

# EFFICIENT REGULARIZED ISOTONIC REGRESSION WITH APPLICATION TO GENE-GENE INTERACTION SEARCH

BY RONNY LUSS  
SAHARON ROSSET AND MONI SHAHAR

*Tel Aviv University*

Isotonic regression is a nonparametric approach for fitting monotonic models to data that has been widely studied from both theoretical and practical perspectives. However, this approach encounters computational and statistical overfitting issues in higher dimensions. To address both concerns we present an algorithm, which we term Isotonic Recursive Partitioning (IRP), for isotonic regression based on recursively partitioning the covariate space through solution of progressively smaller “best cut” subproblems. This creates a regularized sequence of isotonic models of increasing model complexity that converges to the global isotonic regression solution. The models along the sequence are often more accurate than the unregularized isotonic regression model because of the complexity control they offer. We quantify this complexity control through estimation of degrees of freedom along the path. Success of the regularized models in prediction and IRP’s favorable computational properties are demonstrated through a series of simulated and real data experiments. We discuss application of IRP to the problem of searching for gene-gene interactions and epistasis, and demonstrate it on data from genome-wide association studies of three common diseases.

**1. Introduction.** In predictive modeling, we are given a set of  $n$  data observations  $(x_1, y_1), \dots, (x_n, y_n)$ , where  $x \in \mathcal{X}$  (usually  $\mathcal{X} = \mathbb{R}^d$ ) is a vector of covariates or independent variables,  $y \in \mathbb{R}$  is the response, and we wish to fit a model  $\hat{f} : \mathcal{X} \rightarrow \mathbb{R}$  to describe the dependence of  $y$  on  $x$ . The most common approach traditionally used in statistics to accomplish this task is to seek a model that minimizes in-sample squared error ( $l_2$ ) loss, i.e.

$$\hat{f} = \arg \min_{f \in \mathcal{F}} \sum_{i=1}^n (y_i - f(x_i))^2,$$

where  $\mathcal{F}$  is a class of candidate models. The use of squared error loss can be interpreted as maximum likelihood fitting under an additive gaussian noise assumption, or more generally as trying to estimate  $\mathbb{E}(y|x)$  (Gneiting, 2011). Other loss functions can be used to represent other estimation goals.

Isotonic regression is a non-parametric modeling approach which restricts the fitted model to being monotone in all independent variables (Barlow and Brunk, 1972). Define  $\mathcal{G}$  to be the family of isotonic functions, that is,  $g \in \mathcal{G}$  satisfies

$$x_1 \preceq x_2 \Rightarrow g(x_1) \leq g(x_2),$$

where the partial order  $\preceq$  here will usually be the standard Euclidean one, i.e.,  $x_1 \preceq x_2$  if and only if  $x_{1j} \leq x_{2j}$  coordinate-wise. Given these definitions, standard  $l_2$  isotonic regression solves

$$(1) \quad \min_{g \in \mathcal{G}} \sum_{i=1}^n (y_i - g(x_i))^2.$$

---

AMS 2000 subject classifications: Primary 62G08

Keywords and phrases: multivariate isotonic regression, nonparametric regression, regularization path, partitioning

We denote by  $\hat{f}$  the optimal solution to (1). As many authors have noted,  $\hat{f}$  comprises a partitioning of the space  $\mathcal{X}$  into regions with no “holes” that satisfy isotonicity properties defined below, with a constant fitted to  $\hat{f}$  in every region.

In terms of model form, isotonic regression is clearly very attractive in situations where monotonicity is a reasonable assumption, but other common assumptions like linearity or additivity are not. Indeed, this formulation has found useful applications in biology (Obozinski *et al.*, 2008), medicine (Schell and Singh, 1997), statistics (Barlow and Brunk, 1972) and psychology (Kruskal, 1964), among others. In recent years, an exciting new application area has emerged for this approach in genetics: modeling genetic interactions in heritability. Many papers have noted the apparent insufficiency of standard additive modeling approaches in describing the combined effects of genetic factors (e.g., mutations) on phenotypes like height and disease (Goldstein, 2009; Eichler *et al.*, 2010). Some findings in mice have pointed to sub-additive interactions (Shao *et al.*, 2008), while others suggest requiring super-additive assumptions in order to explain heritability (Goldstein, 2009). It is generally accepted, however, that while the effect of one genetic factor on a phenotype can be modulated, enhanced or even eliminated by other genetic factors, it is not expected to reverse direction (Mani *et al.*, 2007; Roth, Lipshitz and Andrews, 2009). In other words, the isotonicity assumption with respect to genetic effects is widely accepted, but the form of epistasis (genetic interaction) between factors is not clear and may vary between phenotypes. Other properties of this application domain also favor the use of isotonic regression as we discuss below. It should be noted, however, that most discussions of epistasis have been theoretical, with extensive systematic efforts at discovering actual epistatic combinations in human disease resulting in very limited (Emily *et al.*, 2009) or no results (Cordell, 2009). These efforts have utilized simple statistical approaches (chi-square tests, logistic regression) to search for low-dimensional interactions, mostly of two mutations at a time. The question of whether their lack of findings is due to limitations of methodology and concentration on low dimension, or lack of real epistatic signal, is a key one, and it can be addressed by isotonic modeling.

Two major concerns arise when considering the practical use of isotonic regression in *modern* situations as the number of observations  $n$ , the data dimensionality  $d$ , and the number of isotonicity constraints  $m = |\{(i, j) : x_i \preceq x_j\}|$  implied by (1) all grow large: statistical overfitting and computational difficulty. The notations  $n$ ,  $d$ , and  $m$  will refer to these quantities throughout the paper.

The first concern is statistical difficulty and overfitting. Beyond very low dimensions, the isotonicity constraints on the family  $\mathcal{G}$  can become inefficient in controlling model complexity and the isotonic regression solutions can be severely overfitted (for example, see Bacchetti (1989) and Schell and Singh (1997)). At the extreme, there may be no isotonicity constraints because no two observations obey the coordinate-wise requirement for the  $\preceq$  ordering. The isotonic solution in this case simply assigns  $\hat{f}(x_i) = y_i$  providing a perfect interpolation of the training data. As demonstrated in the literature (Schell and Singh, 1997) and below, the overfitting concern is clearly well-founded when considering the optimal isotonic regression model implied by (1), even in non-extreme cases with a large number of constraints. In this case, regularization, i.e. fitting isotonic models that are constrained to a restricted subset of  $\mathcal{G}$ , could offer an approach that maintains isotonicity while controlling variance, leading to improved accuracy.

A second concern is computational difficulty. The discussion of isotonic regression originally focused on the case  $x \in \mathbb{R}$ , where  $\preceq$  denoted a complete order (Kruskal, 1964). For this case, the well-known pooled adjacent violators algorithm (PAVA) efficiently solves (1) in computational complexity  $O(n)$ . Low complexities can also be found when the isotonic constraints take a special structure such as a tree ( $O(n \log n)$  in Pardalos and Xue (1999)). Various algorithms have been developed for the partially ordered case, including the classical approach of Dykstra and Robertson (1982) for data on a grid, generalizations of PAVA (Lee, 1983; Block, Qian and Sampson, 1994) and active set methods (de Leeuw, Hornik and Mair, 2009). These approaches offer no polynomial complexity guarantees and are impractical when data sizes exceed a

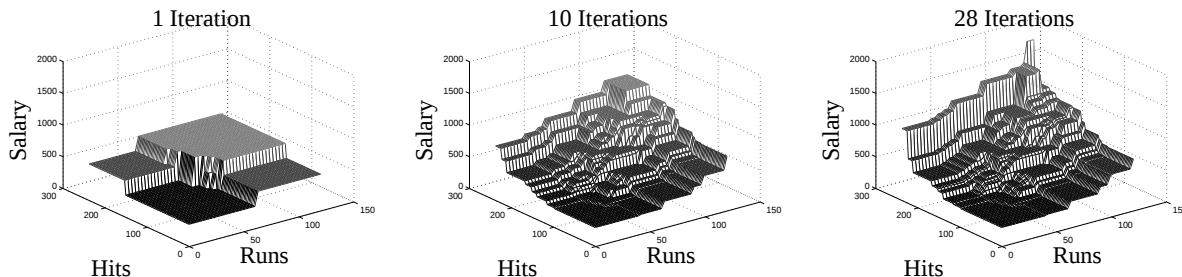
few thousand observations (in some cases much less). Interior point methods offer complexity guarantees of  $O(\max(m, n)^3)$  (Monteiro and Adler, 1989), however they are impractical for large data sizes due to excessive memory requirements. A much more computationally attractive approach can be found in the optimization and operations research literature. The basic idea of this approach is to repeatedly and “optimally” split the covariate space  $\mathcal{X}$  into regions of decreasing size by solving a sequence of specially structured *best cut* problems for which efficient algorithms exist. At most  $n$  partitions are needed, leading to a computational complexity bounded by  $O(n^4)$ , and in some cases even less. From a practical performance perspective, this algorithm can obtain an exact solution of (1) for datasets with tens of thousands of observations in minutes. The first appearance of this approach, to our knowledge, is in the work of Maxwell and Muckstadt (1985) (and similarly Roundy (1986)), who presented a model for finding reorder intervals in a production-distribution system. At the center of their problem was a regression subject to isotonicity constraints, but with a different loss function than that in problem (1), to which they provided an algorithm based on partitioning. Applicability of their algorithm with minimal changes to problem (1) was more recently noticed by several authors (e.g., Spouge, Wan and Wilbur (2003)), who used it to state a similar efficient partitioning scheme for isotonic regression. A similar, highly efficient, algorithm by Hochbaum and Queyranne (2003) also solves problem (1) under the additional constraint that fits take integer values. This partitioning approach does not appear to be well-known in the statistics community, and indeed we have independently developed it before discovering it is already known.

The literature cited above invariably refers to this iterative splitting algorithm merely as an approach for efficiently arriving at the optimal solution of (1). However, as noted before, this solution can be highly overfitted, especially as the dimension  $d$  increases. Our main interest lies in analyzing the iterative approach as a means towards resolving the overfitting problem, as well as the computational issue. We propose to view this iterative algorithm as a *recursive partitioning* approach that generates isotonic models of increasing model complexity, ultimately leading to the solution of (1); the algorithm is termed Isotonic Recursive Partitioning (IRP). We prove that the models generated by the IRP iterations are indeed isotonic (Theorem 4) and consider them as a *regularization path* of increasingly complex isotonic regression models. Models along the path are less complex, and hence likely to be less overfit and offer better predictive performance than the overall solution to (1), while still maintaining isotonicity. This is confirmed by our analysis of the equivalent degrees of freedom along the IRP path, as well as experiments with simulated and real data.

We observe that for very low dimension (typically  $d \leq 2$ ) the non-regularized solution of (1) performs well. As the dimension increases, regularization becomes necessary, and intermediate models on the IRP path perform better than the non-regularized solution. However, eventually overfitting plagues IRP from its first iteration, and the isotonic models fail to perform better than simple linear regression in out-of-sample prediction, even when the linear model is inappropriate. In our simulations, this occurs around dimensions 6-8 even for relatively large data sets.

Progress of IRP is illustrated in Figure 1, where we show an example of applying IRP to the well-known Baseball dataset (He, Ng and Portnoy, 1998) describing the dependence of salary on a collection of player properties. We limit the model to only two covariates to facilitate visualization, and we choose to use the number of runs batted in and hits since they seemed a-priori most likely to comply with the isotonicity assumptions. The increasing model complexity can be seen, moving from iteration 1 (a single split) through 10 iterations of IRP, to the final isotonic model optimally solving (1), comprising a splitting of the covariate space into 29 regions, each of which is fitted with a constant. Note that the single split from iteration 1 creates two flat surfaces (the highest and lowest in the figure), and that the *in-between* surfaces are interpolations in regions with no data points based on the two-surface model. The figures for models after 10 and 28 iterations are quite complex for the same reason - interpolations in the continuous covariate space.

An obvious analogy of IRP can be made to well-known recursive partitioning approaches for regression



**Fig 1:** Illustration of IRP on Baseball data. Salary is modeled by number of runs batted in and hits. Models after iterations 1 and 10 of IRP and the final model are shown.

such as CART (Breiman *et al.*, 1984), where the iterative splitting of the covariate space generates a sequence of models (trees) of increasing model complexity, from which the “best” tree is chosen via cross validation (for example, using the 1-SE rule (Breiman *et al.*, 1984)). As with CART and other similar approaches, IRP performs a greedy search and finds a “local” optimum in every iteration. However, unlike CART, which has no guarantees on the overall model it generates, IRP is proven to terminate in the global solution of the isotonic regression problem (1). Another difference is that IRP splits are not made along one axis at a time, but rather each split is a non-parametric division of one region in  $\mathcal{X}$  into two sub-regions.

Importantly, while our presentation has so far concentrated on isotonic regression with an  $l_2$  loss function, an interesting extension of the partitioning scheme can be used to solve large-scale isotonic regressions under other loss functions that are of great interest in statistics, e.g., logistic and poisson log-likelihoods. Specifically, a well-known result of Barlow and Brunk (1972, Theorem 3.1) implies that, for a large class of loss functions (formally defined in our Discussion section), the solution of isotonic modeling can be found by a simple transformation of the solution to isotonic regression with an  $l_2$  loss function. This will play an important role in our use of isotonic regression on modeling gene-gene interactions in epistasis. Indeed, there we will maximize logistic log-likelihood subject to isotonicity constraints.

The remainder of this paper is organized as follows. We first present and analyze the IRP algorithm in Section 2. We detail the best cut problem solved for splitting at each iteration, and prove that this algorithm is a *no-regret* algorithm, in the sense that it only partitions the data and never merges back previously made partitions and converges to the global solution of (1) (Theorem 2). Furthermore, we prove that the intermediate partitions generated along the IRP path are also isotonic, in the sense that fitting each region to the average gives a model that is in the class  $\mathcal{G}$  of isotonic functions in  $\mathbb{R}^d$  (Theorem 4). Section 3 briefly reviews the theoretical computational guarantees of IRP as reflected in the literature, and develops a simple and realistic case where the overall computation is  $O(n^3)$ . Section 4 discusses the statistical model complexity of models generated along the regularization path. Meyer and Woodroffe (2000) have shown that the number of partitions in the solution of (1) is an unbiased estimator of the (equivalent) degrees of freedom (as defined by Efron (1986)). Since IRP adds one to the number of partitions at each iteration, the number of iterations may be used as a parametrization of this sequence. However, we argue that the number of regions is not a good estimate of degrees of freedom because IRP performs much more fitting in its initial iterations compared to later stages, and demonstrate this effect empirically through simulation. We also show that when the covariates are ternary, i.e. covariate  $x_{ij} \in \{0, 1, 2\}$  (as is natural in our motivating genetic example when dealing with ternary genotype data), the overall number of degrees of freedom and model complexity increase more slowly with dimension, compared to general continuous covariates, resulting in much less overall fitting for each dimension. Section 5 examines IRP’s statistical and computational performance on simulated and real data, specifically pointing out the effect of regularization and

increased dimensionality on predictive performance. We apply IRP to simulations with ternary covariates and sub- and super-additive interactions motivated by the genetic application and demonstrate its favorable performance. Section 6 applies IRP to real-life datasets from the Wellcome Trust Case Control Consortium (WTCCC) and discusses the results (WTCCC, 2007). Section 7 concludes with extensions and connections to previous literature. A Matlab-based software package implementing the IRP algorithm is available at [www.tau.ac.il/~saharon/files/IRPv1.zip](http://www.tau.ac.il/~saharon/files/IRPv1.zip).

We next define terminology to be used throughout the paper.

**1.1. Definitions.** Let  $V = \{x_1, \dots, x_n\}$  be the covariate vectors for  $n$  training points where  $x_i \in \mathbb{R}^d$  and denote  $y_i \in \mathbb{R}$  as the  $i^{\text{th}}$  observed response. We will refer to a general subset of points  $A \subseteq V$  with no holes (i.e.  $x \preceq y \preceq z$  and  $x, z \in A \Rightarrow y \in A$ ) as a *group*. Denote by  $|A|$  the cardinality of group  $A$ . The *weight* of group  $A$  is defined as  $\bar{y}_A = \frac{1}{|A|} \sum_{i: x_i \in A} y_i$ . For two groups  $A$  and  $B$ , we denote  $A \preceq B$  if there exists  $x \in A, y \in B$  such that  $x \preceq y$  and there does not exist  $x \in A, y \in B$  such that  $y \prec x$  (i.e. there is at least one comparable pair of points that satisfy the direction of isotonicity). A set of groups  $\mathcal{V}$  is called isotonic if  $A \preceq B \Rightarrow \bar{y}_A \leq \bar{y}_B$  for all  $A, B \in \mathcal{V}$ . The groups within this set  $\mathcal{V}$  are referred to as isotonic regions. A subset  $\mathcal{L}$  ( $\mathcal{U}$ ) of  $A$  is a *lower set* (*upper set*) of  $A$  if  $x \in A, y \in \mathcal{L}, x \prec y \Rightarrow x \in \mathcal{L}$  ( $x \in \mathcal{U}, y \in A, x \prec y \Rightarrow y \in \mathcal{U}$ ).

A group  $B \subseteq A$  is defined as a block of group  $A$  if  $\bar{y}_{\mathcal{U} \cap B} \leq \bar{y}_B$  for each upper set  $\mathcal{U}$  of  $A$  such that  $\mathcal{U} \cap B \neq \{\}$  (or equivalently if  $\bar{y}_{\mathcal{L} \cap B} \geq \bar{y}_B$  for each lower set  $\mathcal{L}$  of  $A$  such that  $\mathcal{L} \cap B \neq \{\}$ ). A set of blocks  $\mathcal{S} = \{B_1, \dots, B_k\}$  is called *block class* of  $V$  if  $B_i \cap B_j = \{\}$  and  $B_1 \cup \dots \cup B_k = V$ .  $\mathcal{S}$  is an *isotonic block class* if for all  $B_i, B_j \in \mathcal{S}$ ,  $B_i \preceq B_j \Rightarrow \bar{y}_{B_i} \leq \bar{y}_{B_j}$ . A group  $X$  *majorizes* (*minorizes*) another group  $Y$  if  $X \succeq Y$  ( $X \preceq Y$ ). A group  $X$  is a *majorant* (*minorant*) of  $X \cup A$  where  $A = \cup_{i=1}^k A_i$  if  $X \not\prec A_i$  ( $X \not\preceq A_i$ ) for all  $i = 1 \dots k$ .

We denote the optimal solution for maximizing a general function  $f(z)$  in the variable  $z$  by  $z^*$ , i.e.  $z^* = \operatorname{argmax} f(z)$ .

**2. IRP and a regularization path for isotonic regression.** We describe here the partitioning algorithm used to solve the isotonic regression problem (1). Section 2.1 first reformulates the isotonic regression problem and describes the structure of the optimal solution. Section 2.2 motivates and details the IRP algorithm and, in particular, the main partitioning step. Each group created by the partitioning scheme is proven to be the union of blocks in the optimal solution, i.e. all partitions have the no-regret property. An important aspect of the algorithm is the regularization path generated as a byproduct as each partition creates a new feasible solution. Section 2.3 goes on to prove convergence of IRP to the global optimal solution of (1), and most importantly, that each solution along the regularization path is isotonic.

**2.1. Structure of the isotonic solution.** Isotonic regression seeks a monotonic function that fits a given training dataset  $\{(x_i, y_i)\}_{i=1}^n$  and satisfies a set of *isotonicity constraints* which we index by the set  $\mathcal{I} = \{(i, j) : x_i \preceq x_j\}$ . We will usually assume that  $x_i \in \mathbb{R}^d$  and that  $\preceq$  is the standard partial order in  $\mathbb{R}^d$  based on coordinate-wise inequalities. A reformulation of (1) is

$$(2) \quad \min \left\{ \sum_{i=1}^n (\hat{y}_i - y_i)^2 : \hat{y}_i \leq \hat{y}_j \quad \text{for all } (i, j) \in \mathcal{I} \right\}.$$

Problem (2) is a quadratic program with linear constraints. Any solution satisfying the constraints given by  $\mathcal{I}$  is referred to as an isotonic, or feasible, solution. The structure of the optimal solution to (2) is well-known. Observations are divided into  $k$  groups where the fits in each group take the group mean observation value. This can be seen through the following Karush-Kuhn-Tucker (KKT), i.e. optimality, conditions (Boyd and Vandenberghe, 2004) to (2):

- (a)  $\hat{y}_i = y_i - \frac{1}{2} \left( \sum_{j:(i,j) \in \mathcal{I}} \lambda_{ij} - \sum_{j:(j,i) \in \mathcal{I}} \lambda_{ji} \right)$ ,
- (b)  $\hat{y}_i \leq \hat{y}_j$  for all  $(i, j) \in \mathcal{I}$ ,
- (c)  $\lambda_{ij} \geq 0$  for all  $(i, j) \in \mathcal{I}$ ,
- (d)  $\lambda_{ij}(\hat{y}_i - \hat{y}_j) = 0$  for all  $(i, j) \in \mathcal{I}$ ,

where  $\lambda_{ij}$  is the dual variable corresponding to the isotonicity constraint  $\hat{y}_i \leq \hat{y}_j$ . This set of conditions exposes the nature of the optimal solution. Condition (d) implies that  $\lambda_{ij} > 0 \Rightarrow \hat{y}_i = \hat{y}_j$  meaning  $\lambda_{ij}$  can be non-zero only within blocks in the isotonic solution which have the same fitted value. For observations in different blocks,  $\lambda_{ij} = 0$ . Furthermore, the fit within each block is trivially seen to be the average of the observations in the block, as the average minimizes the block's squared loss. A block is thus also referred to as an *optimal group* with respect to an isotonic regression problem. Condition (b) implies isotonicity of the blocks, and thus, we get the familiar characterization of the isotonic regression problem as one of finding a division into an isotonic block class.

**2.2. The partitioning algorithm.** Suppose a current group  $V$  is optimal (i.e.  $V$  is a block) and thus the optimal fits at points in  $V$ , denoted  $\hat{y}_i^*$ , satisfy  $\hat{y}_i^* = \bar{y}_V$  for all  $i \in V$ , which leads to the condition  $\sum_{i \in V} (y_i - \bar{y}_V) = 0$ . Then finding two groups  $A$  and  $B$  within  $V$  such that  $\sum_{i \in B} (y_i - \bar{y}_V) - \sum_{i \in A} (y_i - \bar{y}_A) > 0$  should be infeasible, according to the KKT conditions. The division in IRP looks for two such groups. Denote by  $\mathcal{C}_V = \{(A, B) : A, B \subseteq V, A \cup B = V, A \cap B = \{\}\}$ , and there does not exist  $x \in A, y \in B$  such that  $y \preceq x$  the set of all feasible (i.e. isotonic) partitions defined by observations in  $V$ . We refer to partitioning as making a cut through the variable space (hence our optimal partition is made by an *optimal cut*). The optimal cut is determined by the partition that solves the problem

$$(3) \quad \max_{(A,B) \in \mathcal{C}_V} g^*(A, B) = \max_{(A,B) \in \mathcal{C}_V} \left\{ \sum_{i \in B} (y_i - \bar{y}_V) - \sum_{i \in A} (y_i - \bar{y}_V) \right\} = -|A|(\bar{y}_A - \bar{y}_V) + |B|(\bar{y}_B - \bar{y}_V)$$

where  $A(B)$  is the group on the lower (upper) side of the edges of cut. A more statistically intuitive rule might look for the split that maximizes between-group variance. This partitioning problem solves

$$(4) \quad \max_{(A,B) \in \mathcal{C}_V} \tilde{g}(A, B) = \max_{\{(A,B) \in \mathcal{C}_V\}} \left\{ |A|(\bar{y}_A - \bar{y}_V)^2 + |B|(\bar{y}_B - \bar{y}_V)^2 \right\}.$$

The next proposition makes a connection between the above two maximization problems, and draws a clear conclusion on the relationship between their optimal solutions, namely that the optimal partitions to (3) are always more balanced than the optimal partitions to (4).

**Proposition 1** Denote the optimal solutions of the optimal cut problem (3) and the between-group variance maximization problem (4) by  $(A^*, B^*)$  and  $(\tilde{A}, \tilde{B})$  and their objective functions by  $g^*(A, B)$  and  $\tilde{g}(A, B)$ , respectively. Then

$$(A^*, B^*) = \operatorname{argmax}_{(A,B) \in \mathcal{C}_V} \{|A||B|\tilde{g}(A, B)\}$$

and

$$(|A^*| - |B^*|)^2 \leq (|\tilde{A}| - |\tilde{B}|)^2.$$

We leave the proof to the appendix.

Thus, we can look at the IRP criterion as a modified form of maximizing between-group variance which encourages more balanced splitting. However, while solving the partition problem (4) is difficult, the IRP

partition problem (3) is tractable. Indeed, the optimal partition problem (3) can be reduced to solving the linear program

$$(5) \quad \max \{c^T z : z_i \leq z_j \text{ for all } (i, j) \in \mathcal{I}, -1 \leq z_i \leq 1 \text{ for all } i \in V\},$$

where  $c_i = y_i - \bar{y}_V$ . If the optimal objective value equals zero, then the group  $V$  must be an optimal block.

This group-wise partitioning operation is the basis for our IRP algorithm which is detailed in Algorithm 1<sup>1</sup>. It starts with all observations as one group and recursively splits each group optimally by solving subproblem (5). At each iteration, a list  $\mathcal{C}$  of potential optimal partitions for each group generated thus far is maintained, and the partition among them with the highest objective value is performed. The list  $\mathcal{C}$  is updated with the optimal partitions generated from both sub-groups. Partitioning ends whenever the solution to (5) is trivial (i.e., no split is found because the group is a block). We can think of each iteration  $k$  of Algorithm 1 as producing a model  $M_k$  by fitting the average to each group in its current partition: For a set of groups  $\mathcal{V} = \{V_1, \dots, V_k\}$ , denote  $\bar{y}_{\mathcal{V}} = \{\bar{y}_{V_1}, \dots, \bar{y}_{V_k}\}$ . Then model  $M_k = (\mathcal{V}, \bar{y}_{\mathcal{V}})$  contains the partitioning  $\mathcal{V}$  as well as a fit to each of the observations, which is the mean observation of the group it belongs to in the partition.

---

**Algorithm 1** Isotonic Recursive Partitioning

---

**Require:** Observations  $(x_1, y_1), \dots, (x_n, y_n)$  and partial order  $\mathcal{I}$ .

**Require:**  $k = 0, \mathcal{A} = \{\{x_1, \dots, x_n\}\}, \mathcal{C} = \{(0, \{x_1, \dots, x_n\}, \{\})\}, \mathcal{B} = \{\}, M_0 = (\mathcal{A}, \bar{y}_{\mathcal{A}})$ .

```

1: while  $\mathcal{A} \neq \{\}$  do
2:   Let  $(val, w^-, w^+) \in \mathcal{C}$  be the potential partition with largest  $val$ .
3:   Update  $\mathcal{A} = (\mathcal{A} \setminus (w^- \cup w^+)) \cup \{w^-, w^+\}, \mathcal{C} = \mathcal{C} \setminus (val, w^-, w^+)$ .
4:    $M_k = (\mathcal{A}, \bar{y}_{\mathcal{A}})$ .
5:   for all  $v \in \{w^-, w^+\} \setminus \{\}$  do
6:     Set  $c_i = y_i - \bar{y}_v$  for all  $i$  such that  $x_i \in v$  where  $\bar{y}_v$  is the mean of the observations in set  $v$ .
7:     Solve LP (5) with input  $c$  and get  $z^* = \operatorname{argmax} \text{LP}(5)$ .
8:     if  $z_1^* = \dots = z_n^*$  (group is optimally divided) then
9:       Update  $\mathcal{A} = \mathcal{A} \setminus v$  and  $\mathcal{B} = \mathcal{B} \cup \{v\}$ .
10:    else
11:      Let  $v^- = \{x_i : z_i^* = -1\}, v^+ = \{x_i : z_i^* = +1\}$ .
12:      Update  $\mathcal{C} = \mathcal{C} \cup \{(c^T z^*, v^-, v^+)\}$ 
13:    end if
14:  end for
15:   $k = k + 1$ .
16: end while
17: return  $\mathcal{B}$ , a partitioning of observations corresponding to the optimal groups.
```

---

2.3. *Properties of the partitioning algorithm.* Theorem 2 next states the result which implies that the IRP partitions are *no-regret*. This will lead to our convergence result.

**Theorem 2** Assume group  $V$  is the union of blocks from the optimal solution to problem (2). Then a cut made by solving (3) (using 5) at a particular iteration does not cut through any block in the global optimal solution.

The fact that IRP is a no-regret algorithm can be shown using a connection between the work of Barlow and Brunk (1972) and Maxwell and Muckstadt (1985) (held to the Discussion in Section 7). We prove Theorem 2 directly, but leave it to the Appendix as the theorem is already known to be true (Spouge, Wan and Wilbur, 2003). Remark 7 in the Appendix handles the case for multiple observations. Since Algorithm 1 starts with  $\mathcal{A} = \{\{x_1, \dots, x_n\}\}$  which is a set of one group that is the union of all blocks, we can conclude from this

---

<sup>1</sup>This algorithm appeared in a shorter version of the paper (Luss, Rosset and Shahar, 2010).

theorem that IRP never cuts an optimal block when generating partitions. The following corollary is then a direct consequence of repeatedly applying Theorem 2 in Algorithm 1:

**Corollary 3** *Algorithm 1 converges to the optimal (isotonic block class) solution with no regret.*

Theorem 4 next states our main innovative result that Algorithm 1 provides isotonic solutions at each iteration. This result implies that the path of solutions generated by IRP can be regarded as a regularization path for isotonic regression. Along the path, the model grows in complexity until optimality. These suboptimal isotonic models often result in better predictive performance than the optimal solution, which is susceptible to overfitting as is discussed in Section 5.

**Theorem 4** *Model  $M_k$  generated after iteration  $k$  of Algorithm 1 is in the class  $\mathcal{G}$  of isotonic models.*

**Proof.** The proof is by induction. The base case, i.e. first iteration, where all points form one group is trivial. The first cut is made by solving the linear program (5) which constrains the solution to maintain isotonicity.

Assuming that iteration  $k$  (and all previous iterations) provides an isotonic solution, we prove that iteration  $k+1$  must also maintain isotonicity. Figure 2 helps illustrate the situation described here. Let  $G$  be the group split at iteration  $k+1$  and denote  $A$  ( $B$ ) as the group under (over) the cut. Let  $\mathcal{A} = \{X : X \text{ is a group at iteration } k+1, \text{ there exists } x_i \in X \text{ such that } (i, j) \in \mathcal{I} \text{ for some } x_j \in A\}$  (i.e.  $X \in \mathcal{A}$  border  $A$  from below).

Consider iteration  $k+1$ . Denote  $\mathcal{X} = \{X \in \mathcal{A} : \bar{y}_A < \bar{y}_X\}$  (i.e.  $X \in \mathcal{X}$  violates isotonicity with  $A$ ). The split in  $G$  causes the fit in nodes in  $A$  to decrease. We will prove that when the fits in  $A$  decrease, there can be no groups below  $A$  that become violated by the new fits to  $A$ , i.e. the decreased fits in  $A$  cannot be such that  $\mathcal{X} \neq \{\}$ .

We first prove that  $\mathcal{X} = \{\}$  by contradiction. Assume  $\mathcal{X} \neq \{\}$ . Denote  $i < k+1$  as the iteration at which the last of the groups in  $\mathcal{X}$ , denoted  $D$ , was split from  $G$  and suppose at iteration  $i$ ,  $G$  was part of a larger group  $H$  and  $D$  was part of a larger group  $F$ . It is important to note that  $X \cap (F \cup H) = \{\}$  for all  $X \in \mathcal{X} \setminus D$  at iteration  $i$  because by assumption all groups in  $\mathcal{X} \setminus D$  were separated from  $A$  before iteration  $i$ . Thus, at iteration  $i$ ,  $D$  is the only group bordering  $A$  that violates isotonicity.

Let  $D_U$  denote the union of  $D$  and all groups in  $F$  that majorize  $D$ . By construction,  $D_U$  is a majorant in  $F$ . Hence  $\bar{y}_{D_U} < \bar{y}_{F \cup H}$  by Algorithm 1 and  $\bar{y}_A < \bar{y}_{D_U}$  by definition since  $\bar{y}_{D_U} > \bar{y}_D > \bar{y}_A$ . Also by construction, any set  $X \in H$  that minorizes  $A$  has  $\bar{y}_X < \bar{y}_A$  (each set  $X$  that minorizes  $A$  besides  $D$  such that  $\bar{y}_X < \bar{y}_A$  has already been split from  $A$ ). Hence we can denote  $A_L$  as the union of  $A$  and all groups in  $H$  that minorize  $A$  and we have  $\bar{y}_A > \bar{y}_{A_L}$  and  $A_L$  is a minorant in  $H$ . Since  $A_L \subseteq H$  at iteration  $i$ , we have

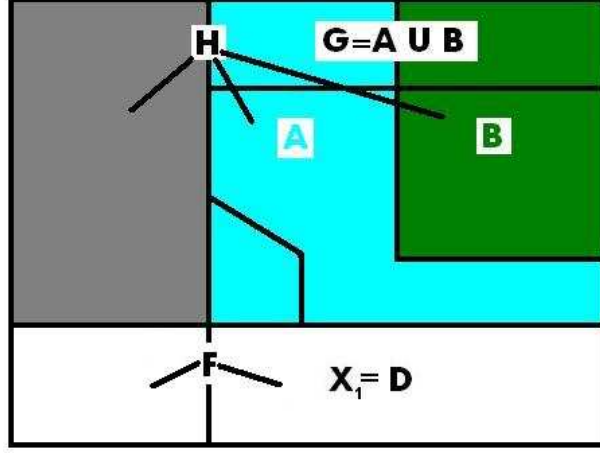
$$\bar{y}_{F \cup H} < \bar{y}_{A_L} < \bar{y}_A < \bar{y}_{D_U} < \bar{y}_{F \cup H}$$

which is a contradiction, and hence the assumption  $\mathcal{X} \neq \{\}$  is false. The first inequality is because the algorithm left  $A_L$  in  $H$  when  $F$  was split from  $H$ , and the remaining inequalities are due to the above discussion. Hence the split at iterations  $k+1$  could not have caused a break in isotonicity.

A similar argument can be made to show that the increased fit for nodes in  $B$  does not cause any isotonic violation. The proof is hence completed by induction. ■

With Theorem 4, the machinery for generating a regularization path is complete. In Section 3, we describe the computational complexity for generating this path followed by a discussion of the statistical complexity of the solutions along the path in Section 4.

**3. Complexity.** We here show that the partitioning step in IRP can be solved efficiently. The computational bottleneck of Algorithm 1 is solving linear program (5) that iteratively partitions each group. Linear



**Fig 2:** Illustration of proof of Theorem 4 showing the defined sets at iteration  $k + 1$ .  $G$  is the set divided at iteration  $k + 1$  into  $A$  (all blue area) and  $B$  (all green area). The group bordering  $A$  from below denoted by  $X_1$  (also referred to as  $D$  in the proof) is in violation with  $A$ . At iteration  $k_0$ ,  $G$  is part of the larger group  $H$  and  $X_1$  is part of the larger group  $F$ . At iteration  $k_0$ , groups  $F$  and  $H$  are separated. The proof shows that when  $A$  and  $B$  are split at iteration  $k + 1$ , no group such as  $X_1$  where  $w_{X_1} > w_A$  could have existed. In the picture,  $X_1$  must have been separated at an iteration  $k_0 < k + 1$ , but the proof, through contradiction, shows that this cannot occur.

program (5) has a special structure that can be taken advantage of in order to solve larger problems faster. Indeed, the dual problem can be written as an optimization problem called a network flow problem that is amenable to very efficient algorithms, as noted by Spouge, Wan and Wilbur (2003) who recognize the network flow problem as the *maximal upper set problem*. We note that our partition problem (5) is very similar to the network flow problem solved in Chandrasekaran et al. (2005) where  $z_i$  there represents the classification performance on node  $i$ . We denote the complexity of solving linear program (5) by  $T(m, n)$  where  $m$  is the number of constraints defined by  $\mathcal{I}$  and  $n$  is the number of observations. Various efficient algorithms for solving this problem exist, giving complexities such as  $T(m, n) = O(mn \log n)$  (Sleator and Tarjan, 1983) along with several algorithms giving  $T(m, n) = O(n^3)$  (Galil and Naamad, 1980). Choosing the more efficient implementation depends on the number of isotonicity constraints  $m$  (e.g.  $n^3 \leq mn \log n$  for the worst case  $m = O(n^2)$ ). A recent result by Stout (2010) shows how to represent an isotonic regression problem by an equivalent problem where both the number of total observations and constraints are of order  $O(n \log^{d-1} n)$ , which greatly reduces the worst case of  $m = O(n^2)$  isotonicity constraints (i.e. by trading off a few additional *shadow* observations for a large reduction in the number of constraints). Since at most  $n$  partitions are made by IRP, complexity is  $O(n^4)$  using  $T(m, n) = O(n^3)$  or reduced to  $O(n^3 \log^{2d-1} n)$  using results of Stout (2010).

In practice, the complexity can be even better by accounting for the fact that IRP solves a sequence of partitioning problems that are decreasing in size (i.e. problems with fewer and fewer observations). Each partition in Algorithm 1 can be divided into different proportions. We generically denote the bigger proportion in a partition by  $p \geq 0.5$ . Proposition 5 next gives a bound on the complexity of Algorithm 1 for this general case (assuming  $T(m, n) = O(n^3)$ ), in terms of the maximal  $p$  over all partitions.

**Proposition 5** Let  $p_{\max} \geq .5$  be the greatest  $p$  over all iterations of Algorithm 1 such that iteration  $k$  partitions a group of size  $n_k$  into two groups of size  $pn_k$  and  $(1 - p)n_k$ . Denote by  $n$  the total number of

observations. Then the complexity of Algorithm 1 is bounded by

$$(6) \quad O(n^3) \frac{1}{1 - p_{\max}^2}.$$

The proof, given in the Appendix, is based on the fact that the sequence of IRP's partition problems are solved on smaller and smaller groups of observations (i.e. while the first partition problem is  $O(n^3)$ , the partition problems for the two created partitions are  $O(p^3 n^3)$  and  $O((1 - p)^3 n^3)$  for some  $p$  where  $0 < p < 1$ ). Even at  $p_{\max} = .99$ , the constant  $1/(1 - p_{\max}^2) \approx 50$ , which is very small when the number of observations is large. Thus, under another reasonable assumption that  $p_{\max}$  is bounded, we can conclude that IRP is of practical complexity  $O(n^3)/(1 - p_{\max}^2)$ . Similar analysis using results of Stout (2010) lead to a practical complexity of  $O(n^2 \log^{2d-1} n)/(1 - p_{\max})$ .

Additional enhancements to the algorithm can be made in order to solve even larger problems than done here. As noted above, the dual problem to our optimal partition problem is a specially structured network flow problem. Luss, Rosset and Shahar (2010) show that the network flow problem (i.e. the dual to problem (5)) can be decomposed and solved through a sequence of smaller network flow problems. This reduction of the optimal partition problem makes IRP a practical tool for even larger data sets than experimented with here; refer to Luss, Rosset and Shahar (2010) for details on the large-scale decomposition which is beyond the scope of this paper.

**4. Degrees of freedom of isotonic regression and IRP.** The concept of degrees of freedom is commonly used in statistics to measure the complexity of a model (or more accurately, a modeling approach). This concept captures the amount of fitting the model performs, as expressed by the optimism of the in-sample error estimates, compared to out-of-sample predictive performance. Here we briefly review the main ideas of this general approach, and then apply them to isotonic regression and IRP.

Following Efron (1986) and Hastie, Tibshirani and Friedman (2001), assume the values  $x_1, \dots, x_n \in \mathbb{R}^d$  are fixed in advance (the *fixed-x* assumption), and that the model gets one vector of observations  $\mathbf{y} = (y_1, \dots, y_n)^\top \in \mathbb{R}^n$  for training, drawn according to  $P(y|x)$  at the  $n$  data points. Denote by  $\mathbf{y}^{\text{new}}$  another independent vector drawn according to the same distribution.  $\mathbf{y}$  is used for training a model  $\hat{f}(x)$  and generating predictions  $\hat{y}_i = \hat{f}(x_i)$  at the  $n$  data points.

We define the *in-sample* mean squared error:

$$\text{MRSS} = \frac{1}{n} \|\mathbf{y} - \hat{\mathbf{y}}\|_2^2$$

and compare it to the expected error the same model incurs on the new, independent copy, denoted in Hastie, Tibshirani and Friedman (2001) by  $\text{ERR}_{\text{in}}$ :

$$\text{ERR}_{\text{in}} = \frac{1}{n} \mathbb{E}_{\mathbf{y}^{\text{new}}} \|\mathbf{y}^{\text{new}} - \hat{\mathbf{y}}\|_2^2.$$

The difference between the two is the *optimism* of the in-sample prediction. As Efron (1986) and others have shown, the expected optimism in MRSS is:

$$(7) \quad E_{\mathbf{y}, \mathbf{y}^{\text{new}}}(\text{ERR}_{\text{in}} - \text{MRSS}) = \frac{2}{n} \sum_i \text{cov}(y_i, \hat{y}_i).$$

For linear regression with homoskedastic errors with variance  $\sigma^2$ , it is easy to show that (7) is equal to  $\frac{2}{n} d \sigma^2$ , where  $d$  is the number of regressors, hence the degrees of freedom. This naturally leads to defining the *equivalent degrees of freedom* of a modeling approach as:

$$(8) \quad df = \sum_i \text{cov}(y_i, \hat{y}_i) / \sigma^2.$$

In non-parametric models, one usually cannot calculate the actual degrees of freedom of a modeling approach, but it is often easier to generate *unbiased estimates*  $\hat{df}$  of  $df$  using Stein’s lemma (Stein, 1981). Meyer and Woodroffe (2000) demonstrate the applicability of this theory in shape-restricted non-parametric regression. Specifically, their Proposition 2, adapted to our notation, implies that if we assume the homoskedastic case  $\text{var}(y_i) = \sigma^2$  for all  $i$ , then the unbiased estimator  $\hat{df}$  for degrees of freedom in isotonic regression is the expected number of pieces  $D$  in the solution  $\hat{y}$  to (2), that is:

$$\mathbb{E}(D) = \sum_i \text{cov}(y_i, \hat{y}_i) / \sigma^2.$$

Considering the IRP algorithm, this puts us in the interesting situation where the number of steps the algorithm takes until it terminates in the globally optimal isotonic solution is equal to the degrees of freedom estimator of this global solution (minus one, since we start with one piece). One might thus be inclined to assume that each iteration of Algorithm 1 adds *about* one degree of freedom, i.e. performs approximately the same amount of fitting in every iteration. A similar idea is represented by the degrees of freedom calculation of Schell and Singh (1997) in their reduced monotonic regression algorithm (which starts from the complete isotonic fit and eliminates pieces).

On more careful consideration, however, it is obvious that this idea is incorrect since the first iteration of IRP finds an optimal cut in the very large space of all possible multivariate isotonic cuts. For comparison, a single deep split in a regression tree has been estimated to consume three or more degrees of freedom (Ye, 1998), and the space of possible splits in initial IRP iterations is much larger than that of a regression tree since IRP splits are not limited to being axis-oriented. Thus, intuitively, the first iteration is expected to use much more than one degree of freedom (the equivalent of fitting one coefficient to a *fixed, pre-determined* regressor). This effect should be exacerbated as the dimension  $d$  of  $x$  increases since the size of the search space for isotonic cuts increases with it. It also inevitably implies that the latter iterations of the IRP algorithm should perform less (ultimately much less) fitting than the equivalent of one degree of freedom in every iteration in order to be consistent with the unbiasedness of  $\hat{df} = D$  as an estimator of  $df$ .

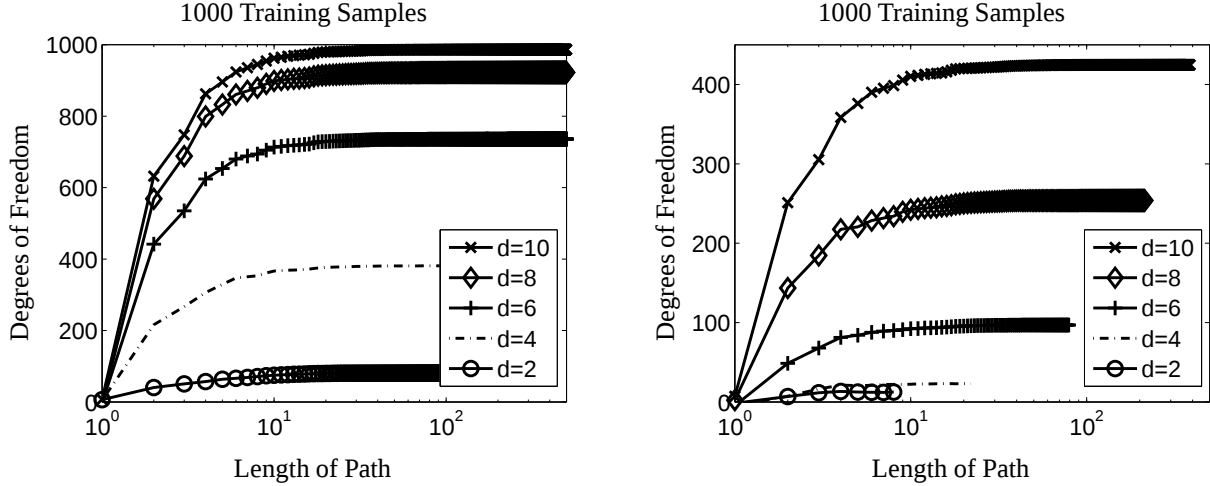
Here we demonstrate empirically that this is indeed the case. We simulate data from a simple additive model

$$(9) \quad x_{ij} \sim \mathcal{U}[1, 2] \text{ i.i.d.},$$

$$(10) \quad y_i = \sum_j x_{ij}^2 + \mathcal{N}(0, 10),$$

where  $x_{ij}$  is dimension  $j$  of the observation  $i$ . We can generate one fixed copy of data using (9), and then repeatedly generate observations using (10), apply IRP, and empirically estimate  $df$  as defined by (8) for every iteration of IRP. Figure 3 (left) shows how  $df$  evolves in this model as the IRP iterations proceed, for increasing dimensions of  $x$ . The covariance in (8) was estimated by drawing values  $X = (x_1, \dots, x_{1000})$  according to the model (9), fixing them, and repeatedly drawing 1000 independent copies of  $\mathbf{y}|X$  according to (10), and applying IRP on each one. 200 simulations were run with one drawing of  $X$  and the results were averaged. As expected, we see that the number of pieces (hence degrees of freedom) in the final isotonic regression increases with the dimension, as does the rate in which the number of degrees of freedom increases in the initial steps of IRP.

In order to emphasize this dependence of the degrees of freedom in initial iterations on the dimension, as well as on the number of observations, Figure 4 presents the evolution of the percentage of the total isotonic regression degrees of freedom along the path (i.e., number of degrees of freedom relative to the number of partitions of the final model) as a function of both the dimension and the amount of data used. As expected, increasing the dimension radically increases the portion of the fitting in the first steps, while

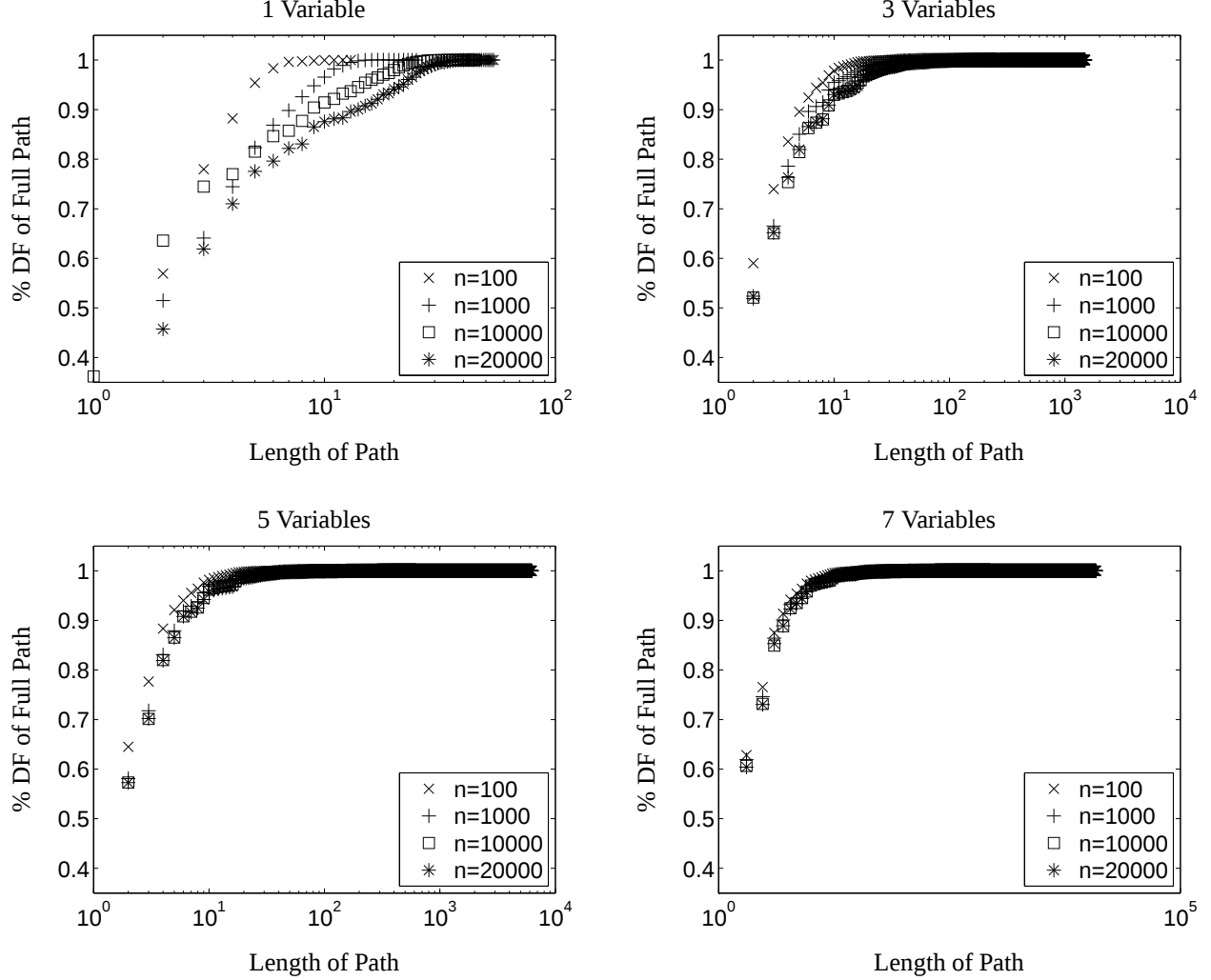


**Fig 3:** Evolution of degrees of freedom for IRP as model complexity increases. Both models use  $y_i = \sum_j x_{ij}^2 + \mathcal{N}(0, 10)$ . Simulation (left) uses  $x_{ij} \sim \mathcal{U}[1, 2]$  and simulation (right) uses  $x_{ij} \in \{0, 1, 2\}$  with probabilities  $\{1/3, 1/3, 1/3\}$ . Each path is the mean over 50 trials with 1000 training samples.

increasing the amount of data decreases this portion (since the overall isotonic fit is generally more complex in these situations). It should be noted that for many of the situations examined, IRP performs more than half of the total isotonic fit, as measured by degrees of freedom, in its first iteration! In dimension 7, even at  $n = 20,000$  observations, almost 60% of the total fitting is associated with the first iteration. Thus, these simulations clearly demonstrate the nature and limitations of IRP’s regularization behavior: the IRP path contains models that are regularized isotonic models compared to the global solution, but IRP’s ability to control model complexity is limited by the concentration of most of the fitting in the initial iterations, especially in higher dimension.

As mentioned in the introduction, an area that combines applications where the isotonic assumptions are reasonable (i.e., low bias) and the overfitting may be of less concern (i.e., variance can be controlled) is in genetics, specifically in modeling gene-gene interactions in phenotype(y)-genotype(X) relationships (Cordell, 2009). The key observation here is that genotypes are ternary ( $x_{ij} \in \{0, 1, 2\}$  copies of the “risk” allele). Thus, each dimension of the predictor space  $\mathcal{X}$  can take only one of three possible observed values in the data. Intuitively, it is clear that this would significantly reduce the space of possible isotonic splits in IRP, and hence reduce the amount of fitting. To demonstrate this empirically, Figure 3 (right) displays an experiment with the same setup, where instead of drawing the  $x$  values from a multivariate uniform, they are drawn independently from  $\{0, 1, 2\}$  with equal probabilities and we use the same model. Figure 3 demonstrates that both the globally optimal isotonic regression and, especially, the first IRP iterations perform much less overall fitting in the ternary case versus the continuous case, as measured by equivalent degrees of freedom. For example, in six dimensions, the continuous case requires almost seven times as many degrees of freedom than in the ternary case to fit the model. However, relative to the final model, a large percentage of the fitting still takes place in the initial iterations.

**5. Performance evaluation.** We demonstrate here the usefulness of isotonic regression on simulation and real data. For each experiment, IRP is run on the training data and produces a path of isotonic models. Each model is used for prediction on the test data and the root mean squared error (RMSE) is recorded. This generates paths of root mean squared errors over the different isotonic models and is illustrated in



**Fig 4:** Percentage of degrees of freedom relative to the full path (i.e. number of partitions in the optimal solution) as model complexity increases. Simulations use  $\mathbf{x}_{ij} \sim \mathcal{U}[1, 2]$  with  $y_i = \sum_j x_{ij}^2 + \mathcal{N}(0, 10)$ . Each path is the mean over 200 trials. Only the first 500 partitions of the paths are displayed in order to make the MSE of the earlier IRP iterations visually clearer.

the figures below. In each table, we record the minimum RMSE along these paths (*IRP Min RMSE*), along with how many partitions were made to generate this minimum RMSE (*IRP Min Path*), and the number of partitions in the global isotonic solution (*IRP Path Length*). IRP, as well as optimal isotonic regression, results are compared to running a least squares regression on the training data and predicting on the testing data with the resulting linear model (corresponding RMSE is called *LS RMSE*), and to the performance of the global isotonic regression solution (*Isotonic Regression RMSE*). Because we are interested in examining the behavior of the entire IRP path for use in selecting the optimal tuning parameter, and to avoid a significant increase in running time, we do not employ cross validation for selecting the best stopping point (number of iterations), but use test sets for this. In practical application, cross validation would be the appropriate approach for selecting the best model for prediction.

**5.1. Simulations.** We first illustrate isotonic regression on simulated data with different distributions. For the first three experiments, the  $i^{th}$  observation's regressors are distributed as  $x_{ij} \sim \mathcal{U}[0, 5]$ ,  $x_{ij} \sim \mathcal{U}[0, 3]$ ,

and  $x_{ij} \sim \mathcal{U}[0, 2]$ , respectively, and in all cases, all  $x_{ij}$  are i.i.d. Responses  $y_i$  for the three simulations are generated as

$$(1) y_i = (\sum_j x_{ij}^2) + \mathcal{N}(0, 4d^2), \quad (2) y_i = (\prod_{j=1}^d x_{ij}) + \mathcal{N}(0, d^2), \quad \text{and} \quad (3) y_i = 2\sum_j x_{ij} + \mathcal{N}(0, d^2),$$

respectively, where  $d$  is the dimension. The first simulation represents an additive (nonlinear) model and the other two simulations are “super-additive” models (i.e., represent strong positive interactions which are hard to approximate with additive effects). This allows us to examine the performance of IRP and isotonic regression in a variety of relevant situations.

The last two experiments are ternary. The  $i^{\text{th}}$  observation’s regressors are distributed as  $x_{ij} \in \{0, 1, 2\}$  with probabilities  $\{.7, .2, .1\}$  and  $\{1/3, 1/3, 1/3\}$  for the fourth and fifth experiments, respectively. The fourth model is subadditive while the fifth model is superadditive; specifically they are

$$(4) y_i = (\sum_j x_{ij})^{1/4} + \mathcal{N}(0, d^2/10) \quad \text{and} \quad (5) y_i = (\prod_j x_{ij}) + \mathcal{N}(0, d^2).$$

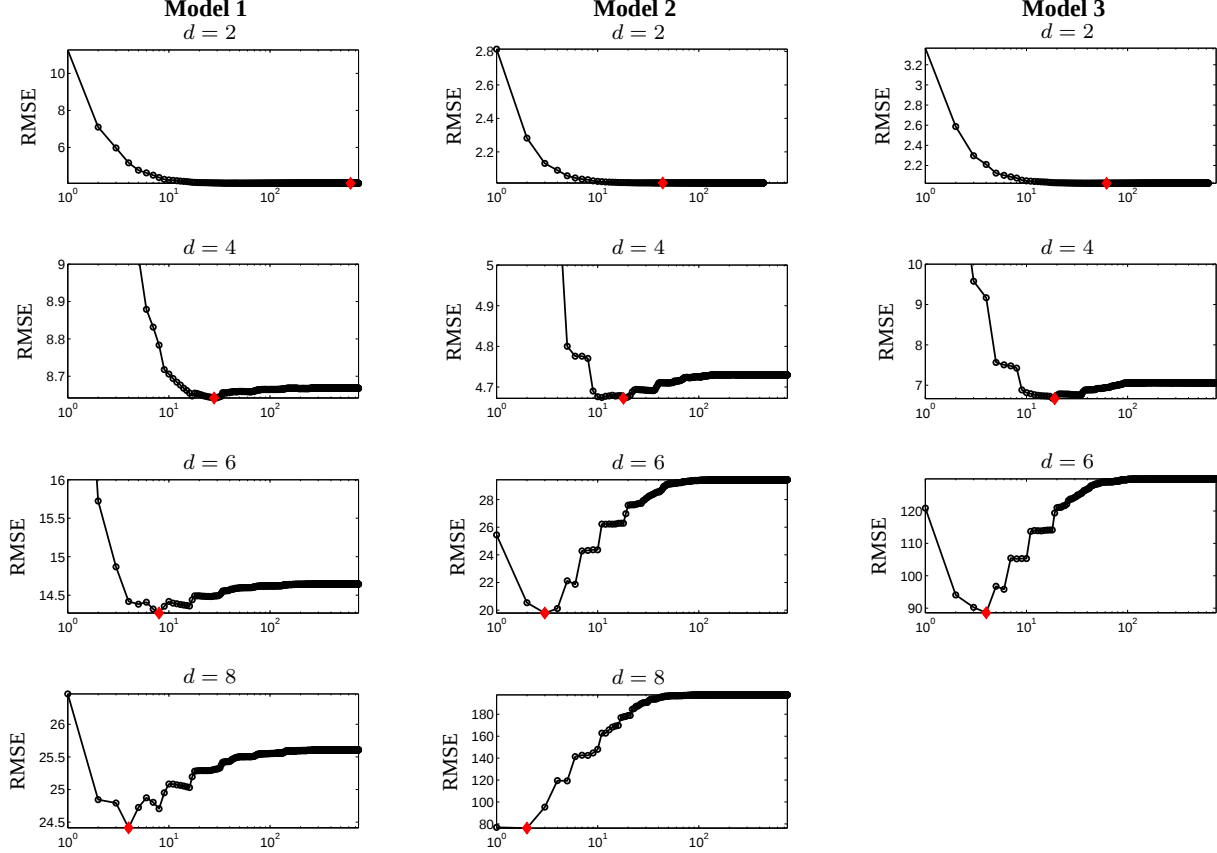
Model 4 is chosen to simulate “sub-additive” genetic interactions as discovered by [Shao et al. \(2008\)](#), where having one risk variant is sufficient to attain most of the effect, and additional variants have little additional influence. Model 5 represents the “super-additive” model, where variants exacerbate each other’s effect, as often speculated to be the case in human disease. Note that, for each of 50 simulations, 12000 training and 3000 testing points were randomly generated and statistics computed (all tests are out-of-sample).

Figure 5 demonstrates testing error for IRP over the regularized path of isotonic solutions for the first three experiments (with continuous covariates). The main observation here is that as the dimension increases, the effect of overfitting of the standard (non-regularized) isotonic regression becomes more significant and causes the skewed u-shaped pattern across the IRP path, where the minimum prediction RMSE is obtained earlier in the path. This is the effect we alluded to in the introduction and it stems from the limitations of the isotonicity constraints in controlling model complexity in high dimension.

Table 1 displays certain statistics on all five simulations as well as a comparison to the results of a least squares regression. We first discuss the case of continuous covariates (first three models). In lower dimensions standard isotonic regression performs well, and regularization through IRP offers no gain (this is seen in dimension  $d = 2$  in all three examples). Here, isotonic regression controls bias by accommodating the non-linearities in the true model and significantly outperforms least squares regression. As the number of covariates increases, regularization through IRP becomes necessary to control variance, and the optimal performance is obtained earlier in the IRP path (dimension  $d = 4$  in our examples). When  $d$  increases further, however, IRP also becomes inefficient at controlling variance, and linear regression dominates. This effect can be traced back to the large amount of fitting performed by IRP already in its initial iterations, as demonstrated in the previous section.

With respect to the models with ternary covariates, isotonic regression outperforms the simple linear regression, however the IRP path does not statistically improve performance. For the sub-additive model (Model 4), performance is better for dimensions 2 and 4, after which again IRP is unsuccessful at controlling variance. However, in the super-additive model (Model 5), IRP dominates for all dimensions. This is not surprising, since super-additive models are more extreme in their deviations from additivity, making the flexibility of IRP more critical.

Thus, our simulations confirm that isotonic regression performs well in low dimension, but requires a lot of data in order to learn good nonlinear models in higher dimensions. In intermediate dimension, IRP can offer a compromise between fitting flexible isotonic models and controlling model complexity, resulting in useful prediction models.



**Fig 5:** Root mean squared error (RMSE) for out-of-sample predictions of simulations with different dimensions  $d$ . The x-axis in each figure corresponds to the number of partitions made by IRP, i.e. the curves show how the RMSE of test data varies as the IRP algorithm progresses. Model 1 uses the function  $y_i = (\sum_j x_{ij}^2) + \mathcal{N}(0, 4d^2)$  with  $x_{ij} \sim \mathcal{U}[0, 5]$ . Model 2 uses the function  $y_i = (\prod_j x_{ij}) + \mathcal{N}(0, d^2)$  with  $x_{ij} \sim \mathcal{U}[0, 3]$ . Model 3 uses the function  $y_i = 2\sum_j x_{ij} + \mathcal{N}(0, d^2)$  with  $x_{ij} \sim \mathcal{U}[0, 2]$ . Fifty simulations were run with 12000 training and 3000 testing points. Only the first 750 partitions of the paths are displayed in order to make the RMSE of the earlier IRP iterations visually clearer.

**5.2. Modeling MPG of Automobiles.** We next illustrate the performance of IRP when predicting the miles-per-gallon of a list of 392 automobiles manufactured between 1970 and 1982 using seven variables (Frank and Asuncion, 2010). Seven regressions are performed in dimensions one through seven, where the variables chosen are from the following order: origin (American, European, or Japanese), model year, number of cylinders, acceleration, displacement, horsepower, and weight. The order of the variables was determined in order of the magnitude of coefficients from a least squares linear regression on all variables. Origin, surprisingly, had the largest magnitude, and in giving discrete variables 1,2,3 to the respective origins, there actually is a monotonic trend in origin (i.e., American cars are least fuel efficient, followed by European cars, with the Japanese being most efficient). While we include origin as a variable here, we note that similar performance for IRP was achieved without origin in an independent experiment.

Since the data set is rather small, we perform leave-one-out cross-validation (i.e. the data is divided into training and testing sets of 391 and 1 instances, respectively, so that each instance is used out-of-sample once). Table 2 displays certain statistics on the IRP, and isotonic regression, performance as well

Model 1:  $y_i = (\sum_j x_{ij}^2) + \mathcal{N}(0, 4d^2)$  with  $x_{ij} \sim \mathcal{U}[0, 5]$

| Number of Variables | IRP Min RMSE               | Isotonic Regression RMSE   | LS RMSE              | IRP Min Path | IRP Path Length |
|---------------------|----------------------------|----------------------------|----------------------|--------------|-----------------|
| 2                   | <b>4.06</b> ( $\pm 0.02$ ) | <b>4.06</b> ( $\pm 0.02$ ) | 4.80 ( $\pm 0.01$ )  | 625          | 999             |
| 4                   | <b>8.64</b> ( $\pm 0.03$ ) | <b>8.67</b> ( $\pm 0.04$ ) | 8.83 ( $\pm 0.03$ )  | 28           | 4861            |
| 6                   | 14.27 ( $\pm 0.06$ )       | 14.64 ( $\pm 0.07$ )       | 12.87 ( $\pm 0.04$ ) | 8            | 8520            |
| 8                   | 24.41 ( $\pm 0.11$ )       | 25.61 ( $\pm 0.12$ )       | 16.89 ( $\pm 0.05$ ) | 4            | 10604           |

Model 2:  $y_i = (\prod_j x_{ij}) + \mathcal{N}(0, d^2)$  with  $x_{ij} \sim \mathcal{U}[0, 3]$

| Number of Variables | IRP Min RMSE               | Isotonic Regression RMSE   | LS RMSE              | IRP Min Path | IRP Path Length |
|---------------------|----------------------------|----------------------------|----------------------|--------------|-----------------|
| 2                   | <b>2.01</b> ( $\pm 0.01$ ) | <b>2.01</b> ( $\pm 0.01$ ) | 2.13 ( $\pm 0.01$ )  | 44           | 437             |
| 4                   | <b>4.67</b> ( $\pm 0.03$ ) | <b>4.73</b> ( $\pm 0.04$ ) | 6.07 ( $\pm 0.03$ )  | 18           | 3711            |
| 6                   | 19.79 ( $\pm 0.24$ )       | 29.43 ( $\pm 0.97$ )       | 19.62 ( $\pm 0.29$ ) | 3            | 7685            |
| 8                   | 76.23 ( $\pm 2.08$ )       | 197.71 ( $\pm 16.49$ )     | 65.40 ( $\pm 2.24$ ) | 2            | 10153           |

Model 3:  $y_i = 2^{\sum_j x_{ij}} + \mathcal{N}(0, d^2)$  with  $x_{ij} \sim \mathcal{U}[0, 2]$

| Number of Variables | IRP Min RMSE               | Isotonic Regression RMSE   | LS RMSE              | IRP Min Path | IRP Path Length |
|---------------------|----------------------------|----------------------------|----------------------|--------------|-----------------|
| 2                   | <b>2.02</b> ( $\pm 0.01$ ) | <b>2.02</b> ( $\pm 0.01$ ) | 2.17 ( $\pm 0.01$ )  | 62           | 628             |
| 4                   | <b>6.67</b> ( $\pm 0.08$ ) | <b>7.06</b> ( $\pm 0.14$ ) | 10.10 ( $\pm 0.09$ ) | 19           | 7558            |
| 6                   | 88.59 ( $\pm 1.13$ )       | 129.91 ( $\pm 4.13$ )      | 70.77 ( $\pm 1.21$ ) | 4            | 11787           |

Model 4:  $y_i = (\sum_j x_{ij})^{1/4} + \mathcal{N}(0, d^2/10)$  with  $x_{ij} \in \{0, 1, 2\}$  with probabilities  $\{.7, .2, .1\}$

| Number of Variables | IRP Min RMSE               | Isotonic Regression RMSE   | LS RMSE             | IRP Min Path | IRP Path Length |
|---------------------|----------------------------|----------------------------|---------------------|--------------|-----------------|
| 2                   | <b>0.63</b> ( $\pm 0.00$ ) | <b>0.63</b> ( $\pm 0.00$ ) | 0.67 ( $\pm 0.00$ ) | 8            | 8               |
| 4                   | <b>1.27</b> ( $\pm 0.01$ ) | <b>1.27</b> ( $\pm 0.01$ ) | 1.29 ( $\pm 0.01$ ) | 9            | 30              |
| 6                   | 1.91 ( $\pm 0.01$ )        | 1.91 ( $\pm 0.01$ )        | 1.92 ( $\pm 0.01$ ) | 4            | 85              |
| 8                   | 2.55 ( $\pm 0.01$ )        | 2.56 ( $\pm 0.01$ )        | 2.54 ( $\pm 0.01$ ) | 5            | 267             |

Model 5:  $y_i = (\prod_j x_{ij}) + \mathcal{N}(0, d^2)$  with  $x_{ij} \in \{0, 1, 2\}$  with probabilities  $\{1/3, 1/3, 1/3\}$

| Number of Variables | IRP Min RMSE               | Isotonic Regression RMSE   | LS RMSE              | IRP Min Path | IRP Path Length |
|---------------------|----------------------------|----------------------------|----------------------|--------------|-----------------|
| 2                   | <b>2.00</b> ( $\pm 0.01$ ) | <b>2.00</b> ( $\pm 0.01$ ) | 2.11 ( $\pm 0.01$ )  | 4            | 5               |
| 4                   | <b>4.00</b> ( $\pm 0.02$ ) | <b>4.00</b> ( $\pm 0.02$ ) | 4.47 ( $\pm 0.02$ )  | 6            | 25              |
| 6                   | <b>6.04</b> ( $\pm 0.02$ ) | <b>6.04</b> ( $\pm 0.02$ ) | 7.27 ( $\pm 0.05$ )  | 87           | 103             |
| 8                   | <b>8.30</b> ( $\pm 0.04$ ) | <b>8.30</b> ( $\pm 0.04$ ) | 10.90 ( $\pm 0.21$ ) | 67           | 430             |

TABLE 1

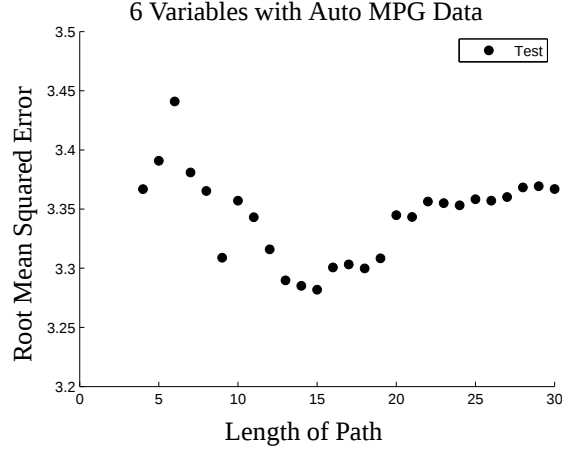
Statistics for simulations generated by the three different models as labeled above. A path of root mean squared errors (RMSE) for each model along the regularization path was computed. IRP Min RMSE refers to the minimum RMSE along these paths. Values for RMSE are given along with a conservative 95% confidence interval. LS RMSE is the RMSE from using least squares regressions. IRP Min Path is the number of partitions made to generate the minimum RMSE and IRP Path Length is the number of partitions in the global isotonic solution. Bolded RMSE values for IRP and isotonic regression indicate that they are lower than the RMSE of the least squares regression with 95% confidence.

as a comparison to the results of a least squares regression. Figure 6 displays MSE on out-of-sample data for IRP over the regularized path of isotonic solutions for a regression with six variables, exemplifying that overfitting occurs after 15 iterations of IRP (seen by the U-shaped curve with minimum at 15 iterations).

| Number of Variables | IRP Min RMSE                        | Isotonic Regression RMSE            | LS RMSE             | IRP Min Path | IRP Path Length |
|---------------------|-------------------------------------|-------------------------------------|---------------------|--------------|-----------------|
| 1                   | 6.46 ( $\pm 0.60$ )                 | 6.50 ( $\pm 0.43$ )                 | 6.46 ( $\pm 0.49$ ) | 9            | 17              |
| 2                   | <b>4.91 (<math>\pm 0.82</math>)</b> | <b>4.95 (<math>\pm 0.37</math>)</b> | 5.24 ( $\pm 0.37$ ) | 7            | 26              |
| 3                   | <b>3.73 (<math>\pm 1.17</math>)</b> | <b>3.76 (<math>\pm 0.36</math>)</b> | 4.02 ( $\pm 0.38$ ) | 9            | 37              |
| 4                   | 3.83 ( $\pm 1.13$ )                 | 3.91 ( $\pm 0.37$ )                 | 4.04 ( $\pm 0.38$ ) | 7            | 62              |
| 5                   | <b>3.32 (<math>\pm 1.41</math>)</b> | <b>3.36 (<math>\pm 0.38</math>)</b> | 3.92 ( $\pm 0.40$ ) | 15           | 109             |
| 6                   | <b>3.28 (<math>\pm 1.43</math>)</b> | 3.37 ( $\pm 0.37$ )                 | 3.77 ( $\pm 0.38$ ) | 15           | 114             |
| 7                   | 3.29 ( $\pm 1.43$ )                 | 3.36 ( $\pm 0.37$ )                 | 3.37 ( $\pm 0.33$ ) | 8            | 128             |

TABLE 2

Statistics for auto mpg data. Miles-per-gallon is regressed on a seven potential variables: origin (American, European, or Japanese), model year, number of cylinders, acceleration, displacement, horsepower, and weight. Row  $k$  uses the first  $k$  variables from this list in the regression. A path of root mean squared errors for each model along the regularization path was computed. Bold demonstrates statistical significance of either IRP or isotonic regression over a least squares regression with 95% confidence, as determined by a paired t-test using 392 observed squared losses obtained from leave-one-out cross-validation.



**Fig 6:** Root mean squared error (RMSE) for auto data with six variables using IRP illustrates that overfitting occurs after 15 iterations of IRP, i.e. RMSE decreases until 15 partitions are made, after which RMSE begins to increase. The figure displays only the first 30 partitions so that the U-shape is clear.

**6. The Search for Gene-Gene Interactions.** As mentioned in Section 1, the search for gene-gene interactions (epistasis) is a major endeavor in the genetics community, with the goal to identify the mechanisms that connect genotypes and phenotypes of interest, including height and disease.

It is generally acknowledged that the search has so far yielded very limited actual findings of major and impactful epistasis in human phenotypes (Cordell, 2009). Where significant interactions have been found, like recently by Zhang, Zhang and Liu (2011), these were often limited to close-by regions in the genome, where statistical and biological interactions are hard to differentiate (because mutation distributions are correlated due to linkage disequilibrium). If our interest is focused on biological interactions (e.g., two mutations influencing each other's causative effects on the phenotype), we should be particularly interested in finding epistasis between non-correlated mutations.

The limited findings in the literature may be due to two distinct reasons: First, the difficulty of searching through the space of all possible combinations of mutations. For example, in typical genome wide association studies (GWAS), one genotype has hundreds of thousands of single nucleotide polymorphisms (SNPs), yielding billions of potential two-way interactions, and much larger numbers of higher order interactions.

Heuristics for searching, like limiting search to interactions between SNPs that are individually associated with the phenotype, may miss combinations with weak main effects but strong interactions. Second, the limitation to simple statistical models like chi-square tests and logistic regression with explicit interactions may not allow identifying complex high-order interactions even within the searched space. IRP offers a remedy to this second concern, in allowing the modeling of complex interactions subject to isotonicity only. As our simulations have shown, good performance on ternary data is expected up to dimension six and even beyond when truly strong epistatic signal is present.

To empirically examine epistasis in human disease using IRP, we obtained data from the Wellcome Trust Case Control Consortium (WTCCC, 2007), encompassing around 2000 cases for each of three diseases: Crohn’s disease (CD), Type-I diabetes (T1D), and Type-II diabetes (T2D). These are compared to 3000 healthy controls. All of these samples were genotyped for around 500000 SNPs, and each phenotype (disease) was analyzed for association between each SNP status and case-control status using chi-square and logistic regression. For each disease, between five and nine significant SNPs were discovered after careful multiple comparison corrections (WTCCC, 2007).

We discuss in detail our modeling of the CD dataset. It comprises 2000 cases and 3000 controls. Our covariates include nine unlinked SNPs from seven different chromosomes, which were discovered as significant after a multiple comparison correction in the WTCCC GWAS and at least one other study (Hindorff *et al.*, 2011). These are genotypes, i.e. ternary variables in  $\{0, 1, 2\}$ . This is a classification problem with binary response (i.e.  $y_i \in \{0, 1\}$  for control vs case), and we thus model the risk of Crohn’s disease with an isotonic logistic regression, rather than an  $l_2$  isotonic regression as we have done for continuous regressions. Specifically, we fit isotonic models by maximizing the in-sample logistic log-likelihood rather than the sum of squares:

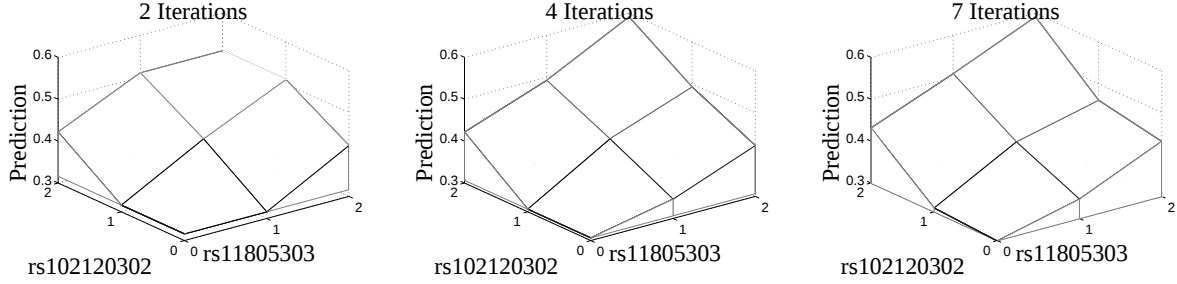
$$(11) \quad \max \left\{ \sum_{i=1}^n y_i \log(\hat{p}_i) + (1 - y_i) \log(1 - \hat{p}_i) : \hat{p}_i \leq \hat{p}_j \text{ for all } (i, j) \in \mathcal{I} \right\}.$$

As explicated by Bacchetti (1989) and Auh and Sampson (2006), the solution to (11) is identical to the solution from solving the  $l_2$  squared loss isotonic regression problem (2) when the values  $y_i \in \{0, 1\}$ . Since we solve  $l_2$  isotonic regression (2) using the IRP algorithm, we can use it to solve the isotonic logistic regression problem.

To evaluate isotonic model performance, we compare cross-validated area under the ROC curve (AUC) for each model in the IRP path to that of the linear logistic regression model with the same covariates that assumes no interaction between the SNPs. We employ 150-fold cross validation, and use the approach of DeLong, DeLong and Clarke-Pearson (1988) to calculate p values on the holdout AUC difference.

Figure 7 illustrates the approach on a subset of two SNPs. We can see the results of IRP after 2,4 and 7 iterations (the complete path has 8 iterations). The model at iteration 7 gives AUC of 0.5562, compared to 0.5501 for logistic regression, yielding a p-value of 0.0282 (which is only mildly indicative of a possible interaction given multiple comparison issues, as we hand-picked the two mutations and the number of iterations). The fit at iteration 7 seems to support a super-additive interaction: presence of two copies each from both SNPs confers a jump in CD risk.

We applied IRP on all subsets of three and four SNPs, and also the entire set of nine SNPs. Several models yielded AUC improvements over logistic regression which were significant at the nominal 0.05 level (like the two-variable model above), but none would withstand a multiple comparison correction. With all SNPs, the full path length was 96 iterations. The minimum model AUC for isotonic logistic regression was 0.6457 at 55 iterations and for logistic regression was 0.6449. Because AUC was highly variable between folds for this high-dimensional model, this difference was insignificant: p-value was 0.4. The final model had an AUC of 0.6456 and the p-value was 0.42.



**Fig 7:** Illustration of IRP with logistic log-likelihood loss function for modeling Crohn’s disease, as a function of two SNPs: rs11805303 on chromosome 1 and rs102120302 on chromosome 2. Models after iterations 2, 4, and 7 of IRP are shown. The model at iteration 7 gives the best predictive performance along the path.

Our results on the T1D and T2D datasets were qualitatively similar: no high order interactions of significance were identified by IRP. As the IRP approach possesses the flexibility and power to identify such interactions if they are strong enough (as illustrated in our simulations), we can conclude that the strong univariate signals identified by the WTCCC studies do not yield significant higher order interactions. This is of course in line with previous findings (Cordell, 2009; Emily *et al.*, 2009), but now also confirmed by our more sophisticated and flexible approach.

**7. Discussion and extensions.** The IRP algorithm offers solutions to both the statistical and computational difficulties of isotonic regression. Algorithmically, IRP solves (2) as a sequence of easier binary partitioning problems that are efficiently solved using network flow algorithms. From the statistical perspective, IRP generates a path of isotonic models, each defining a partitioning of the space  $\mathcal{X}$  into isotonic regions. The averages of observations in these regions comply with the isotonicity constraints (Theorem 4). In this view, IRP provides isotonic solutions along its path that are regularized versions of the globally optimal isotonic regression solution.

Our discussion so far has focused on using the sum of squares loss function in (2) for fitting “standard” isotonic regression subject to squared error loss as well the logistic log-likelihood which we noted is an identical problem. A well-known result of Barlow and Brunk (1972) implies that the solution of a whole variety of loss functions subject to isotonicity constraints can be obtained by solving standard isotonic regression, as long as the loss can be written as minimizing  $\sum_{i=1}^n w_i (\Psi(z_i) - x_i z_i)^2$  in  $z \in \mathbf{R}^n$  for some convex differentiable  $\Psi$  and some data-dependent values  $x$  and weights  $w$ . These results imply that many other loss functions subject to isotonicity constraints can optimally be solved by IRP via a reformulation to a problem of the form (2). Barlow and Brunk (1972) note that such transformations can be applied to many maximum likelihood estimation problems (subject to isotonicity constraints), including Bernoulli (as described in Section 6), multinomial, poisson, and gamma distributed problems. We plan to investigate the applicability of the resulting regularization algorithms in future work.

We can now also formalize the connection between IRP and the well-known work of Maxwell and Muckstadt (1985) (and similarly Roundy (1986)) that was mentioned in the introduction. Maxwell and Muckstadt (1985) solved an operations research problem (related to scheduling reorder intervals for a production system) by reducing it to the optimization of a convex objective subject to isotonicity constraints. In our notation, their objective (i.e. loss function) is  $\sum_{i=1}^n (c_i/\hat{y}_i + b_i\hat{y}_i)$ , where  $c_i, b_i$  are data-dependent nonnegative constants determined by their problem formulation. To apply the theory of Barlow and Brunk (1972), we reformulate their problem as minimizing  $\sum_{i=1}^n c_i (z_i - (-b_i/c_i))^2$  in  $z \in \mathbf{R}^n$ , i.e. a standard weighted isotonic regression, and recovering  $\hat{y}_i^* = \sqrt{-z_i^*}$  (note that the isotonic regression fits nonpositive observa-

tions  $-b_i/c_i$ ). Indeed, the algorithm of [Maxwell and Muckstadt \(1985\)](#) is completely equivalent to applying IRP on this modified problem! It should be emphasized, however, that [Maxwell and Muckstadt \(1985\)](#) were interested in this algorithm purely as a means to reach the optimal solution, and were uninterested in statistical considerations which led us to consider intermediate IRP solutions as regularized isotonic models of independent interest. [Spouge, Wan and Wilbur \(2003\)](#) also used [Maxwell and Muckstadt \(1985\)](#) to inspire the partitioning algorithm for the standard isotonic regression problem, however they do not make the connection using [Barlow and Brunk \(1972\)](#), and also have no statistical interests in mind.

As our analysis and experiments have demonstrated, computation is not a significant concern with IRP, at least for moderate to large data sizes. However, overfitting is still a major concern as dimensionality grows. As demonstrated in Sections 4 and 5, while IRP offers partial protection from overfitting through its regularization behavior, even the first step in the IRP path could already suffer from high variance in dimension as low as six. A key question pertains to identifying factors affecting this overfitting behavior, specifically characterizing situations in which the initial IRP iterations are less prone to overfitting.

## 8. Appendix.

### Proposition 1:

**Proof.** We first rewrite both the IRP partition problem (3) and the maximal between-group variance partition problem (4). Assume  $V = A \cup B$  and  $A \cap B = \{\}$ . Then it is easy to show  $|V|\bar{y}_V = |A|\bar{y}_A + |B|\bar{y}_B$  which gives  $(\bar{y}_A - \bar{y}_V) = -|B|(\bar{y}_B - \bar{y}_V)/|A|$ . The objective function to (3) can be written  $|B|(\bar{y}_B - \bar{y}_V) - |A|(\bar{y}_A - \bar{y}_V)$  and using the previous relationship can again be rewritten  $2|B|(\bar{y}_B - \bar{y}_V)$ . An obvious property of the optimal IRP cut is that  $\bar{y}_B \geq \bar{y}_V$ . If we add this as a redundant constraint to the IRP partition (3), then we can find the same optimal partition by maximizing the square of the objective, i.e. maximize  $4|B|^2(\bar{y}_B - \bar{y}_V)^2$  subject to the appropriate constraints. The objective of the between-group variance partition (4) can be rewritten using the above relationship as  $(|B| + |B|^2/|A|)(\bar{y}_B - \bar{y}_V)^2$ . Then denoting the IRP and maximal between-group variance objectives by  $g^*(A, B)$  and  $\tilde{g}(A, B)$  respectively, we have  $g^*(A, B) = 4|A||B|\tilde{g}(A, B)/n$  since  $|A| + |B| = n$  is constant. Eliminating the constant  $4/n$  gives the first result.

In order to prove the second statement, notice that optimality of (3) and (4) gives  $|A^*||B^*|\tilde{g}(A^*, B^*) \geq |A^*||B^*|\tilde{g}(\tilde{A}, \tilde{B})$  and  $\tilde{g}(\tilde{A}, \tilde{B}) \geq \tilde{g}(A^*, B^*)$  which implies  $|A^*||B^*| \geq |\tilde{A}||\tilde{B}|$ . This along with the relation

$$(|A^*| + |B^*|)^2 = |A^*|^2 + 2|A^*||B^*| + |B^*|^2 = |\tilde{A}|^2 + 2|\tilde{A}||\tilde{B}| + |\tilde{B}|^2 = (|\tilde{A}| + |\tilde{B}|)^2$$

gives  $|A^*|^2 + |B^*|^2 \leq |\tilde{A}|^2 + |\tilde{B}|^2$ . We use this to get the relation

$$\begin{aligned} (|A^*| - |B^*|)^2 &= |A^*|^2 - 2|A^*||B^*| + |B^*|^2 \\ &\leq |\tilde{A}|^2 - 2|A^*||B^*| + |\tilde{B}|^2 \leq |\tilde{A}|^2 - 2|\tilde{A}||\tilde{B}| + |\tilde{B}|^2 = (|\tilde{A}| - |\tilde{B}|)^2 \end{aligned}$$

which gives the second result of the proposition. ■

### Theorem 2:

**Proof.** Divide the blocks in  $V$  into three subsets:

1.  $\mathcal{L}$ : union of all blocks in  $V$  that are “below” the algorithm cut.
2.  $\mathcal{U}$ : union of all blocks in  $V$  that are “above” the algorithm cut.
3.  $\mathcal{M}$ : union of  $K$  blocks in  $V$  that get broken by the cut (note that blocks in  $\mathcal{M}$  may be separated by blocks in  $\mathcal{L}$  or  $\mathcal{U}$ ).

Define  $M_1$  ( $M_K$ ) to be the minorant (majorant) block in  $\mathcal{M}$ . For each  $M_k$  define  $M_k^L$  ( $M_k^U$ ) as the groups in  $M_k$  below (above) the algorithm cut. Define  $A_K^L \subseteq \mathcal{L}$  ( $A_1^U \subseteq \mathcal{U}$ ) as the union of blocks along

the algorithm cut such that  $A_K^L \succ M_K^L$  ( $A_1^U \prec M_1^U$ ). Refer to Figure 8 for an example of these definitions where  $A_1^U = A_1^L = A_K^U = A_K^L = \{\}$  for simplicity.

We use the above definitions and assumptions to state the following two consequences that cause a contradiction:

- I  $\bar{y}_{M_1} < \bar{y}_{M_K}$  by optimality (i.e. according to KKT conditions) and isotonicity.
- II  $\bar{y}_{M_1} > \bar{y}_V$  and  $\bar{y}_{M_K} < \bar{y}_V$ . This is proven below.

(II) implies  $\bar{y}_{M_1} > \bar{y}_{M_K}$  which contradicts (I) and we are left to prove (II). Optimality of blocks  $M_1$  and  $M_K$  gives

- (a)  $\bar{y}_{M_1^L} > \bar{y}_{M_1^U}$
- (b)  $\bar{y}_{M_K^L} > \bar{y}_{M_K^U}$ .

The proof for  $\bar{y}_{M_1} > \bar{y}_V$  is as follows with two cases:

1.  $A_1^U = \{\}$ :  $\bar{y}_{M_1^U} > \bar{y}_V$  because using the algorithm cut in (5), we have

$$\sum_{i: x_i \in M_1^U} (y_i - \bar{y}_V) > 0 \Rightarrow \sum_{i: x_i \in M_1^U} y_i > |M_1^U| \bar{y}_V \Rightarrow \bar{y}_{M_1^U} > \bar{y}_V.$$

The first inequality is true about the cut because there exist no block below  $M_1^U$  to affect isotonicity. Then using (a), we get

$$\bar{y}_{M_1^L} > \bar{y}_{M_1^U} > \bar{y}_V \Rightarrow \bar{y}_{M_1} > \bar{y}_V$$

2.  $A_1^U \neq \{\}$ :  $\bar{y}_{M_1} > \bar{y}_{A_1^U} > \bar{y}_V$ . The first inequality is due to optimality and the second is again because the algorithm cut in (5) gives

$$\sum_{i: x_i \in A_1^U} (y_i - \bar{y}_V) > 0 \Rightarrow \sum_{i: x_i \in A_1^U} y_i > |A_1^U| \bar{y}_V \Rightarrow \bar{y}_{A_1^U} > \bar{y}_V,$$

which again is possible because no block exists below  $A_1^U$  to affect isotonicity.

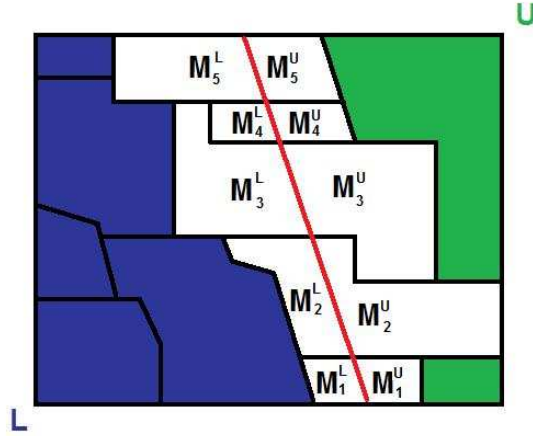
The proof for  $\bar{y}_{M_K} < \bar{y}_V$  is a similar argument and hence gives (II). The case  $K = 1$  is also trivially covered by the above arguments. We conclude that the algorithm cannot cut any block. ■

The following remark is necessary for completeness of the proof of Theorem 2.

**Remark 6** The case of two connected optimal groups having equal means need not be discussed in Theorem 2. In this event, the optimal solution to isotonic regression is not unique. It is trivial that  $M_1$  would not have been split by Algorithm 1 if  $\bar{y}_{M_1^L} = \bar{y}_{M_1^U} \neq \bar{y}_V$ . Otherwise, consider the case  $\bar{y}_{M_1^L} = \bar{y}_{M_1^U} = \bar{y}_V$  and assume  $M_1$  is a block broken by the cut in  $V$ .  $M_1^L$  and  $M_1^U$  are also possible blocks whereby  $M_1^L \in \mathcal{L}$  and  $M_1^U \in \mathcal{U}$ , and hence  $M_1 = M_1^L \cup M_1^U \notin \mathcal{M}$ . The same remarks apply to  $M_K$ . Thus, the proof still holds if there are multiple isotonic solutions.

**Remark 7** The case of multiple observations at the same coordinates can be disregarded. To see this, let  $J$  be a set of nodes with the same coordinates. From the constraints,  $y_i = y_j$ , for all  $i, j \in J$  and thus the number of observations can be reduced and all observations in  $J$  fit to the same value  $\hat{y}$ . Then

$$\sum_{j: x_j \in J} (\hat{y} - y_j)^2 = |J| (\hat{y} - \bar{y}_J)^2 + \sum_{j: x_j \in J} y_j^2 - \bar{y}_J^2$$



**Fig 8:** Illustration of proof of Theorem 2. Black lines separate blocks. The diagonal red line through the center demonstrates a cut of Algorithm 1.  $\mathcal{L}$  is the union of blue blocks below the cut and  $\mathcal{U}$  is the union of green blocks above the cut. White blocks are blocks that are potentially split by Algorithm 1. These blocks are split into  $M_1^L, \dots, M_5^L$  below the cut and  $M_1^U, \dots, M_5^U$  above the cut. In the proof,  $M_i = M_i^L \cup M_i^U$  for all  $i = 1 \dots 5$ . The proof shows, for example, that if the algorithm splits  $M_1$  into  $M_1^L$  and  $M_1^U$  according to the defined cut in (5), then there must be no isotonicity violation when creating blocks from  $M_1^L$  and  $M_1^U$ . However, since  $M_1$  is assumed to be a block, there must exist an isotonicity violation between  $M_1^L$  and  $M_1^U$ , providing a contradiction.

so that the sum of squared differences over  $J$  can be reduced to be a single weighted squared difference. Problem (2) becomes the weighted isotonic regression problem

$$(12) \quad \min \left\{ \sum_{i=1}^n w_i (\hat{y}_i - y_i)^2 : \hat{y}_i \leq \hat{y}_j \text{ for all } (i, j) \in \mathcal{I} \right\},$$

for which the KKT conditions imply that observations are again divided into  $k$  groups where the fits in each group take the weighted group mean  $\bar{y}_V^w = \sum_{i: x_i \in V} (w_i y_i) / \sum_{i: x_i \in V} w_i$  rather than the group mean. The optimal cut problem (5) changes to have  $z_i = w_i (y_i - \bar{y}_V^w)$  and the above results on IRP generalize easily noting that now the weighted algorithm cut implies  $\bar{y}_A^w > \bar{y}_V^w$  for a group  $A$  on the upper side of the cut such that no group exists below  $A$  that could affect isotonicity.

### Proposition 5:

**Proof.** Any final partition can be represented by a simple tree. Consider level  $k$  of the tree. Let  $p_k \geq .5$  be the greatest  $p$  over levels  $1, \dots, k-1$  such that a partition of group size  $n_k$  into two groups of size  $pn_k$  and  $(1-p)n_k$  where  $n_k$  is the corresponding size of the partitioned group. Denote by  $L_k$  the largest group partitioned at iteration  $k$  whose size can be bounded by  $|L_k| \leq np_k^k$ . We next note that the complexity of solving a problem with  $n$  observations is higher than solving 2 problems with  $pn$  and  $(1-p)n$  observations. Indeed,  $n^3 = pn^3 + (1-p)n^3 > p^2n^3 + (1-p)^2n^3$ . Thus, we assume that at iteration  $k$ , we solve only problems of the largest possible size (rather than several problems of small size). The number of groups at iteration  $k$  can also be bounded by  $n/|L_k|$ . Denote by  $T_{p_k}(k)$  the complexity of partitioning all groups at level  $k$ . Then

$$T_{p_k}(k) \leq O\left(\frac{n}{|L|}|L|^3\right) = O(n|L|^2) \leq O(n(np_k^k)^2) = O(n^3)p_k^{2k}.$$

Then denote by  $K$  the total number of levels in the partition tree. We have

$$\sum_{k=1}^K T_{p_k}(k) \leq \sum_{k=1}^K O(n^3) p_{\max}^{2k} \leq \sum_{k=1}^{\infty} O(n^3) p_{\max}^{2k} = O(n^3) \frac{1}{1 - p_{\max}^2}.$$

■

**9. Acknowledgements.** The authors are grateful to Quentin Stout for drawing our attention to additional relevant references. We thank the editor, associate editor and referee for their thoughtful and useful comments. This study makes use of data generated by the Wellcome Trust Case-Control Consortium. A full list of the investigators who contributed to the generation of the data is available from [www.wtccc.org.uk](http://www.wtccc.org.uk). R. Luss and S. Rosset were partially supported by Israeli Science Foundation grant 1227/09 and an IBM Open Collaborative Research grant.

## References.

- AUH, S. and SAMPSON, A. R. (2006). Isotonic logistic discrimination. *Biometrika* **93** 961-972.
- BACCHETTI, P. (1989). Additive Isotonic Model. *Journal of the American Statistical Association* **84** 289-294.
- BARLOW, R. E. and BRUNK, H. D. (1972). The Isotonic Regression Problem and Its Dual. *Journal of the American Statistical Association* **67** 140-147.
- BLOCK, H., QIAN, S. and SAMPSON, A. (1994). Structure Algorithms for Partially Ordered Isotonic Regression. *Journal of Computational and Graphical Statistics* **3** 285-300.
- BOYD, S. and VANDENBERGHE, L. (2004). *Convex Optimization*. Cambridge University Press.
- BREIMAN, L., FRIEDMAN, J., STONE, C. J. and OLSHEN, R. A. (1984). *Classification and Regression Trees*. Chapman and Hall/CRC.
- CHANDRASEKARAN, R., RYU, Y. U., JACOB, V. S. and HONG, S. (2005). Isotonic Separation. *INFORMS Journal on Computing* **17** 462-474.
- CORDELL, H. J. (2009). Detecting genegene interactions that underlie human diseases. *Nature Reviews Genetics* **10** 392.
- DE LEEUW, J., HORNIK, K. and MAIR, P. (2009). Isotone Optimization in R: Pool-Adjacent-Violators Algorithm (PAVA) and Active Set Methods. UC Los Angeles: Department of Statistics, UCLA. Retrieved from: <http://cran.r-project.org/web/packages/isotone/vignettes/isotone.pdf>.
- DELONG, E. R., DELONG, D. M. and CLARKE-PEARSON, D. L. (1988). Comparing the Areas under Two or More Correlated Receiver Operating Characteristic Curves: A Nonparametric Approach. *Biometrics* **44** 837-845.
- DYKSTRA, R. L. and ROBERTSON, T. (1982). An Algorithm for Isotonic Regression for Two or More Independent Variables. *Annals of Statistics* **10** 708-716.
- EFRON, B. (1986). How Biased is the Apparent Error Rate of a Prediction Rule? *Journal of the American Statistical Association* **81** 461-470.
- EICHLER, E. E., FLINT, J., GIBSON, G., KONG, A., LEAL, S. M., MOORE, J. H. and NADEAU, J. H. (2010). Missing heritability and strategies for finding the underlying causes of complex disease. *Nature Reviews Genetics* **11** 446-450.
- EMILY, M., MAILUND, T., HEIN, J., SCHAUER, L. and SCHIERUP, M. H. (2009). Using biological networks to search for interacting loci in genome-wide association studies. *Eur J Hum Genet.* **17** 1231.
- FRANK, A. and ASUNCION, A. (2010). UCI Machine Learning Repository. Auto MPG Data Set available at <http://archive.ics.uci.edu/ml>.
- GALIL, Z. and NAAMAD, A. (1980). An  $O(EV \log^2 V)$  Algorithm for the Maximal Flow Problem. *Journal of the Computer and System Sciences* **21** 203-217.
- GNEITING, T. (2011). Making and Evaluating Point Forecasts. *JASA* **106** 746-762.
- GOLDSTEIN, D. B. (2009). Common Genetic Variation and Human Traits. *NEJM* **360** 1696-1698.
- HASTIE, T., TIBSHIRANI, R. and FRIEDMAN, J. (2001). *The Elements of Statistical Learning*. Springer.
- HE, X., NG, P. and PORTNOY, S. (1998). Bivariate Quantile Smoothing Splines. *Journal of the Royal Statistical Society. Series B* **60** 537-550.
- HINDORFF, L. A., JUNKINS, H. A., HALL, P. N., MEHTA, J. P. and MANOLIO, T. A. (2011). A Catalog of Published Genome-Wide Association Studies. Available at: [www.genome.gov/gwastudies](http://www.genome.gov/gwastudies).
- HOCHBAUM, D. S. and QUEYRANNE, M. (2003). Minimizing a Convex Cost Closure Set. *SIAM Journal of Discrete Mathematics* **16** 192-207.

- KRUSKAL, J. B. (1964). Multidimensional scaling by optimizing goodness of fit to a nonmetric hypothesis. *Psychometrika* **29**.
- LEE, C. I. C. (1983). The Min-Max Algorithm and Isotonic Regression. *The Annals of Statistics* **11** 467–477.
- LUSS, R., ROSSET, S. and SHAHAR, M. (2010). Decomposing Isotonic Regression for Efficiently Solving Large Problems. Proceedings of the Neural Information Processing Systems Conference, 2010.
- MANI, R., ONGE, R. P. S., HARTMAN, J. L., GIAEVER, G. and ROTH, F. P. (2007). Defining genetic interaction. *PNAS* **105** 3461–3466.
- MAXWELL, W. L. and MUCKSTADT, J. A. (1985). Establishing Consistent and Realistic Reorder Intervals in Production-Distribution Systems. *Operations Research* **33** 1316–1341.
- MEYER, M. and WOODROOFE, M. (2000). On the Degrees of Freedom in Shape-Restricted Regression. *Annals of Statistics* **28** 1083–1104.
- MONTEIRO, R. D. C. and ADLER, I. (1989). Interior path following primal-dual algorithms. Part II: Convex Quadratic Programming. *Mathematical Programming* **44** 43–66.
- OBOZINSKI, G., LANCKRIET, G., GRANT, C., JORDAN, M. I. and NOBLE, W. S. (2008). Consistent probabilistic outputs for protein function prediction. *Genome Biology* **9** 247–254. Open Access.
- PARDALOS, P. M. and XUE, G. (1999). Algorithms for A Class of Isotonic Regression Problems. *Algorithmica* **23** 211–222.
- ROTH, F. P., LIPSHITZ, H. D. and ANDREWS, B. J. (2009). Q&A: Epistasis. *Journal of Biology* **8**. Article 35.
- ROUNDY, R. (1986). A 98%-Effective Lot-Sizing Rule For a Multi-Product, Multi-Stage Productoin/Inventory System. *Mathematics of Operations Research* **11** 699–727.
- SCHELL, M. J. and SINGH, B. (1997). The Reduced Monotonic Regression Method. *Journal of the American Statistical Association* **92** 128–135.
- SHAO, H., BURRAGE, L. C., SINASAC, D. S., HILL, A. E., ERNEST, S. R., O'BRIEN, W., COURTLAND, H.-W., JEPSEN, K. J., KIRBY, A., KULBOKAS, E. J., DALY, M. J., BROMANG, K. W., LANDER, E. S. and NADEAU, J. H. (2008). Genetic architecture of complex traits: Large phenotypic effects and pervasive epistasis. *PNAS* **105** 11910–11914.
- SLEATOR, D. and TARJAN, R. E. (1983). A data structure for dynamic trees. *Journal of Computer and System Sciences* **26** 362–391.
- SPOUGE, M., WAN, H. and WILBUR, W. J. (2003). Least Squares Isotonic Regression in Two Dimensions. *Journal of Optimization Theory and Applications* **117** 585–605.
- STEIN, C. M. (1981). Estimation of the Mean of a Multivariate Normal Distribution. *The Annals of Statistics* **9** 1131–1151.
- STOUT, Q. (2010). An Approach to Computing Multidimensional Isotonic Regressions. Submitted. Available at: <http://www.eecs.umich.edu/qstout/pap/MultidimIsoReg.pdf>.
- WTCCC, (2007). Genome-wide association study of 14,000 cases of seven common diseases and 3,000 shared controls. *Nature* **447** 661–678.
- YE, J. (1998). On Measuring and Correcting the Effects of Data Mining and Model Selection. *Journal of the American Statistical Association* **93** 120–131.
- ZHANG, Y., ZHANG, J. and LIU, J. S. (2011). Block-based Bayesian Epistasis Association Mapping with Application to WTCCC Type 1 Diabetes Data. *Annals of Applied Statistics*, to appear.

SCHOOL OF MATHEMATICAL SCIENCES, TEL AVIV UNIVERSITY  
 RAMAT AVIV, TEL AVIV, 69978 ISRAEL  
 E-MAIL: [ronnyluss@gmail.com](mailto:ronnyluss@gmail.com)  
[saharon@post.tau.ac.il](mailto:saharon@post.tau.ac.il)

FACULTY OF ENGINEERING, TEL AVIV UNIVERSITY  
 RAMAT AVIV, TEL AVIV, 69978 ISRAEL  
 E-MAIL: [moni@eng.tau.ac.il](mailto:moni@eng.tau.ac.il)