients, we define *M-composition descriptor* as follows.

Definition 2.1 (M-composition descriptor). We define $\mathbf{J} \prec M(\mathbf{X})$ to be an M -composition descriptor if it is constructed via M compositions of operators on some primary features $\mathbf{X}$: $\mathbf{J} \prec M(\mathbf{X}) = \mathbf{o} \prec M(\mathbf{X}) = \mathbf{O} \mathbf{O} \mathbf{O} \prec M-1(\mathbf{X}) = \mathbf{O} \mathbf{O} \mathbf{O} \prec M-1 \circ \cdots \circ \mathbf{O}(\mathbf{X})$, where $\mathbf{O}_m \in \mathcal{O}$ is the m th composition operator(s) for $1 \leq m \leq M$, and $\mathbf{f}(1)(\mathbf{X}), \dots, \mathbf{J} \prec M-1(\mathbf{X})$ are the necessary intermediate descriptors for constructing the descriptor $\mathbf{J} \prec M(\mathbf{X})$.

Note that if the m th composition operator is a binary operator, there may exist two $(m-1)$ th composition operators but we suppress the notation in the definition above for simplicity. Furthermore, if an M -composition descriptor $\mathbf{J} \prec M(\mathbf{X})$ only depends on a subset of primary features $\mathbf{X}_S$, where $S \subseteq [p] = \{1, \dots, p\}$, we also write it as $\mathbf{J} \prec M(\mathbf{X}_S)$ and call it an (M, S) -descriptor.

Definition 2.2 ((M, S)-composition OIS model). An (M, S) -OIS model assumes

$$\mathbf{Y} = \beta_0 \mathbf{1} + \sum_{k=1}^K \mathbf{L}_k \mathbf{f}_k(\mathbf{M}_k(\mathbf{X}_{S_k})) + \boldsymbol{\varepsilon}, \quad (4)$$

where $M = \max_{k=1, \dots, K} M_k$ denotes the highest order of operator compositions, K denotes the number of additive descriptors, $\mathbf{X}_{S_k}$ is the set of primary features used in the k th descriptors, and $S = \bigcup_{k=1}^K S_k$ is the set of all active primary feature indices.

Throughout the article, we assume the data follows an (M, S) -OIS model in (4). We use M_0 for the oracle highest composition complexity and S_0 for the oracle set of active primary feature indices. Descriptors in the (M, S) -OIS model and their intermediate descriptors are called *active*. We next use

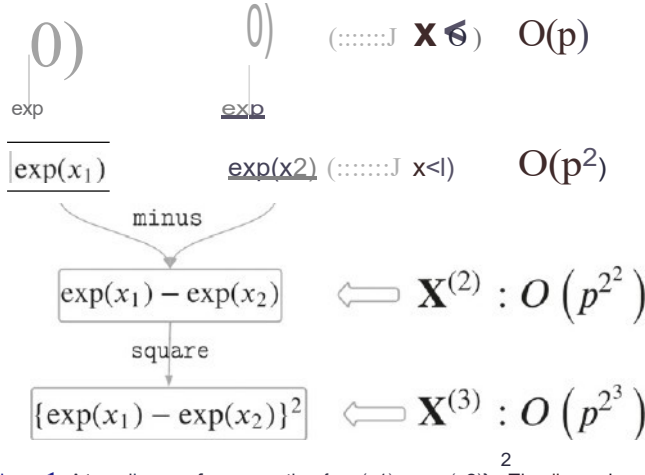


Figure 1. A tree diagram for generating $\{\exp(x_1) - \exp(x_2)\}^2$. The dimension of descriptor space increases double exponentially with the composition complexity.

a toy example to illustrate the introduced concepts in OIS and the challenges posed by descriptor selection. Suppose that the data-generating model is

$$y = \beta_0 + \beta_1 f_1^{(M_1)}(X_{S_1}) + \beta_2 f_2^{(M_2)}(X_{S_2}) + \epsilon,$$

where $M_1 = 3, M_2 = 2, S_1 = \{1, 2\}, S_2 = \{3, 4\}, j_1(M_1(X)) = \{\exp(x_1) - \exp(x_2)\}^2$, and $j_2(M_2(X)) = \sin(\exp(x_3 x_4))$. Here $\{\exp(x_1) - \exp(x_2)\}^2$ is a 3-composition descriptor or a $(3, \{1, 2\})$ -descriptor, and $\sin(\exp(x_3 x_4))$ is a 2-composition descriptor or a $(2, \{3, 4\})$ -descriptor. Both descriptors arise from applying O iteratively three times on the primary features: $\mathbf{X}^{(3)} = O \circ O \circ O(\mathbf{X})$. The composition of operators resembles a tree-like structure for generating descriptors; Figure 1 describes the tree-like workflow for generating $\{\exp(x_1) - \exp(x_2)\}^2$.

To see how the descriptor space dimension increases with the number of iterations, let C_u and C_b denote the number of unary and binary operators, respectively, and let p_j denote the dimension of the j th descriptor space $\mathbf{X}^{(j)}$. Note that the dimension of $\mathbf{X}^{(1)}$ is $p_1 = C_u P + C_b P(P-1)/2$, which is on the order of $O(p^2)$; the dimension of $\mathbf{X}^{(2)}$ is $p_2 = C_u p_1 + C_b p_1(p_1 - 1)/2$; $\dots$; $O(p^{2 \cdot 2}) = O(p^4)$; the dimension of $\mathbf{X}^{(3)}$ is $p_3 = C_u p_2 + C_b p_2(p_2 - 1)/2$; $\dots$; $O(p^{2 \cdot 2 \cdot 2}) = O(p^8)$. The dimension of descriptor space will increase double exponentially with the number of binary operator compositions, for example, with M compositions of binary operators on X the resulting descriptor space has a dimension of order $O(p^{2^M})$. Similarly, the double exponential expansion applies to the number of binary operators C_b . Excluding redundant descriptors does not prevent this double exponential expansion. As shown in Section 3.4, building $\mathbf{X}^{(2)}$ from $p = 59$ primary features results in an astronomical descriptor space containing over $5.5 \times 10^7 = 0(59^{22})$ descriptors even after removing redundant and nonphysical descriptors. In addition, the number of active (intermediate) descriptors will be nonincreasing with M as shown in Figure 1: there are four active descriptors $\exp(x_1), \exp(x_2), x_3, x_4$ in $\mathbf{x}^{(1)}$, but only two active descriptors $\exp(x_1) - \exp(x_2)$ and $\exp(x_3 x_4)$ in $\mathbf{X}^{(2)}$, and two active descriptors $\{\exp(x_1) - \exp(x_2)\}^2$ and $\sin(\exp(x_3 x_4))$ in $\mathbf{X}^{(3)}$.

2.2. PAN Descriptor Selection for OIS Model

We propose an iterative descriptor construction and selection procedure PAN for the OIS model, which generates descriptors by iteratively applying operators and selecting the potentially useful intermediate descriptors between each iteration of descriptors synthesis. The iterative descriptor selection procedure excludes irrelevant intermediate descriptors from the descriptor generating step, achieving a *progressive* variable selection and enabling variable selection based on ab initio primary features. This reduces the dimension of the subsequent descriptor space $\mathbf{x}^{(m)}$ and mitigates collinearity among the descriptors in comparison to the one-shot descriptor construction approaches.

To describe the method in its most general form, we allow different sets of operators O_m, S, O for each iteration $m = 1, \dots, M$, leading to the descriptor space

$$O_M \circ O_{M-1} \circ \dots \circ O_2 \circ O_1(\mathbf{X}). \quad (5)$$

The framework of our iterative descriptor selection procedure is as follows.

PAN descriptor selection procedure:

1. **Repeat** the following until at least one descriptor exhibits a strong linear association with the response variable y ($i = 0$):
 - (a) Use a nonparametric variable selection procedure to perform descriptor selection on $\mathbf{x}^{(i)}$ and obtain the selected descriptors $\mathbf{x}^{(i)}$;
 - (b) Apply the i th operator set O_i on all of the previously selected descriptors, $\text{Um}(\mathbf{x}^{(i)})$, yielding a new descriptor space, $\mathbf{x}^{(i+1)} = O_i(\text{Um}(\mathbf{x}^{(i)}))$, where O_i can be different for each iteration i ;
2. Once there exist descriptor(s) that exhibit a strong linear association with response variable y , use a linear parametric variable selection procedure to perform descriptor selection on $\mathbf{x}^{(i)}$, and obtain the selected descriptors, $\mathbf{X}^* \subseteq \mathbf{X}^{(i)}$.

We keep all the selected descriptors in the main loop to facilitate the creation of high-order complexity descriptors with the help of low-order complexity descriptors. For instance, constructing $\mathbf{x}^{(i+1)}$ using O defined in (1) requires us to keep $\mathbf{x}_1 \in \mathbf{X}^{(0)}$ selected at the base iteration and $\mathbf{x}_i \in \mathbf{x}^{(i)}$ selected at the first iteration.

To see how this iterative procedure helps reduce the dimension of descriptor space significantly, let $s_i = |\mathbf{X}^{(i)}|$ be the number of descriptors selected in the i th iteration and $p_i = |\mathbf{X}^{(i)}|$ be the dimension of the i th descriptor space. Suppose that the number of selected descriptors is sparse in each iteration, that is, $s_i \ll p_i$. Assuming binary operators were used, the dimension of the $(i+1)$ th descriptor space in PAN is on the order of $O(s_i^2) \ll O(p_i^2)$, and this holds for all iteration $i \geq 0$. If we further assume $s_i \ll O(p_i)$ for all $i \geq 0$, where $p = |\mathbf{X}|$ is the number of primary features, then the dimension of the $(i+1)$ th descriptor space for PAN is on the order of $O(p^2)$, compared to $O(p^{2^{i+1}})$ for the one-shot methods. Note these assumptions are reasonable according to the discussion in Section 2.1. In the simulation study in Section 3.4 with $p = 10$ primary features, we observed that $s_i \ll O(10^1)$ and $p_i \ll O(10^2)$ for all iterations of

the PAN procedure. On the contrary, a one-shot method, such as SISSO, generates a descriptor space containing 9.26×10^9 descriptors in the same setting.

The use of a nonparametric variable selection procedure in Step 1(a) is necessary because the intermediate descriptors may not have a strong linear association with the response variable. Thus, a method that accounts for model misspecification, such as a nonparametric method, is more suitable for preliminary screening of the intermediate descriptors. In addition, a suitable nonparametric module for PAN needs to account for the geometry embedded in OIS and the unconventional goal of selecting operators along with variables; the next section discusses performance requirements for this step and introduces an implementation of PAN, iBART, that is particularly suitable for OIS. In Step 2, the oracle descriptors are linearly associated with the response. Hence, we employ linear parametric variable selection methods, such as LASSO (Tibshirani 1996), to reduce false positives and select the final descriptors.

2.3. Choosing Nonparametric Module in PAN and iBART

In OIS, the ultimate goal is to recover or approximate the true functional relationship between the response and the primary features. This goal together with PAN entails new performance requirements for the nonparametric module in PAN. To illustrate such a need, let us consider a simple OIS model

$$y = f^{(M)}(X) + \varepsilon = \sqrt{x_1 + x_2} + \varepsilon. \quad (6)$$

In traditional nonparametric variable selection problems, the regression model only sees the primary features, X , and the desired performance is successful identification of the active index set, $S_0 = \{1, 2\}$. The base iteration of PAN has a similar goal as we would want the selected index set S to be a superset of S_0 . However, the intermediate descriptor space $x^{(m)}$ in the m th iteration consists of nonlinear transformations of the selected primary features X_8 , and the active index set is no longer well-defined.

Due to the iterative structure of PAN, a good nonparametric selection method for PAN must be able to generate and identify the m -composition descriptors $f^{(m)}(X)$ at the m th iteration (i.e., all active intermediate descriptors). To this end, a suitable nonparametric module should satisfy the following *PAN criterion*:

It selects *all* of the m -composition descriptors that are necessary for constructing the true M -composition descriptor, for all $0 \leq m \leq M$ iterations.

Taking model (6) as an example, failing to select either $f_1(x) = x_1$ or $f_2(x) = x_2$ in the base iteration will preclude the generation of $f^{(1)}(x) = (x_1 + x_2)$ and $j^{(2)} = x_1 + x_2$ in subsequent iterations. The PAN criterion thus favors a "conservative" nonparametric method, in which false positives during selection are allowed but false negatives *must not* occur in any iteration. Here false positives are defined within each intermediate descriptor space $x^{(m)}$.

The literature has provided a rich menu of nonparametric selection methods; however, the asymmetric effect of false positives and false negatives on descriptor selection illustrated

by the PAN criterion motivates our choice of tree-based nonparametric approaches. To see this, suppose the true descriptor is $f(x_1) = \log(x_1)$ and the design matrix $O(X)$ consists of I -composition transformations of the primary features $X = (x_1, \dots, x_p)$. A typical nonparametric method aims to identify the oracle primary predictors (x_i in this case) that associate with y through an unknown function $f(\cdot)$, without considering transformations of x_i as possible candidate predictors. However, the goal in the presence of OIS is to identify the true primary predictors (x_i) and the correct operator composition ($\log(\cdot)$). This goal, in the presence of many non-signal but highly correlated unary transformations of x_1 , namely, $\sqrt{x_1}$, $\log(x_1)$, etc., is shown to be difficult for many nonparametric methods; see supplementary material Section A.2.1 for detailed analyses. Tree-based methods are invariant to monotonic transformations and thus tend to be robust to related transformations that are often piecewise monotonic. Consequently, they may select unary transformations of x_1 in addition to $f(x_1) = \log(x_1)$ —such false positives, although increasing the candidate search space, are favorably compatible with the PAN criterion. We next provide a review of BART (Chipman, George, and McCulloch 2010) and describe BART-G.SE (Bleich et al. 2014)—the default nonparametric module in the PAN framework.

BART is a Bayesian nonparametric ensemble tree method for modeling $y = f(X) + e$, the unknown relationship between a response vector y and a set of predictors $x_1, \dots, x_p$. More specifically, BART models the regression function f by a sum of regression trees

$$Y = \sum_{i=1}^T g_i(x_1, \dots, x_p; T_i, I_{i,:}) + \varepsilon, \quad \varepsilon \sim \mathcal{N}(\mathbf{0}, \sigma^2 I_n).$$

Each binary regression tree g_i consists of a tree structure T_i partitioning observations into B ; terminal nodes, and a set of terminal parameters $I_{i,:} = \{\mu_{1i}, \dots, \mu_{Bi}\}$ attached to these nodes. Observations within a given terminal node b are constrained to have the same terminal parameter μ_b . The prior distributions for $(T_i, I_{i,:})$ constrain each tree to be small so that each tree contributes to approximate f in a small and distinct fashion. Readers are referred to Chipman, George, and McCulloch (2010) for the full details of BART and posterior sampling.

The primary usage of BART is prediction, and the predicted values $\hat{y}_i$ for y serve no purpose in variable selection. However, a variable inclusion rule can be defined based on the variable inclusion proportion q_i of x_i , which can be easily estimated from the posterior samples. To this end, we adopt the permutation-based selection threshold based on the permutation null distribution of $q = (q_1, \dots, q_p)$ proposed by Bleich et al. (2014).

Specifically, B permutations of the response vector $y_i, \dots, y_n$ are generated, and a BART model is fitted for each of the permuted response vectors with the same predictors $x_1, \dots, x_p$. The variable inclusion proportions from the permuted BART models $q_1, \dots, q_i$ are then used to create a permutation null distribution for the non-permuted variable inclusion proportion q . The predictor x_i is selected if $q_i > m_i + C^* \cdot s_i$, where m_i and s_i are the mean and standard deviation of the permuted variable inclusion proportion $q_i = (q_{i1}, \dots, q_{iB})$, and $C^* = \inf_{C \in \mathbb{R}^+} \{V_i, f_{II} = 1 H(q_{iB} \dots m_i + C \cdot s_i) > 1 - \alpha\}$ is the smallest global standard error multiplier (G.SE) that gives

a simultaneous 1 - α coverage across the permutation null distributions of q_i for all predictors. We refer to BART with the permutation-based selection procedure described above as BART-G.SE.

Various BART-related methods have been recently developed (Linero 2018; Horiguchi, Pratola, and Santner 2021; Liu, Rockova, and Wang 2021). They often incorporate sparsity-inducing priors into BART and prove to be highly effective in various tasks. However, it is unclear whether the excellent performance of them developed in traditional settings carries over to being compatible with the PAN criterion. Indeed, we have found that methods aiming at optimally choosing non-parametric variables in traditional settings may incur fewer false positives but have a higher chance to miss the true descriptors in intermediate iterations—such false negatives are devastating in the context of PAN and OIS variable selection. The PAN criterion provides a useful guide in choosing not only the regression method but also the selection rule, for which we recommend BART-G.SE. In our numerical experiments, we vary the non-parametric module in PAN by comparing several recent BART-related methods and other nonparametric selection methods, and find that the proposed iBART, PAN with BART-G.SE as the nonparametric module, is particularly well suited for OIS variable selection and tends to give the best overall performance; see the supplementary material for numerical results and a comprehensive discussion.

2.4. Practical Consideration and the Algorithm

The operators in O in (1) can be classified into the unary operators $O_u = \{J, \exp, \log, j \cdot J, \cdot 1, ^2, \sin(\cdot), \cos(\cdot)\}$ and the binary operators $O_b = \{+, -, \times, /, 1 - j\}$, each posing different challenges to descriptor selection. The unary operators O_u induce strong collinearity among the engineered descriptors; for instance, $\text{cor}(x, jx) > 0.9$ when $x_1 \sim N(0, 1)$ with $n = 200$. The binary operators O_b increase the descriptor space dimension double exponentially and generate complex nonlinear descriptors. These two issues are compounded when the two operator sets are used together. Therefore, we propose to decouple the two operator sets and alternate them, leading to two special cases of (5):

$$O_{A_u}^{(M)}(X) = \underbrace{\cdots O_b \circ O_u \circ O_b \circ O_u}_{M \text{ times}}(X), \quad (7)$$

$$O_{[;]}(X) = \underbrace{\cdots O_u \circ O_b \circ O_u \circ O_b}_{M \text{ times}}(X) = \text{ot} \rightarrow O_b(X). \quad (8)$$

In addition to the binary operators identified earlier, we include an additional binary operator, $\text{rr}_1 : \mathbb{R}^n \rightarrow \mathbb{R}^n$ defined by $\text{rr}_1(a, b) = a$, which allows intermediate descriptors to be passed down unchanged. Note the two alternating descriptor spaces $O_{[;]}(X)$ and $O_{A_u}(X)$ are not equivalent to the full descriptor space $O_{<M>}(X)$ with the same composition complexity M . However, one can show that $O_{[;]}(X)$ and $O_{A_u}(X)$ can recover $O_{<M>}(X)$ with some $M_u > M$ and $M_b > M$, respectively. This is formally described in the theorem below and its proof is available in supplementary material Section D.

Theorem 2.1. Let $X = (x_1, \dots, x_p) \in \mathbb{R}^{n \times p}$ be a primary feature space and O be a set of operators such that it can be partitioned into a unary operator set O_u and a binary operator set O_b . Suppose that $I \in O_u$ and $\text{rr}_1 \in O_b$. Then for any $M \in \mathbb{N}$, there exists M_u and M_b such that $O_{<M_u>}(X) \supseteq O_{<M>}(X)$ and $O_{<M_b>}(X) \supseteq O_{<M>}(X)$, respectively.

For instance, the 2-composition descriptor space $O_{<2>}(X)$ generated using (3) contains descriptors such as $f_1 = (x_1 + x_2)$ and $f_2 = (x_1 + x_2)^2$. These two descriptors can be generated using $O_{[;]}(X)$ in (8): $f_1 = \text{add}(\text{square} \circ \text{rr}_1(x_1, x_2))$, $\text{square} \circ \text{rr}_1(x_1, x_2)$ or $f_2 = \text{square} \circ \text{add}(x_1, x_2)$. Using (7), f_1 and f_2 can also be generated as $f_1 = \text{add}(\text{square}(x_1), \text{square}(x_2))$ or $f_2 = \text{square}(\text{add}(x_1, x_2))$. As such, under the M -composition OIS model in (4), one may consider using $O_{<M>}$, $O_{<M_u>}(X)$, or $O_{<M_b>}(X)$ as long as they contain the true descriptors. In what follows we will focus on $O_{[;]}(X)$ and $O_{A_u}(X)$ instead of $O_{<M>}(X)$ for their aforementioned advantages.

We adopt the descriptor generating strategies in (7) and (8) for different scenarios. In particular, if the primary features X are believed to generate the model through their intricate interactions that will be captured by binary operators and compositions of such binary operators, we recommend (8). This is often the case in real-world applications, such as in Section 4, where the domain scientists have chosen a large set of potentially useful primary features, and unary transformations of these primary features are less interpretable and thus less desired. If such prior knowledge is not available and the relevant functional form of primary features is unknown, as in Section 3, we would recommend (7) to first identify such functional forms by selecting between unary descriptors.

The stopping criterion in Step 1 can be easily modified depending on practical needs. For example, we can specify the maximum composition complexity $M_{\max}$, like SISSO, or implement a data-driven criterion that terminates Step 1 when there exists a descriptor such that its absolute correlation with response variable y exceeds a pre-specified threshold $P_{\max}$ that is close to 1. We note that iBART allows larger $M_{\max}$ than SISSO as iBART does not rely on one-shot feature engineering, and the use of $P_{\max}$ allows early termination.

It is also common in practice that one may only want to select k descriptors for easy interpretation, such as $k \leq 5$. If the cardinality of the selected descriptors X^* is greater than k , then an ℓ_0 -penalized regression may be used to choose the best k descriptors.

We allow user-specified options for these considerations when implementing iBART, summarized in Algorithm 1. The following default settings will be used in our experiments. We use the BART-G.SE thresholding variable selection procedure implemented in the R package `bartMachine` to perform nonparametric variable selection for Step 1(a) in PAN and use LASSO implemented in the R package `glmnet` to perform parametric variable selection for Step 2 in PAN. For BART-G.SE, we set the number of trees to 20 that is recommended by Bleich et al. (2014), the number of burn-in samples to

10,000, the number of posterior samples to 5000, the number of permutations of the response to 50, and the rest of the parameters to the default values. With these values, our Markov chain Monte Carlo (MCMC) runs appear to have reached a sufficient number of iterations in our experiments. We choose the penalty term λ in LASSO by minimizing the mean squared error loss through a 10-fold cross-validation procedure and set the other parameters to the default values. Unreported results using real data indicate that other tuning methods to debias LASSO yield similar performance. The algorithm terminates if the composition complexity M reaches $M_{\max} = 4$ or the maximum absolute correlation $|P|$ reaches $P_{\max} = 0.95$, whichever occurs first. Setting $M_{\max}$ to 4 appears sufficient for the considered materials GWAS application as descriptors with more complexity become challenging to interpret.

Algorithm 1: iBART

Input: $P_{\max} \in (0, 1)$ = maximum absolute correlation with the response variable;
 $M_{\max}$ = maximum composition complexity;
 L_{zero} = whether to perform fa-penalized regression;
 k = number of selected descriptors by fa-penalized regression. Only required when $L_{\text{zero}} == \text{TRUE}$.
Output: $\mathbf{XZ}$: selected descriptors
Data: $\mathbf{X}$: primary features; $\mathbf{Ou}$: set of unary operators; $\mathbf{Ob}$: set of binary operators; $\mathbf{y}$: response vector
 $M = 0$
 $p = \max_{\mathbf{X}} \text{EX}(M > \text{cor}(\mathbf{x}, \mathbf{y}))$
while $M < M_{\max}$ **or** $|P| < P_{\max}$ **do**
 $\mathbf{x} < M$ $\leftarrow$ BART-G.SE selected descriptors on $\mathbf{x} < M$
 $\mathbf{x} < M+1$ $\leftarrow$ OM+1 ($\mathbf{u}; \mathbf{x}; >$)
 $M \leftarrow M+1$
 $p \leftarrow \max_{\mathbf{X}} \text{EX}(M > \text{cor}(\mathbf{x}, \mathbf{y}))$
end
 $\mathbf{X}^* \leftarrow$ LASSO selected descriptors on $\mathbf{x} < M$
if $L_{\text{zero}} == \text{TRUE}$ **and** $|\mathbf{X}^*| > k$ **then**
 $\mathbf{XZ} \leftarrow$ best k descriptors from fa-penalized regression
end
else
 $\mathbf{XZ} \leftarrow \mathbf{X}^*$
end

3. Simulation

3.1. Outline

In Sections 3.2 and 3.3 we demonstrate that the employed BART-G.SE tends to satisfy the PAN criterion when selecting unary and binary operators, respectively, and evaluate its false positives. Section 3.4 assesses iBART relative to several existing descriptor selection methods in view of the PAN criterion and OIS variable selection accuracy using a complex simulation setting. Section 3.5 showcases the robustness of iBART in the initial input dimension p . Section 3.6 examines the performance of iBART when the operator set $\mathcal{O}$ can not generate the ground truth model. We use Algorithm 1 and the default settings described in Section 2.4 when implementing iBART, unless stated other-

wise. Each simulation is replicated 100 times. We also compare variants of iBART by varying the nonparametric module; see the supplementary material for details.

3.2. Unary Operators

We consider all unary transformations of five primary features $\mathbf{X} = (x_1, \dots, x_5)$, that is, the descriptor space is $\mathbf{Ou}(\mathbf{X}) = \bigcup_{i=1}^5 \{x_i, x_i^{-1}, x_i^2, jx_i, \log(x_i), \exp(x_i), \text{lxd}, \sin(jx_i), \cos(., x_i)\}$. For each unary operator $u_j \in \mathbf{Ou}$, we generate the response vector by

$$\mathbf{y} = 10u_j(\mathbf{x}_1) + \boldsymbol{\varepsilon}, \quad \boldsymbol{\varepsilon} \sim \mathcal{N}_n(\mathbf{0}, \mathbf{I}), \quad (9)$$

with sample size $n = 200$, yielding nine independent models in total. We draw the primary features $\mathbf{X}$ independently from the standard normal distribution if the domain of u_j is $\mathbb{R}$, and the Lognormal(2, 0.5) distribution if the domain is $\mathbb{R}^+$.

The left plot in Figure 2 shows the number of true positives (TP) of BART-G.SE. For each of the nine models in (9), the true descriptor is selected 100/100 times. This suggests that BART-G.SE is capable of identifying the true descriptor with high probability when the descriptor space is populated with unary operators $\mathbf{Ou}$. Other nonparametric methods did not select all TP 100/100 times for all nine scenarios, failing the PAN criterion; see the supplementary material for further results and discussion.

The left plot in Figure 2 shows the number of FP of BART-G.SE for each simulation setting. Although BART-G.SE is capable of capturing the true descriptor with high probability, it also selects some non-signal descriptors in some cases. The number of FP is especially high when the true descriptor is x_i , lxd , $\exp(x_1)$, $\log(x_1)$, x_1^2 , or jx_i . This is due to high collinearity among these six descriptors. In particular, the empirical Pearson correlation between $\log(x_1)$ and x_1^2 is over 0.99 under simulation setting (9) and thus some false positives are expected. Selecting inactive descriptors in one iteration is less of a concern in variable selection with OIS as they do not constitute misspecified models, and can be further screened out during Step 2 of PAN.

3.3. Binary Operators

In this simulation study, we apply binary operators $\mathbf{Ob} \{+, -, \times, /, 1 - 1, 7\}$ on the five primary features $\mathbf{X} (x_1, \dots, x_5)$. The descriptor space $\mathbf{Ob}(\mathbf{X}) \in \mathbb{R}^{200 \times 55}$ contains all binary transformations of all possible pairs of the five primary features. For each binary operator $b_j \in \mathbf{Ob}$, we generate the response vector by

$$\mathbf{y} = 10b_j(\mathbf{x}_1, \mathbf{x}_2) + \boldsymbol{\varepsilon}, \quad \boldsymbol{\varepsilon} \sim \mathcal{N}_n(\mathbf{0}, \mathbf{I}),$$

with sample size $n = 200$, yielding a total of six independent models. We generate the primary features $\mathbf{X}$ following $x_1, \dots, x_5 \sim \text{iid } \mathcal{N}(0, 1)$.

The right plot in Figure 2 shows the number of TP for each of the six models in (10). BART-G.SE is able to identify the true binary descriptor 100/100 times in all six settings, similar to earlier observations in Section 3.2. The right plot in Figure 2 shows that BART-G.SE may also select some irrelevant descriptors but

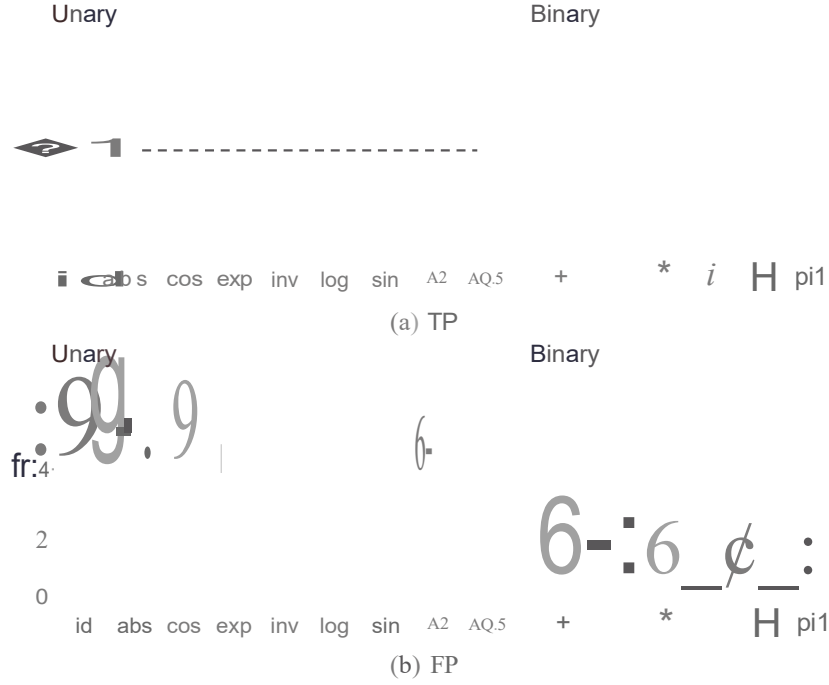


Figure 2. Boxplots of TP and FP over 100 simulation replicates under (9) and (10).

considerably fewer than in Section 3.2. In particular, when the true descriptor is $(x_1 - x_i)$, BART-G.SE correctly selects $(x_1 - x_2)$ but also selects $|x_i - x_j|$ 100/100 times. Unlike in Section 3.2, the inclusion of $|x_i - x_j|$ is not due to high correlation between $|x_i - x_2|$ and $(x_1 - x_2)$. In fact, their empirical Pearson correlation is merely 0.07. We suspect that $|x_i - x_2|$ was selected due to its piecewise monotonicity in $(x_1 - x_2)$ that tends to be invariant to tree-based methods. Similarly, BART-G.SE selects both $(x_1 - x_2)$ and $|x_i - x_j|$ with high probability when the true descriptor is $|x_i - x_j|$. When the true descriptor is $(x_1 + x_i)$, $x_1 + x_i$, x_i , x_1 , $x_i x_2$, or $rr_1(x_1, x_2)$, the FP of BART-G.SE is nearly zero. This shows that BART-G.SE is not too conservative when high collinearity does not exist.

Our investigation using unary and binary descriptors suggests that the permutation threshold of BART-G.SE tends to satisfy the PAN criterion without being too conservative, making it a good candidate for PAN and OIS variable selection. We next use a more complex simulation to assess iBART in view of the PAN criterion and OIS selection accuracy.

3.4. Complex Descriptors with High-Order Compositions

In this section, we compare iBART with existing approaches under a complex model and demonstrate superior performance of iBART. We use the following model

$$y = 15\{\exp(x_1) - \exp(x_i)\}i + 20\sin(rrX3X4) + e, \quad e \sim Nn(0, a^2I), \quad (11)$$

with $n = 250$, $p = 10$, and $a = 0.5$. Here the number of primary features p is set to a relatively small number since competing one-shot methods cannot complete the simulation when $p > 20$ due to the ultra-high dimension of their descriptor spaces. In Section 3.5, we demonstrate that iBART scales well

in p and gives a robust performance. We use the operator set CJ defined in (1) with rr_1 , and the primary feature vectors x_i are drawn independently from a uniform distribution, namely, $x_1, \dots, x_p \stackrel{iid}{\sim} Un(-1, 1)$. Section B of supplementary material demonstrates the effect of dependent primary features on iBART and other methods using the same functional relationship in (11).

iBART is implemented in R based on Algorithm 1 in Section 2.4 using the default settings. We chose the descriptor generating process (7) in Section 3 since the iid primary features do not show strong collinearity. Two versions of iBART are considered: iBART without and with l_0 -penalized regression, labeled as "iBART" and "iBART+ l_0 ", respectively. The l_0 -penalization finds the best subset of k variables from the set of selected variables using the Akaike information criterion (AIC) with $k \in \{1, 2, 3, 4\}$. We compare the performance of iBART and iBART+ l_0 with SISSO and LASSO. SISSO is implemented using the Fortran 90 program published by Ouyang et al. (2018) on GitHub with the following settings: the descriptor magnitude allowed in the descriptor space is set to $[1 \times 10^{-6}, 1 \times 10^5]$; the size of the SIS-selected subspace is set to 20; the operator composition complexity M is set to 3; the maximum number of operators in a descriptor is set to 6; and the number of selected descriptors $k \in \{1, 2, 3, 4\}$ is tuned by AIC. The R package glmnet is used to implement LASSO and the penalty parameter λ is tuned via 10-fold cross-validation to minimize the mean squared error loss. Since the size of $CJ^{(3)}(X)$ exceeds the limit of an R matrix, we give LASSO an advantage by reducing the descriptor space to $CJ_u \circ O \circ O(X)$ that aligns with the true data-generating process. We also use the l_0 -penalized regression step as in SISSO and iBART+ l_0 for LASSO, leading to LASSO+ l_0 .

For each method, we calculate the number of TP selections, FP selections, and false negative (FN) selections. We use the

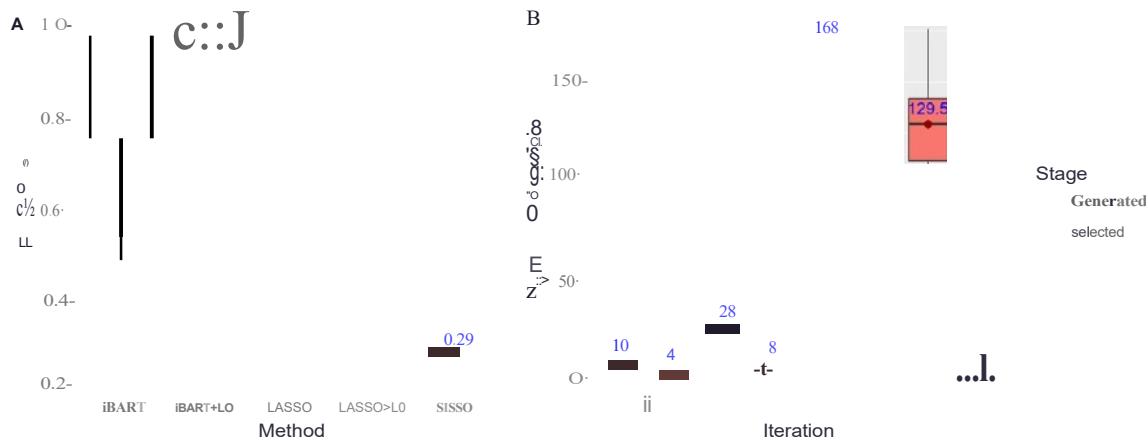


Figure 3. Left: boxplots of F1 scores over 100 simulations for different methods under Model (11). Right Boxplots of iBART generated and selected descriptors in each iteration.

F_1 score as an overall metric to quantify the performance of each method: $F_1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$, where Precision = $\frac{\text{TP}}{\text{TP} + \text{FP}}$ and Recall = $\frac{\text{TP}}{\text{TP} + \text{FN}}$. The value $F_1 = 1$ means correct identification of the true model without having any FP and FN. In this simulation, the two TPs are $\exp(x_1) - \exp(x_2)$ and $\sin(\exp(x_1) - \exp(x_2))$.

As shown in Figure 3A, both iBART-based methods achieve very high F_1 scores, with a median F_1 of 0.8 and 1 for iBART and iBART+ ℓ_0 , respectively. In particular, both iBART and iBART+ ℓ_0 have 2 TPs in all simulations while incurring an average FP of 0.93 and 0.3, respectively. This demonstrates that iBART satisfies the PAN criterion in all iterations as it is able to identify the 2 TPs 100/100 times in simulations. Furthermore, iBART+ ℓ_0 gives a very low FP, scoring a perfect F_1 score 75/100 times. LASSO-based methods and SISSO have lower F_1 scores than iBART and iBART+ ℓ_0 but for different reasons. LASSO selects 37.65 descriptors on average, and the 2 TPs are always selected. This means that LASSO also enjoys the PAN criterion but its F_1 score is hindered by the large number of FP. With the help of the ℓ_0 -penalized regression, LASSO+ ℓ_0 reduces the average number of FP from 35.65 to 2 while maintaining 2 TPs 100/100 times, substantially increasing the F_1 score to 0.67, the third-highest score but still considerably lower than iBART and iBART+ ℓ_0 . SISSO, on the other hand, has only 1 TP but 3 FPs in all simulations, which unfortunately does not satisfy the PAN criterion. In particular, SISSO can identify $\sin(\exp(x_1) - \exp(x_2))$ 100/100 times but always selects a false signal, $\exp(x_1) - \exp(x_2)$, a descriptor that has an absolute correlation over 0.9 with the TP, $\{\exp(x_1) - \exp(x_2)\}^2$. In summary, both iBART-based and LASSO-based methods satisfy the PAN criterion but LASSO-based methods incur more FPs. SISSO, however, fails to identify one of the two descriptors 100/100 times but selects a highly correlated counterpart, indicating its relatively weakened ability to distinguish the true descriptor when there are descriptors highly correlated with the TP. Results not reported here show that replacing BART-G.SE with LASSO in the PAN framework missed TPs with high probability even at the base iteration, indicating the importance of nonparametric variable selection in PAN.

To gain insight into the scalability of iBART, Figure 3(B) shows the boxplots of the number of iBART generated and

selected descriptors in each iteration. Throughout the 100 simulation replicates, iBART generates no more than $2p^2$ descriptors, which is significantly less than the number of descriptors generated by SISSO (9.26×10^9) and LASSO (1.2×10^6). Such dimension reduction not only reduces runtime and memory usage but also enables iBART to tackle data with much larger p , as we show in Sections 3.5 and 4.

3.5. Large p

In this section, we demonstrate that the performance of iBART is robust to increase in input dimension p while one-shot descriptor selection methods are not. Under Model (11) with $p = 20$, SISSO with the same parameter settings in the preceding section generates more than 3.8×10^{11} descriptors and failed to complete one simulation within 24 hr on a server with 1.3TB of memory and 40 CPU cores available to SISSO. LASSO, on the other hand, failed at $p = 20$ since the `glmnet` function in the R package `glmnet` cannot handle a matrix of size $250 \times (1.57 \times 10^7)$. The proposed method iBART instead scales well in the input dimension p partly because of the efficient dimension reduction via the PAN strategy. In particular, when $p = 20$, iBART took an average of 497 sec to finish, and its memory usage peaks out at 10.85 GB in 100 simulation replicates.

We implement iBART for $p \in \{10, 20, 50, 100, 200\}$ using the same simulation settings in Section 3.4. Figure 4 shows that iBART's F_1 score is robust to the increase in p . In fact, benefiting from the *ab initio* mechanism through the PAN strategy, the performance of iBART would be identical for varying p as long as it selects x_1, x_2, x_3, x_4 in the first iteration, which is often the case based on Figure 4 at least under the current simulation setting. We acknowledge that the stability of the F_1 scores across different values of p may be attributed, in part, to the relatively small noise standard deviation. In contrast, one-shot descriptor selection methods suffer significantly from just a small increase in the input dimension p since the descriptor space increases double exponentially in p . One way for one-shot description selection methods to circumvent the ultra-high dimension issue is to reduce the maximum composition complexity $M_{\max}$. However, this undesirably rules out descriptors with the correct composition complexity. For example, the descriptor in (11) $\exp(x) =$

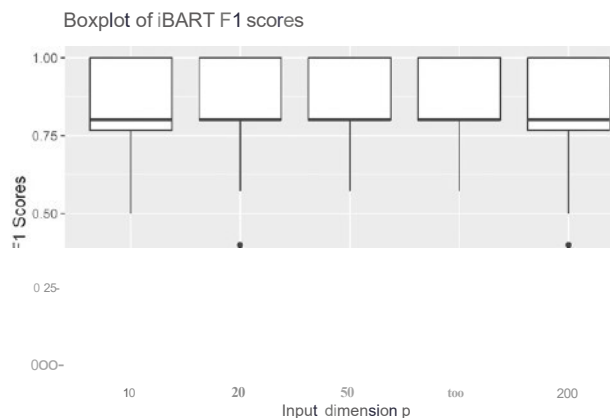


Figure 4. Boxplot of F1 scores for iBART under Model (11) with $p \in \{10, 20, 50, 100, 200\}$.

$\{\exp(x_1) - \exp(x_2)\}^2$ requires at least three compositions of operators and reducing $M_{\max}$ to 2 means that $J_i(x)$ would never be generated and hence would never be selected. Thus, in a complex scenario, one-shot descriptor selection methods cannot generate and select the complex descriptors unless p is small. We considered a small $p = 10$ in the preceding section to accommodate such limitation of one-shot methods.

3.6. Model Misspecification

In this section, we examine the behavior of iBART when the operator set $\mathcal{O}$ is not sufficient to generate the data-generating function. We generate data from the model $y = xJ^{0.7} + e$ with $e \sim \text{Nn}(\mathbf{0}, \mathbf{I})$, $n = 250$, $p = 10$, use the same operator set $\mathcal{O}$ as in Section 3.4, and draw the primary features from a uniform distribution, that is, $x_1, \dots, x_p \sim \text{Un}(0, 3)$. In this simulation, iBART iterates once after screening the primary features. Since iBART only selects x_1 in the first iteration, 100/100 times, it does not apply the binary operators in $\mathcal{O}$.

In contrast to the preceding sections, the computation of true and false positives, as well as the F1 scores, is not feasible here due to the model specification by design. To evaluate the performance of our method, Figure 5(A) shows the frequency of unique descriptors selected by iBART and the average Pearson correlation between the selected descriptor and $x^{0.7}$ in parentheses. Note that multiple descriptors might be selected in one replication, and therefore the sum of frequencies exceeds 100. We can see that, although the true descriptor $xJ^{0.7}$ is not in the candidate descriptor space, iBART selects highly correlated descriptors from the candidate space, including xJ for 98% of the time and x_1 for 87% of the time. Figure 5(B) shows that the RMSE of the iBART selected models does not deviate much from the RMSE of the true model, indicating a similar explanatory power. These observations suggest that when the employed operator space is not sufficient to generate the ground truth, iBART is capable of selecting a model that closely resembles the true model, at least in the considered simulation setting. This is reassuring as in many applications of OIS, such as the one in Section 4, the goal is precisely to find an accurate but interpretable model that well approximates complex real-world physical systems.

4. Application to Single-Atom Catalysis

We apply the proposed method to analyze a single-atom catalysis dataset (O'Connor et al. 2018) in which the goal is to identify physical descriptors that are associated with the binding energy of metal-support pairs calculated by density functional theory (DFT). Single-atom catalysts are popular in modern materials science and chemistry as they offer high reactivity and selectivity while maximizing utilization of the expensive active metal component (Yang et al. 2013; O'Connor et al. 2018; Wang, Li, and Zhang 2018). However, single-atom catalysts suffer from a lack of stability caused by the tendency for single metal atoms to agglomerate in a process called sintering. To prevent sintering in single-atom catalysts, one can tune the binding strength between single metal atoms and oxide supports. While first principle simulations can calculate the binding energy for given metal-support pairs, modeling their association requires explicit statistical modeling and is key to aid the design of single-atom catalysts that are robust against sintering. Feature engineering leads to physical descriptors constructed using mathematical operators and physical properties of the supported metal and the support, and gains popularity in materials informatics as the obtained descriptors are interpretable and provide insights into the underlying physical relationship. The key challenge is to select the most relevant physical properties among large-scale candidate predictors that have explanatory and predictive power to the binding energy; often the sample size is small as first principle simulations are computationally intensive.

The data comprise binding energy of $n = 91$ metal-support pairs and $p = 59$ physical properties of these metal-support pairs, in which the binding energy serves as the response y while the 59 physical properties constitute the primary features $X = (x_1, \dots, x_{59})$. We use the operator set given in (1) but exclude $\sin(n \cdot)$ and $\cos(n \cdot)$ because they are not physically meaningful in this data application.

In this analysis, we compare the best k descriptors ($k = 1, \dots, 5$) of iBART+fo (referred to as iBART for brevity), SISSO, and the method proposed in O'Connor et al. (2018). The method proposed in O'Connor et al. (2018) combines LASSO with extensive domain knowledge; in particular, they rule out a substantial proportion of less promising descriptors using expert knowledge of single-atom catalysis. Such expert-guided, tailored dimension reduction is not needed for SISSO or iBART, but is necessary for LASSO-type methods because the astronomical size of $\mathcal{C}^2(X)$ in this application far exceeds the maximum matrix size allowed in many modern programming languages. We refer to O'Connor et al. (2018) as LASSO" to distinguish it from LASSO+fo implemented in Section 3.4. To enhance interpretability, we eliminate nonphysical descriptors, such as $\text{volume} + \text{speed}$, by automatically comparing the units of the constructed descriptors. Results without this constraint are reported in the supplementary material, showing that this constraint does not alter our conclusions when comparing methods but substantially improves the interpretability of the identified model.

For iBART, we follow the settings described in Section 2.4 and use the descriptor generating process (8). The SISSO parameters are set as followed. The maximum number of descriptors

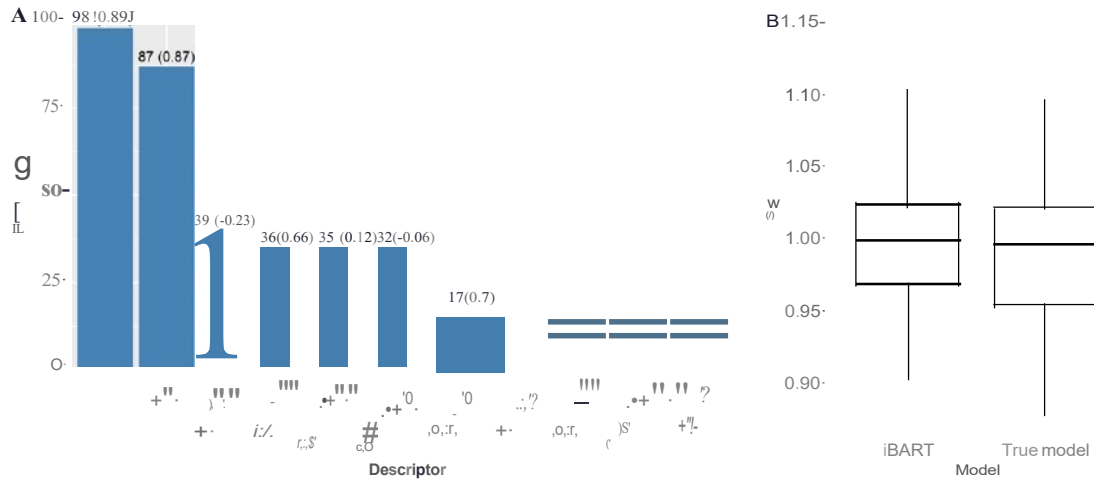


Figure 5. Left: frequency (average Pearson correlation with x_{J-7}) of iBART selected descriptors. Right: root means squares error (RMSE) of iBART models and the true models.

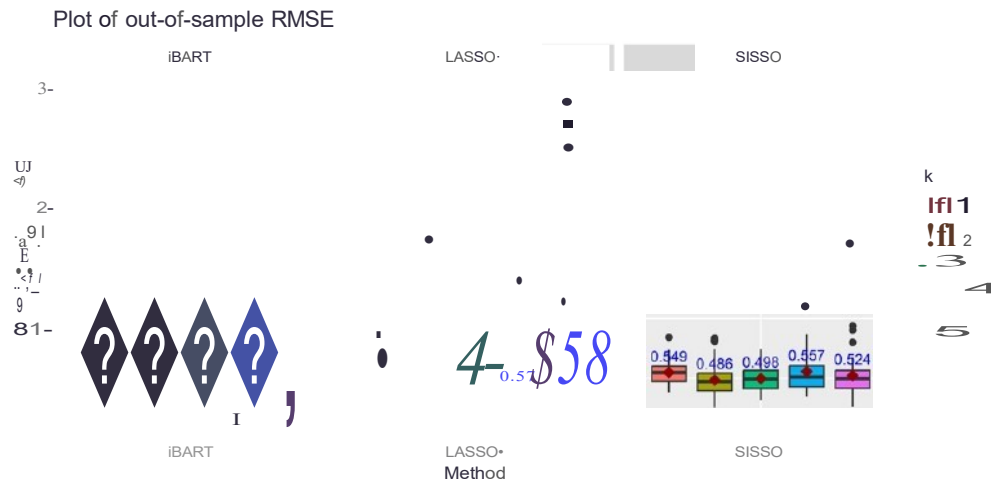


Figure 6. Boxplot of the out-of-sample RMSE for each method across SO random partitions with $k = 1$ to 5 (left to right in each plot). The blue numbers and the red rhombuses indicate the average out-of-sample RMSE.

allowed in a model is set to 5; the descriptor magnitude allowed in the descriptor space is set to $[1 \times 10^{-8}, 1 \times 10^8]$; the size of the SIS-selected subspace is set to 40; the composition complexity M is cap at 2 because $0 < 3$ (X) and higher complexity space exceed the maximum number of elements allowed in a Fortran 90 array. The LASSO* procedure is implemented using the MATLAB code published by O'Connor et al. (2018) with all parameters left as default.

Each method's performance is assessed using out-of-sample RMSE, runtime, and the number of generated descriptors. To calculate out-of-sample RMSE, we randomly partition the $n = 91$ observations into a 90% training set (82 samples) and a 10% testing set (9 samples) and repeat this process SO times. For brevity, we herein refer to out-of-sample RMSE as RMSE.

From Figure 6, the smallest average RMSE is attained at $k = 3$ for iBART (0.41), $k = 1$ for LASSO* (0.471), and $k = 2$ for SISSO (0.486), that is, iBART reduces the RMSE by 13% relative to LASSO* and 16% relative to SISSO; also see the first row of Table 1. iBART outperforms LASSO* and SISSO with smaller average RMSE and reduced variability for $k = 2$. The larger average RMSE of iBART at $k = 1$ may be partially due to its smaller descriptor space as a result of its alternating descriptor

Table 1. Performance comparison of three methods: out-of-sample RMSE, runtime, and the number of generated descriptors, averaged over 50 cross-validations.

	iBART	LASSO*	SISSO
RMSE	0.41	0.471	0.486
Runtime	225 sec	5511 sec $\times$ 20	6943 sec
Number of generated descriptors	627	3.3×10^5	5.5×10^7

generation process, and our implementation details that give advantages to LASSO*. Multiple descriptors are often needed to approximate complex physical systems, and the predictive performance of iBART reassuringly improves with more descriptors in the model. On the contrary, the RMSE of LASSO* increases with more descriptors in the model, indicating overfitting. For all k , we observe iBART does not report large deviations from the average performance, suggesting robustness to training sets compared to LASSO* and SISSO.

In addition to the performance gain in RMSE, iBART also leads to a substantial reduction in computing time and memory usage. Table 1 shows that iBART leads to over 30-fold ($6943/225 = 30.86$) speedup compared to SISSO, and over 480-fold ($5511 \times 20/225 = 489$) speedup compared

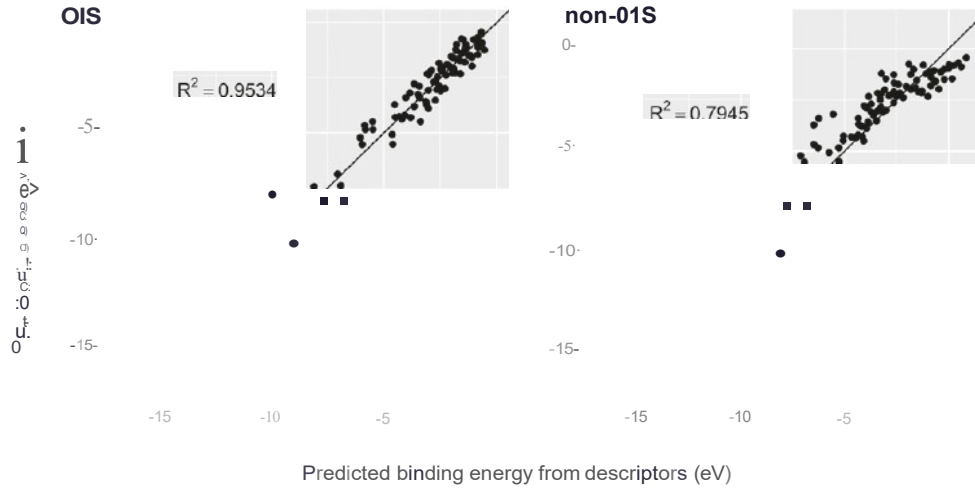


Figure 7. DFT binding energies versus predicted values using linear models with and without OIS. The black line is $y = x$. Each model has $k = 3$ descriptors.

Table 2. Selected linear models by iBART for $k \in \{1, 2, 3, 4, 5\}$.

k	Selected descriptors
1	$\frac{\Delta H_{\text{sub}} - \Delta H_{\text{f,ox,bulk}}}{\Delta E_{\text{vac}}}$
2	$\frac{EA^5 \cdot (\Delta H_{\text{sub}} - \Delta H_{\text{f,ox,bulk}})}{N_{\text{val}}^m / CN_{\text{bulk}}^m}, \left \frac{\Delta H_{\text{sub}} - \Delta H_{\text{f,ox,bulk}}}{\Delta E_{\text{vac}}} \right $
3	$EA^5 \cdot \Delta H_{\text{f,ox,bulk}}, \left \frac{\Delta H_{\text{sub}} - \Delta H_{\text{f,ox,bulk}}}{\Delta E_{\text{vac}}} \right , \left \frac{(\eta^{1/3})^m}{N_{\text{val}}^m \cdot \Delta E_{\text{vac}}} \right $
4	$EA^5 \cdot \Delta H_{\text{f,ox,bulk}}, \frac{(\eta^{1/3})^m}{\phi^5}, \frac{(\eta^{1/3})^m \cdot EA^5}{(N_{\text{val}}^m)^2}, \left \frac{EA^5 \cdot (\Delta H_{\text{sub}} - \Delta H_{\text{f,ox,bulk}})}{\Delta E_{\text{vac}}} \right $
5	$\frac{(\eta^{1/3})^m}{\phi^5}, \frac{\Delta H_{\text{f,ox,bulk}}}{\Delta E_{\text{vac}}}, \frac{EA^5 \cdot (\Delta H_{\text{sub}} - \Delta H_{\text{f,ox,bulk}})}{N_{\text{val}}^m / CN_{\text{bulk}}^m}, \frac{(\eta^{1/3})^m \cdot EA^5}{(N_{\text{val}}^m)^2}, \left \frac{\Delta H_{\text{sub}} - \Delta H_{\text{f,ox,bulk}}}{\Delta E_{\text{vac}}} \right $

to LASSO*, tested on an Intel Xeon Gold 6230 CPU @ 2.10 GHz using either 1 or 20 CPU cores. We use the published code by authors for competing methods, and the runtime reported here does not isolate the effect of various programming languages (R for iBART, MATLAB for LASSO*, and Fortran 90 for SISSO). We did not obtain an accurate single-core runtime of LASSO* because of its poor scalability, and instead multiplied its 20-core runtime by 20 for comparison (i.e., $5511 \times 20 \times 50/3600 = 1530$ hr); the exact speedup of iBART compared to the single-core runtime of LASSO* may vary due to the reduced communication cost among multiple cores and other factors. We remark that the LASSO* method implemented here is given an advantage with an additional dimension reduction step, and an exploration of higher complexity descriptor space as in iBART is computationally prohibitive for LASSO*. The excellent scalability of iBART transforms into memory efficiency as the descriptor space in iBART is orders of magnitude smaller than that of competing methods. Table 1 shows that iBART generates a descriptor space of size $627 < O(p^2)$ in the last iteration on average. Thanks to this significantly smaller descriptor space, we were able to run iBART on a laptop with only 16GB of memory; in contrast, LASSO* and SISSO failed at the descriptor generation step owing to the enormous descriptor space they try to generate, and require server-grade computing facilities.

Table 2 reports the selected descriptors by iBART with various k using the full dataset, and Table 3 reports the physical meanings of the selected primary features. We can see that some descriptors are recurrent in various models, such

Table 3. Descriptions of the selected primary features by iBART.

Primary feature	Physical meaning
ti.H,ub	Heat of sublimation
ti.Htox,bulk	Oxidation energy of the bulk metal
ti.fv,c*	Oxygen vacancy energy
EA'	Electron affinity of support
all..	Number of valence electrons in metal adatom
(N'b lk	Coordination number of the surface metal atom in the bulk phase
	Discontinuity in electron density of metal adatom
	Chemical potential of the electrons in support
	Fourth ionization energy of support with the bulk metal in the 4+ oxidation state

as $\frac{ti.H,ub - ti.H,ox,bulk}{f.Evac}$. This reassuringly suggests the stability of the proposed descriptor selection method across k . Readers are referred to Liu et al. (2022) for in depth analysis and physical interpretation of these selected descriptors from the perspective of catalyst design.

Figure 7 demonstrates a clear advantage of the OIS model over the non-OIS model. The OIS model is the iBART model with $k = 3$, the optimum model suggested by RMSE. The non-OIS model is a simple least squares model with X as the design matrix and no feature engineering step. For ease of comparison, the non-OIS model also has $k = 3$ predictors determined by best subset selection. The OIS model yields an R^2 of 0.9534, indicating high explanatory power. In contrast, the non-OIS model gives an $R^2 = 0.7945$ and shows poor fitting performance as the scatterplot deviates from the diagonal line $y = x$ considerably. Both the OIS and the non-OIS models have analytical forms, but the OIS model gains more explanatory power and gives an insightful description of the response through nonlinearity of its predictors. In particular, the non-OIS model is $\hat{Y}_{\text{non-OIS}} = -4.7 - 0.4 \times ti.Hf,ox + 0.3 \times N!a1 + 1.0 \times ti.Evac$, while the OIS model is

$$\hat{Y}_{\text{OIS}} = -0.01 + 0.4 \times (EA^5 \cdot ti.Hf,ox,bulk) \cdot \frac{f.Hsub - f.Hf,ox,bulk}{f.Evac} - 0.6 \times \frac{1.Evac}{(771/3r)} - 19.6 \times \frac{1}{f.Evac}$$

This OIS model pinpoints targeted descriptors to guide further investigation into the underlying physical discovery.

5. Discussion

In this article, we study variable selection in the presence of feature engineering that is widely applicable in many scientific fields to provide interpretable models. Unlike in classical variable selection, candidate predictors are engineered from primary features and composite operators. While this problem has become increasingly important in science, such as the emerging field of materials informatics, the induced new geometry has not been studied in the statistical literature. We propose a new strategy "parametrics assisted by nonparametrics": or PAN, to efficiently explore the descriptor space and achieve nonparametric dimension reduction for linear models. Using BART-G.SE as the nonparametric module, the proposed method iBART iteratively constructs and selects complex descriptors. Compared to one-shot descriptor construction approaches, iBART does not operate on an ultra-high dimensional descriptor space and thus substantially mitigates the "curse of dimensionality" and high correlations among descriptors. We introduce the OIS framework, define a PAN criterion, and assess iBART through the lens of this criterion. This sets a foundation for future research in interpretable model selection with feature engineering.

Other than the methodological contributions, we use extensive experiments to demonstrate appealing empirical features of iBART that may be crucial for practitioners. iBART automates the feature generating step; one-shot methods such as SISSO often require user intervention as otherwise they are not scalable to handle the overwhelmingly large descriptor space—this descriptor space increases double exponentially in the number of primary features and binary operators as a result of compositions. iBART has excellent scalability and leads to robust performance when the dimension of primary features increases to a level that renders existing methods computationally infeasible. Compared to SISSO, which is widely perceived as state-of-the-art in the field of materials genomes, our data application shows that iBART reduces the out-of-sample RMSE by 16% with an over 30-fold speedup and a fraction of memory demand. Overall, the proposed method accomplishes traditionally "server-required" tasks using a regular laptop or desktop with improved accuracy. Beyond materials genomes illustrated in Section 4, iBART can be applied to a vast domain where interpretable modeling is of interest.

There are interesting next directions building on our iBART approach. First, OIS provides a useful perspective for structured variable selection by introducing operators and compositions. Certain operator sets may be particularly useful depending on the application; for example, the composite multiply operator leads to high-order interactions. It is interesting to examine the performance of iBART with a diverse choice of operators. Also, we have focused on finite sample performance when assessing selection accuracy through simulations partly in view of the limited sample size typically available in the motivating example. It is nevertheless interesting to theoretically investigate the necessary conditions for the PAN criterion to hold for selected nonparametric variable selection methods such as BART when the sample size diverges.

Supplementary Materials

Code and data for replicating the study results are available at <https://github.com/xylimeng/OIS>. An R package that implements iBART is available at <https://github.com/mattshengli/iBART>. The supplementary materials include additional simulations and proof. Supplementary A contains a detailed comparison of variants of iBART by varying the nonparametric module. Supplementary B contains additional simulation results using correlated primary features. Supplementary C presents additional real data application results without enforcing unit consistency. Supplementary D contains the proof of Theorem 2.1.

Acknowledgments

We thank the editor, associate editor, and reviewers for constructive comments that helped to improve the article.

Disclosure Statement

The authors report there are no competing interests to declare.

Funding

This research was partially supported by the Big-Data Private-Cloud Research Cyberinfrastructure MRI-award funded by NSF under grant CNS-1338099 and by Rice University's Center for Research Computing. Shengbin Ye's research was supported by NIH grant T32CA096520. Meng Li's research was partially supported by NSF grants DMS-2015569 and DMS/NIGMS-2153704. Thomas P. Senftle's research was supported by NSF grant CBET-2143941.

ORCID

Shengbin Ye  <http://orcid.org/0000-0001-8767-2595>
Thomas P. Senftle  <http://orcid.org/0000-0002-5889-5009>
Meng Li  <http://orcid.org/0000-0003-2123-2444>

References

- Bleich, J., Kapelner, A., George, E. I., and Jensen, S. T. (2014), "Variable Selection for BART: An Application to Gene Regulation," *Annals of Applied Statistics*, 8, 1750-1781. [82,85,86]
- Chipman, H. A., George, E. I., and McCulloch, R. E. (2010), "BART: Bayesian Additive Regression Trees," *Annals of Applied Statistics*, 4, 266-298. [85]
- Fan, J., and Lv, J. (2008), "Sure Independence Screening for Ultrahigh Dimensional Feature Space," *Journal of the Royal Statistical Society, Series B*, 70, 849-911. [82]
- Feurer, M., Klein, A., Eggenberger, K., Springenberg, J., Blum, M., and Hutter, F. (2015), "Efficient and Robust Automated Machine Learning," in *Advances in Neural Information Processing Systems* (Vol. 28), Curran Associates, Inc. [83]
- Foppa, L., Ghiringhelli, L. M., Girgsdies, F., Hashagen, M., Kube, P., Havecker, M., Carey, S. J., Tarasov, A., Kraus, P., Rosowsky, F., Schlogl, R., Trunschke, A., and Scheffler, M. (2021), "Materials Genes of Heterogeneous Catalysis from Clean Experiments and Artificial Intelligence," *MRS Bulletin*, 46, 1016-1026. [81]
- Ghiringhelli, L. M., Vybiral, J., Levchenko, S. V., Draxl, C., and Scheffler, M. (2015), "Big Data of Materials Science: Critical Role of the Descriptor," *Physical Review Letters*, 114, 105503. [81]
- Hart, G. L., Mueller, T., Toher, C., and Curtarolo, S. (2021), "Machine Learning for Alloys," *Nature Reviews Materials*, 6, 730-755. [81]
- Horiguchi, A., Pratola, M. T., and Santner, T. J. (2021), "Assessing Variable Activity for Bayesian Regression Trees," *Reliability Engineering & System Safety*, 207, 107391. [86]

- Keith, J. A., Vassilev-Galindo, V., Cheng, B., Chnriela, S., Gastegger, M., Miiller, K.-R., and Tkatchenko, A. (2021), "Combining Machine Learning and Computational Chemistry for Predictive Insights into Chemical Systems," *Chemical Reviews*, 121, 9816-9872. [81]
- Khurana, U., Samulowitz, H., and Turaga, D. (2018), "Feature Engineering for Predictive Modeling Using Reinforcement Learning," in *The Thirty-Second AAAI Conference on Artificial Intelligence (AAAI-18)*. [83]
- Lin, Y., Saboo, A., Frey, R., Sorkin, S., Gong, J., Olson, G. B., Li, M., and Niu, C. (2021), "CALPHAD Uncertainty Quantification and TDBX," *The Journal of The Minerals, Metals & Materials Society*, 73, 116-125. [81]
- Linero, A. R. (2018), "Bayesian Regression Trees for High-Dimensional Prediction and Variable Selection," *Journal of the American Statistical Association*, 113, 626-636. [86]
- Liu, C.-Y., Ye, S., Li, M., and Senftle, T. P. (2022), "A Rapid Feature Selection Method for Catalyst Design: Iterative Bayesian Additive Regression Trees (iBART)," *The Journal of Chemical Physics*, 156, 164105. (82,92]
- Liu, C.-Y., Zhang, S., Martinez, D., Li, M., and Senftle, T. P. (2020), "Using Statistical Learning to Predict Interactions between Single Metal Atoms and Modified MgO(100) Supports," *npj Computational Materials*, 6, 102. (82]
- Liu, X., Wu, Y., Irving, D. L., and Li, M. (2021), "Gaussian Graphical Models with Graph Constraints for Magnetic Moment Interaction in High Entropy Alloys," ArXiv e-prints. [81]
- Liu, Y., Rockova, V., and Wang, Y. (2021), "Variable Selection with ABC Bayesian Forests," *Journal of the Royal Statistical Society, Series B*, 83, 453-481. [86]
- Markovitch, S., and Rosenstein, D. (2002), "Feature Generation Using General Constructor Functions," *Machine Learning*, 49, 59-98. (83]
- O'Connor, N. J., Jonayat, A. S. M., Janik, M. J., and Senftle, T. P. (2018), "Interaction Trends between Single Metal Atoms and Oxide Supports Identified with Density Functional Theory and Statistical Learning," *Nature Catalysis*, 1, 531-539. (82,90,91]
- Ouyang, R., Curtarolo, S., Ahrnetcik, E., Scheffler, M., and Ghiringhelli, L. M. (2018), "SISSO: A Compressed-Sensing Method for Identifying the Best Low-Dimensional Descriptor in an Immensity of Offered Candidates," *Physical Review Materials*, 2, 083802. (82,83,88]
- Tibshirani, R. (1996), "Regression Shrinkage and Selection via the Lasso," *Journal of the Royal Statistical Society, Series B*, 58, 267-288. (82,85]
- Wang, A., Li, J., and Zhang, T. (2018), "Heterogeneous Single-Atom Catalysis," *Nature Reviews Chemistry*, 2, 65-81. (90]
- Yang, X.-F., Wang, A., Qiao, B., Li, J., Liu, J., and Zhang, T. (2013), "Single-Atom Catalysts: A New Frontier in Heterogeneous Catalysis," *Accounts of Chemical Research*, 46, 1740-1748. (90]
- Zhong, M., Tran, K., Min, Y., Wang, C., Wang, Z., Dinh, C.-T., De Luna, P., Yu, Z., Rasouli, A. S., Brodersen, P., Sun, S., Voznyy, O., Tan, C.-S., Askerka, M., Che, F., Liu, M., Seifitokaldani, A., Pang, Y., Lo, S.-C., Ip, A., Ulissi, Z., and Sargent, E. H. (2020), "Accelerated Discovery of CO₂ Electrocatalysts Using Active Machine Learning," *Nature*, 581, 178-183. (81]