
Fast Calculation of Feature Contributions in Boosting Trees

Zhongli Jiang¹

Min Zhang¹

Dabao Zhang¹

¹Department of Epidemiology & Biostatistics, University of California, Irvine, CA, USA

Abstract

Recently, several fast algorithms have been proposed to decompose predicted value into Shapley values, enabling individualized feature contribution analysis in tree models. While such local decomposition offers valuable insights, it underscores the need for a global evaluation of feature contributions. Although coefficients of determination (R^2) allow for comparative assessment of individual features, individualizing R^2 is challenged by the underlying quadratic losses. To address this, we propose Q-SHAP, an efficient algorithm that reduces the computational complexity of calculating Shapley values for quadratic losses to polynomial time. Our simulations show that Q-SHAP not only improves computational efficiency but also enhances the accuracy of feature-specific R^2 estimates.

1 INTRODUCTION

Models built with tree ensembles are powerful but often complicated, making it challenging to understand the influence of inputs. Feature importance plays a critical role in demystifying these models and enhancing their interpretability by assigning each input feature a score. This is crucial in domains like healthcare and biomedicine, where trust and interpretation of the model are essential [Stiglic et al., 2020, Bussmann et al., 2021]. Common feature importance measures like gain can be inconsistent [Lundberg and Lee, 2017b] while permutation importance lacks theoretical foundations [Ishwaran, 2007].

Shapley values, derived from cooperative game theory and introduced by Shapley [1953], offer a robust method for the fair distribution of payoffs generated by a coalition of players. This can be analogously applied to assess the contribution of each feature in a machine learning model. It ensures that each feature’s contribution is assessed by con-

sidering all possible combinations of features, thereby providing a comprehensive understanding of feature impacts. Recent applications of Shapley values have focused on local interpretation [Lundberg and Lee, 2017a,b, Chau et al., 2022], where they are employed to examine the influence of individual features on specific predictions. Nonetheless, there are numerous scenarios where global importance is preferred, such as analysis of the role of a feature across the entire dataset [Molnar, 2020, Covert et al., 2020].

Among the works that compute Shapley values in a global context, a popular approach is to use model variance decomposition. Lipovetsky and Conklin [2001] decomposed R^2 in linear regression, offering consistent interpretations even in the presence of multicollinearity. Owen and Prieur [2017] also conducted a conceptual analysis of Shapley values for the variance. However, computation remains a significant challenge, as the calculation of Shapley values grows exponentially with the number of features. To address this issue, several Monte Carlo-based methods have been proposed to effectively reduce the computational burden [Song et al., 2016, Covert et al., 2020, Williamson and Feng, 2020].

Although Monte Carlo-based, model-agnostic methods are more efficient than brute-force approaches, they are still computationally intensive, especially when dealing with high-dimensional data that requires extensive feature permutation sampling to ensure consistency [Lundberg and Lee, 2017a, Lundberg et al., 2020]. This challenge has prompted the development of methods that leverage the specific structures of tree-based models. However, much of the focus has been on explaining individual predictions, as seen with TreeSHAP [Lundberg and Lee, 2017b], FastTreeSHAP [Yang, 2021], LinearTreeSHAP [Bifet et al., 2022] and FourierSHAP [Gorji et al., 2024]. Bénard et al. [2022] considered population-level importance using R^2 , specifically tailored for random forests [Breiman, 2001].

Lundberg et al. [2020] suggests that explaining the loss function for a “path-dependent” algorithm is challenging. To the best of our knowledge, there is no available method to calcu-

late Shapley values of quadratic losses by leveraging structures of decision trees for fast computation. In this paper, we propose Q-SHAP, which can decompose quadratic terms of predicted values of a decision tree into each feature's attribute in polynomial time. It leads to fast computation of feature-specific R^2 for a decision tree. We also extend our approach to Gradient Boosted Decision Trees.

The rest of the paper is structured as follows. In Section 2, we provide a brief overview of Shapley values of R^2 . In Section 3, we present our proposed algorithm Q-SHAP to calculate Shapley values of R^2 in polynomial time for single trees, and then extend the approach for tree ensembles in Section 4. We justify the efficacy and efficiency of the algorithm using extensive simulations in Section 5 and real data analysis in high dimension in Section 6. We conclude the paper with a discussion in Section 7.

2 SHAPLEY VALUES OF R^2 FOR INDIVIDUAL FEATURES

2.1 MODEL SPECIFICATION

Here we investigate a specific label Y and its explainability by a full set of p features $X = (X_1, X_2, \dots, X_p)$. For any subset $F \subseteq \mathcal{P} = \{1, 2, \dots, p\}$, we define the corresponding set of features as $X_F = (X_j)_{j \in F}$.

Suppose that, for any set of features X_F , an oracle model m_F can be built such that, for any value $x = (x_j)_{j \in F}$,

$$m_F(x) = E[Y|X_F = (x_j)_{j \in F}].$$

The Shapley value of j -th feature, in terms of its contribution to the total variation, is defined as

$$\phi_{\rho^2, j} = \frac{1}{p \text{var}(m_\emptyset)} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} \times (\text{var}(m_{F \cup \{j\}}) - \text{var}(m_F)), \quad (1)$$

where $|F|$ is the number of features in F . The term $\text{var}(m_{F \cup \{j\}})$ is the variance explained by feature set $F \cup \{j\}$ and the term $\text{var}(m_F)$ is the variance explained solely by set F . This definition is analogous to Covert et al. [2020] and Williamson and Feng [2020]. By averaging over all possible feature combinations, the Shapley values are the only solution that satisfies the desired properties of symmetry, efficiency, additivity, and dummy [Shapley, 1953].

2.2 EMPIRICAL ESTIMATION

Suppose we have a set of data with sample size n observed for both label and features as

$$\begin{aligned} \mathbf{Y} &= (y_1, y_2, \dots, y_n), \\ \mathbf{X} &= (\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_p) = (\mathbf{x}_{1.}^T, \mathbf{x}_{2.}^T, \dots, \mathbf{x}_{n.}^T)^T. \end{aligned}$$

Accordingly, we denote the observed data of features in subset F as

$$\mathbf{X}_{\cdot F} = (\mathbf{X}_{\cdot j})_{j \in F}.$$

Suppose that, for each subset F of features, a single optimal model $\hat{m}_F$ is built on data $(\mathbf{Y}, \mathbf{X}_{\cdot F})$. Then the i -th label can be predicted with

$$\hat{y}_i(\mathbf{X}_{\cdot F}) = \hat{m}_F(\mathbf{x}_{i.}).$$

2.3 FROM R^2 TO A QUADRATIC LOSS

We will establish the connection of R^2 to a quadratic loss through equation (1). We define the quadratic loss on the optimal model $\hat{m}_F$ as

$$Q_F = \sum_{i=1}^n (y_i - \hat{m}_F(\mathbf{x}_{i.}))^2 \quad (2)$$

for any set of features F . With $m_\emptyset(\mathbf{x}_{i.}) = \bar{y}$, we have $Q_\emptyset = \sum_{i=1}^n (y_i - \bar{y})^2$. Following the law of total variance, we can estimate $\text{var}(m_F)$ by

$$\widehat{\text{var}}(m_F) = (Q_\emptyset - Q_F) / n.$$

Thus, an empirical estimate of (1) is

$$\phi_{R^2, j} = -\frac{1}{pQ_\emptyset} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} (Q_{F \cup \{j\}} - Q_F),$$

which is proportional to a Shapley value for the sum of squared errors, i.e., the quadratic loss in (2).

2.4 FROM QUADRATIC LOSS TO Q-SHAP

We now further reduce Shapley values of the sum of squared errors to Shapley values of linear and quadratic terms of predicted values. Expanding the loss function in (2), we can rewrite,

$$\begin{aligned} \phi_{R^2, j} &= -\frac{1}{pQ_\emptyset} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} \sum_{i=1}^n (\hat{m}_{F \cup j}^2(\mathbf{x}_{i.}) \\ &\quad - \hat{m}_F^2(\mathbf{x}_{i.}) - 2(\hat{m}_{F \cup j}(\mathbf{x}_{i.}) - \hat{m}_F(\mathbf{x}_{i.}))y_i). \end{aligned}$$

To calculate this, we define the Shapley value for each sample i as,

$$\begin{aligned} \phi_{R^2, j}(\mathbf{x}_{i.}) &= -\frac{1}{pQ_\emptyset} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} (\hat{m}_{F \cup j}^2(\mathbf{x}_{i.}) - \hat{m}_F^2(\mathbf{x}_{i.})) \\ &\quad + \frac{2y_i}{pQ_\emptyset} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} (\hat{m}_{F \cup j}(\mathbf{x}_{i.}) - \hat{m}_F(\mathbf{x}_{i.})), \end{aligned}$$

which is a linear combination of two sets of Shapley values, i.e., Shapley values of predicted value $\hat{m}_F$, which are ready to be calculated [Lundberg and Lee, 2017b, Yang, 2021, Bifet et al., 2022], and Shapley values of the quadratic term of predicted value $\hat{m}_F^2$, i.e.,

$$\phi_{\hat{m}^2, j}(\mathbf{x}_{i.}) = \frac{1}{p} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} \times (\hat{m}_{F \cup j}^2(\mathbf{x}_{i.}) - \hat{m}_F^2(\mathbf{x}_{i.})), \quad (3)$$

for which we will develop the algorithm Q-SHAP to calculate. For the rest of the paper, we will focus on computing the Shapley values in Equation (3) in polynomial time for tree-based models and carrying it over to calculate feature-specific R^2 .

3 THE ALGORITHM Q-SHAP FOR SINGLE TREES

Here we will focus on calculating the Shapley values in (3) for a single tree. By definition, each Shapley value (3) requires evaluating all feature subsets, leading to an NP-hard computation in general. However, binary trees restrict predictions to a finite set of values attached to leaf nodes, and we will show that summarizing differences of feature subsets are related to operations on polynomial coefficients. Leveraging our results of polynomial identity in Proposition 1 and Proposition 2, we can calculate this Shapley value over pairs of leave nodes, shown in Theorem 1, instead of all feature subsets.

3.1 NOTATIONS

We assume the underlying decision tree has the maximum depth at D and a total of L leaves, and use l to denote a specific leaf. We further introduce a dot product for polynomials for subsequent calculation. For two polynomials $A(z) = \sum_{i=0}^n a_i z^i$ and $B(z) = \sum_{i=0}^n b_i z^i$, we define their dot product as $A(z) \cdot B(z) = \sum_{i=0}^n a_i b_i$.

3.2 DISTRIBUTING THE PREDICTION TO LEAVES

Decision trees match each data point to one leaf for prediction. However, for our prediction defined on any subset F , a data point $\mathbf{x}_{i.}$ can fall into multiple leaves due to the uncertainty by unspecified features $\mathcal{P} \setminus F$. We can calculate $\hat{m}_F(\mathbf{x}_{i.})$, following TreeSHAP, as the empirical mean by aggregating the weighted prediction on each leaf,

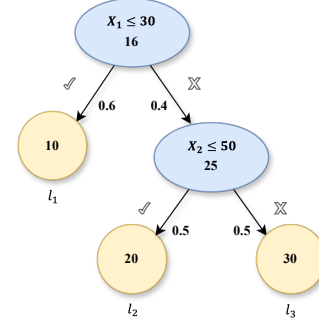
$$\hat{m}_F(\mathbf{x}_{i.}) = \sum_l \hat{m}_F^l(\mathbf{x}_{i.}) \triangleq \sum_l w(\mathbf{x}_{i.}, l, F) \times \hat{m}^l, \quad (4)$$

where $\hat{m}^l$ is the predicted value at leaf l based on the model built on all features.

Given an oracle tree built on all available features, we try to recover the oracle tree for a subset of features without rebuilding, following Bifet et al. [2022] and Karczmarz et al. [2022]. For example, suppose the tree involves two features, X_1 and X_2 , as shown in Figure 1.a. However, when feature X_2 is excluded from the structure, we would use only X_1 to build the tree as shown in Figure 1.b.

Therefore, we replace it with a pseudo internal node to preserve the structure of the original full oracle tree and pave the way for further formulation.

(a) Decision tree built on X_1 and X_2



(b) Hypothetical tree with X_1 only

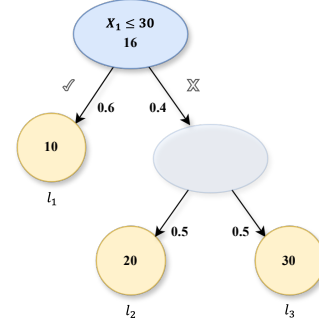


Figure 1: Illustration of (a) a decision tree built with both features X_1 and X_2 and (b) its hypothetical tree with only feature X_1 .

With a data point $\mathbf{x}_{i.} = (40, 25)$, we illustrate the calculation in Equation (4) by first calculating the predicted value for the tree in Figure 1.b, $\hat{m}_{\{1\}}(\mathbf{x}_{i.}) = \hat{m}_{\{1\}}^{l_1}(\mathbf{x}_{i.}) + \hat{m}_{\{1\}}^{l_2}(\mathbf{x}_{i.}) + \hat{m}_{\{1\}}^{l_3}(\mathbf{x}_{i.}) = 0 \times 10 + 0.5 \times 20 + 0.5 \times 30$.

When the tree is built with an additional feature $j = 2$ as shown in Figure 1.a, we have the predicted value $\hat{m}_{\{1,2\}}(\mathbf{x}_{i.}) = \hat{m}_{\{1,2\}}^{l_1}(\mathbf{x}_{i.}) + \hat{m}_{\{1,2\}}^{l_2}(\mathbf{x}_{i.}) + \hat{m}_{\{1,2\}}^{l_3}(\mathbf{x}_{i.}) = \mathbf{1} \times 0 \times 10 + \mathbf{0.5}^{-1} \times 0.5 \times 20 + \mathbf{0} \times 0.5 \times 30$

where the bold numbers reweight $\hat{m}_{\{1\}}^l(\mathbf{x}_{i.})$ for $\hat{m}_{\{1,2\}}^l(\mathbf{x}_{i.})$. Let us take a closer look at these weights, each corresponding to one leaf. For leaf l_1 , the weight is 1 since the newly added feature X_2 is not involved in its path and the reweighted prediction remains as zero. For leaf l_2 , the reweighted prediction is lifted up with the weight inversely

proportional to the previous probability because $\mathbf{x}_i$ follows its path to the leaf with probability 1. On the other hand, although the path to leaf l_3 includes the newly added feature, $\mathbf{x}_i$ doesn't follow this path, resulting in a weight at 0. Next, we will generalize such a reweighting strategy to calculate (4) for trees with different sets of features.

Denote F^l the features involved in the path to leaf l and $F^l(\mathbf{x}_i)$ the subset of F^l whose decision criteria are satisfied by $\mathbf{x}_i$. Note that each feature $j \in F^l$ may appear multiple times in the path to leaf l so we denote $n_{j,c}^l$ the number of samples passing through the node which is attached to the c -th appearance. We similarly define $n_{j,c}^l(\mathbf{x}_i)$ for each feature $j \in F^l(\mathbf{x}_i)$.

For any feature $j \in \mathcal{P}$, we can define the weight function based on a partition of $\mathcal{P}$ into three subsets $F^l(\mathbf{x}_i)$, $F^l \setminus F^l(\mathbf{x}_i)$, and $\mathcal{P} \setminus F^l$,

$$w_j^l(\mathbf{x}_i) \triangleq \begin{cases} \prod_c \frac{n_{j,c}^l}{n_{j,c}^l(\mathbf{x}_i)}, & \text{if } j \in F^l(\mathbf{x}_i); \\ 0, & \text{if } j \in F^l \setminus F^l(\mathbf{x}_i); \\ 1, & \text{if } j \in \mathcal{P} \setminus F^l. \end{cases}$$

Therefore, for $j \notin F$, we have

$$\hat{m}_{F \cup j}^l(\mathbf{x}_i) = w_j^l(\mathbf{x}_i) \hat{m}_F^l(\mathbf{x}_i).$$

Recursive application of the above formula leads to

$$\hat{m}_F^l(\mathbf{x}_i) = \prod_{k \in F} w_k^l(\mathbf{x}_i) \hat{m}_\emptyset^l,$$

where $\hat{m}_\emptyset^l = \hat{m}^l \frac{n^l}{n}$ with n^l the sample size at leaf l , n the total sample size.

When $F = \emptyset$, the above result reduces to $\hat{m}_\emptyset(\mathbf{x}_i) = \sum_l \hat{m}^l \frac{n^l}{n}$, so the optimal prediction is just the mean for all data points, which is consistent with $\hat{m}_\emptyset(\mathbf{x}_i) = \bar{y}$.

We can rewrite (3) by aggregating over leaves as

$$\begin{aligned} \phi_{\hat{m}^2,j}(\mathbf{x}_i) &= \frac{1}{p} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} \\ &\quad \times \left(\sum_l (w_j^{l_2}(\mathbf{x}_i) - 1) \hat{m}_\emptyset^{l_2} \prod_{k \in F} w_k^{l_2}(\mathbf{x}_i) \right) \\ &\quad + \frac{2}{p} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} \\ &\quad \times \left(\sum_{l_1 \neq l_2} (w_j^{l_1}(\mathbf{x}_i) w_j^{l_2}(\mathbf{x}_i) - 1) \hat{m}_\emptyset^{l_1} \hat{m}_\emptyset^{l_2} \right. \\ &\quad \times \left. \prod_{k \in F} w_k^{l_1}(\mathbf{x}_i) w_k^{l_2}(\mathbf{x}_i) \right) \\ &\triangleq T_{1,j}(\mathbf{x}_i) + 2T_{2,j}(\mathbf{x}_i). \end{aligned} \quad (5)$$

We further define, for leaves l_1 and l_2 ,

$$\begin{aligned} T_j^{l_1 l_2}(\mathbf{x}_i) &= \frac{1}{p} \sum_{F \subseteq \mathcal{P} \setminus \{j\}} \binom{p-1}{|F|}^{-1} \left((w_j^{l_1}(\mathbf{x}_i) w_j^{l_2}(\mathbf{x}_i) - 1) \right. \\ &\quad \times \left. \hat{m}_\emptyset^{l_1} \hat{m}_\emptyset^{l_2} \prod_{k \in F} w_k^{l_1}(\mathbf{x}_i) w_k^{l_2}(\mathbf{x}_i) \right), \end{aligned} \quad (6)$$

and we have $T_{2,j}(\mathbf{x}_i) = \sum_{l_1 \neq l_2} T_j^{l_1 l_2}(\mathbf{x}_i)$, $T_{1,j}(\mathbf{x}_i) = \sum_l T_j^{ll}(\mathbf{x}_i)$. Therefore, we will focus on the calculation of $T_j^{l_1 l_2}(\mathbf{x}_i)$ in (6) the rest of this section.

We can reduce the calculation of $\prod_{k \in F} w_k^{l_1}(\mathbf{x}_i) w_k^{l_2}(\mathbf{x}_i)$ in (6) by only calculating $\prod_{k \in F_-} w_k^{l_1}(\mathbf{x}_i) w_k^{l_2}(\mathbf{x}_i)$ with $F_- = F \cap (F^{l_1} \cup F^{l_2})$, because $\prod_{k \in F \setminus F_-} w_k^{l_1}(\mathbf{x}_i) w_k^{l_2}(\mathbf{x}_i) = 1$. In combination with the proposition below, computation in (6) can be dramatically reduced from the full feature set $\mathcal{P}$ to a set only related to the corresponding leaves in a tree.

Proposition 1 For any well-defined $p, n, |F|$,

$$\sum_{k=0}^{p-n} \frac{\binom{p-n}{k}}{p \binom{p-1}{|F|+k}} = \frac{1}{n \binom{n-1}{|F|}}.$$

We leave the proof of Proposition 1 in Appendix A. Further denote $n_{12} = |F^{l_1} \cup F^{l_2}|$ and a polynomial of z , $P^{l_1 l_2}(z) = \prod_{k \in F^{l_1} \cup F^{l_2} \setminus j} (z + w_k^{l_1}(\mathbf{x}_i) w_k^{l_2}(\mathbf{x}_i))$. We then define a coefficient polynomial $C_{n_{12}}(z) = \frac{1}{\binom{n_{12}-1}{0}} z^0 + \frac{1}{\binom{n_{12}-1}{1}} z^1 + \dots + \frac{1}{\binom{n_{12}-1}{n_{12}-1}} z^{n_{12}-1}$.

Theorem 1 The Shapley value in (3) can be calculated as,

$$\phi_{\hat{m}^2,j}(\mathbf{x}_i) = \sum_l T_j^{ll}(\mathbf{x}_i) + 2 \sum_{l_1 \neq l_2} T_j^{l_1 l_2}(\mathbf{x}_i), \quad (7)$$

with

$$\begin{aligned} T_j^{l_1 l_2}(\mathbf{x}_i) &= \frac{1}{n_{12}} (w_j^{l_1}(\mathbf{x}_i) w_j^{l_2}(\mathbf{x}_i) - 1) \hat{m}_\emptyset^{l_1} \hat{m}_\emptyset^{l_2} \\ &\quad \times [C_{n_{12}}(z) \cdot P^{l_1 l_2}(z)]. \end{aligned} \quad (8)$$

Proof. With Proposition 1, we can write (6) as

$$\begin{aligned} T_j^{l_1 l_2}(\mathbf{x}_i) &= \frac{1}{n_{12}} (w_j^{l_1}(\mathbf{x}_i) w_j^{l_2}(\mathbf{x}_i) - 1) \hat{m}_\emptyset^{l_1} \hat{m}_\emptyset^{l_2} \\ &\quad \times \sum_{t=0}^{n_{12}-1} \frac{1}{\binom{n_{12}-1}{t}} \sum_{F \subseteq F^{l_1} \cup F^{l_2} \setminus j}^{|F|=t} \prod_{k \in F} w_k^{l_1}(\mathbf{x}_i) w_k^{l_2}(\mathbf{x}_i). \end{aligned}$$

We notice that $\sum_{F \subseteq F^{l_1} \cup F^{l_2} \setminus j}^{|F|=t} \prod_{k \in F} w_k^{l_1}(\mathbf{x}_i) w_k^{l_2}(\mathbf{x}_i)$ is the coefficient of z^t in polynomial $P^{l_1 l_2}(z)$, hence the equation (8) holds with $C_{n_{12}}(z)$ adjusting the weight based on the size of set F . The calculation in (7) follows (5). ■

We only need to consider feature $j \in |F^{l_1} \cup F^{l_2}|$ as, otherwise, we have $T_j^{l_1 l_2}(\mathbf{x}_{i.}) = 0$ following the definition of $w_j^{l_1}(\mathbf{x}_{i.})$. Note that, when there is a feature in set F that doesn't belong to $F^{l_1}(\mathbf{x}_{i.}) \cap F^{l_2}(\mathbf{x}_{i.}) \setminus j$, we have $\prod_{k \in F} w_k^{l_1}(\mathbf{x}_{i.}) w_k^{l_2}(\mathbf{x}_{i.}) = 0$. Thus we can further simplify the term to

$$\begin{aligned} & T_j^{l_1 l_2}(\mathbf{x}_{i.}) \\ &= \frac{1}{n_{12}} (w_j^{l_1}(\mathbf{x}_{i.}) w_j^{l_2}(\mathbf{x}_{i.}) - 1) \hat{m}_\emptyset^{l_1} \hat{m}_\emptyset^{l_2} \sum_{t=0}^{n_{12}-1} \frac{1}{\binom{n_{12}-1}{t}} \\ & \quad \times \sum_{|F|=t} \prod_{F \subseteq F^{l_1}(\mathbf{x}_{i.}) \cap F^{l_2}(\mathbf{x}_{i.}) \setminus j} w_k^{l_1}(\mathbf{x}_{i.}) w_k^{l_2}(\mathbf{x}_{i.}). \end{aligned}$$

Consequently, the evaluation of $P^{l_1 l_2}(z)$ can be reduced to a much smaller set.

3.3 THE ALGORITHM

In this section, we will introduce a fast and stable evaluation for the dot product of a coefficient polynomial $C(z)$ where we know the coefficients and a polynomial $P(z)$ with a known product form, involved in Theorem 1.

Proposition 2 *Let ω be a vector of the complex n -th roots of unity whose element is $\exp(\frac{2k\pi i}{n})$ for $k = 0, 1, \dots, n-1$, c the coefficient vector of $C(z)$, and $IFFT$ the Inverse Fast Fourier Transformation. Then*

$$C(z) \cdot P(z) = P(\omega)^T IFFT(c).$$

The proof of Proposition 2 is shown in Appendix A. We facilitate the computation via the complex roots of unity because of their numerical stability and fast operations in matrix multiplications. Due to the potential issue of ill condition, especially at large degrees, our calculation avoids inversion of the Vandermonde matrices, although it has been proposed to facilitate the computing by Bifet et al. [2022]. In addition, for each sample size n , we only need to calculate $IFFT(c)$ once, up to order D in $O(n \log(n))$ operations, and the results can be saved for the rest of calculation through Q-SHAP. Note that term k and term $n - k$ in $P(w)$ are complex conjugates, and, for a real vector c , $IFFT(c)$ also has the conjugate property for paired term k and term $n - k$. Consequently, the dot product of $P(\omega)$ and $IFFT(c)$ inherits the conjugate property and its imaginary parts are canceled upon addition. Therefore, we only need evaluate the dot product at half of the n complex roots.

We can aggregate the values of leaf combinations to derive the Shapley values of squared predictions using Q-SHAP as in Algorithm 1 and then calculate the Shapley values of R^2 using RSQ-SHAP as in Algorithm 2. The calculation of feature-specific R^2 uses the iterative Algorithm 1 instead

of a recursive one. As detailed in Appendix D, the time complexity of the algorithm is $O(L^2 D^2)$ for a single tree, which doesn't depend on the dimension p and is extremely fast when the maximum tree depth is not too large.

Algorithm 1 Q-SHAP

```

Q-SHAP( $\mathbf{x}_{i.}$ )
Initialize  $T[j] = 0$  for  $j = 1, \dots, p$ 
for  $l_1 \in$  index set  $0, \dots, L-1$  do
  for  $l_2 \in$  index set  $l_1, \dots, L-1$  do
    Let  $n_{12} = |F^{l_1} \cup F^{l_2}|$ 
    for  $j \in F^{l_1} \cup F^{l_2}$  do
      Let  $t[j] = \frac{1}{n_{12}} [w_j^{l_1}(\mathbf{x}_{i.}) w_j^{l_2}(\mathbf{x}_{i.}) - 1] \times$ 
         $\hat{m}_\emptyset^{l_1} \hat{m}_\emptyset^{l_2} [C_{n_{12}}(z) \cdot P^{l_1 l_2}(z)]$ 
      if  $l_1 \neq l_2$  then
         $T[j] = T[j] + 2t[j]$ 
      else
         $T[j] = T[j] + t[j]$ 
      end if
    end for
  end for
end for
return  $T = (T[1], T[2], \dots, T[p])$ 

```

Algorithm 2 RSQ-SHAP

$$\text{RSQ-SHAP}(j) = -\frac{1}{Q_\emptyset} \sum_{i=1}^n \{ \text{Q-SHAP}(\mathbf{x}_{i.})[j] - 2y_i \text{SHAP}(\mathbf{x}_{i.})[j] \}$$

4 THE ALGORITHM Q-SHAP FOR TREE ENSEMBLES FROM BOOSTING

Tree ensembles from Gradient Boosted Machines (GBM) [Friedman, 2001] greatly improve predictive performance by aggregating many weak learners [Chen and Guestrin, 2016, Ke et al., 2017, Prokhorenkova et al., 2018]. Each tree, say tree k , is constructed on the residuals from the previous tree, i.e., tree $k-1$. We assume that there are a total of K trees in the ensemble and the quadratic loss by the first k trees, with all features in $\mathcal{P}$, is $Q_{\mathcal{P}}^{(k)}$. Denoting $Q_{\mathcal{P}}^{(0)} = Q_\emptyset$, the k -th tree reduces the loss by

$$\Delta Q_{\mathcal{P}}^{(k)} = Q_{\mathcal{P}}^{(k-1)} - Q_{\mathcal{P}}^{(k)}, \quad (9)$$

with the tree ensemble reducing the total loss by

$$Q_\emptyset - Q_{\mathcal{P}}^{(K)} = \sum_{k=1}^K \Delta Q_{\mathcal{P}}^{(k)}.$$

Per our interest in feature-specific R^2 , we resort to the quadratic loss defined as the sum of squared errors in (2).

On the other hand, the k -th tree provides the prediction $\hat{m}_{\mathcal{P}}^{(k)}(\mathbf{x}_{i.})$. Therefore, the prediction by the first k trees

can be recursively calculated as $\hat{y}_i^{(k)}(\mathbf{X}) = \hat{y}_i^{(k-1)}(\mathbf{X}) + \alpha \hat{m}_{\mathcal{P}}^{(k)}(\mathbf{x}_i)$, where α is the learning rate and $\hat{y}_i^{(0)}(\mathbf{X}) \equiv \bar{y}$. Note that the residuals after building $(k-1)$ tree are $\{r_i^{(k-1)} = y_i - \hat{y}_i^{(k-1)}(\mathbf{X}) : i = 1, 2, \dots, n\}$, which are taken to build the k -th tree. Thus,

$$\begin{aligned} \Delta Q_{\mathcal{P}}^{(k)} &= \sum_{i=1}^n (r_i^{(k-1)})^2 - \sum_{i=1}^n (r_i^{(k-1)} - \alpha \hat{m}_{\mathcal{P}}^{(k)}(\mathbf{x}_i))^2 \\ &= - \sum_{i=1}^n (\alpha^2 \hat{m}_{\mathcal{P}}^{(k)2}(\mathbf{x}_i) - 2\alpha r_i^{(k-1)} \hat{m}_{\mathcal{P}}^{(k)}(\mathbf{x}_i)). \end{aligned}$$

Thus, decomposition of $\Delta Q_{\mathcal{P}}^{(k)}$ in (9) for feature-specific Shapley values can be conducted via the decomposition of two sets of values, i.e., SHAP on the predicted value $\hat{m}_{\mathcal{P}}^{(k)}(\mathbf{x}_i)$ and Q-SHAP on its quadratic term $\hat{m}_{\mathcal{P}}^{(k)2}(\mathbf{x}_i)$. Summing up these Shapley values over all trees leads to Shapley values for the tree ensemble.

5 SIMULATION STUDY

One of the challenges in assessing methods that explain predictions is the typical absence of a definitive ground truth. Therefore, to fairly demonstrate the fidelity of our methodology, we must rely on synthetic data that allows for the calculation of the theoretical Shapley values. Here we consider three different models,

- a. $Y = 4X_1 - 5X_2 + 6X_3 + \epsilon$;
- b. $Y = 4X_1 - 5X_2 + 6X_3 + 3X_1X_2 - X_1X_3 + \epsilon$;
- c. $Y = 4X_1 - 5X_2 + 6X_3 + 3X_1X_2 - X_1X_2X_3 + \epsilon$.

All three features involved in the models are generated from Bernoulli distributions with probabilities of 0.6, 0.7, and 0.5, respectively.

We also simulate additional nuisance features independently from $Bernoulli(0.5)$ to make the total number of features $p = 100$ and $p = 500$, respectively. The error term ϵ is generated from $N(0, \sigma_\epsilon^2)$ with σ_ϵ at three different levels, i.e., 0.5, 1, and 1.5. The theoretical values of total R^2 and feature-specific R^2 are shown in Table 6 of Appendix B.

We evaluate the performance of three different methods, our proposed Q-SHAP, SAGE by Covert et al. [2020], and SPVIM by Williamson and Feng [2020], in calculating the feature-specific R^2 for the above three models with data sets of different sample sizes at $n = 500, 1000, 2000$, and 5000. We use package *sage-importance* for SAGE and package *vimpy* for SPVIM to calculate feature-specific Shapley values of total explained variance, which are divided by the total variance for corresponding feature-specific R^2 values.

For each setting, we generated 1,000 data sets. For each data set, we built a tree ensemble using XGBoost [Chen and Guestrin, 2016] with tuning parameters optimized via 5-fold

cross-validation and grid search in a parameter space specified with the learning rate in $\{0.01, 0.05, 0.1\}$ and number of estimators in $\{50, 100, 200, 300, \dots, 1000\}$. We fixed the maximum depth of models a, b, and c at 1, 2, and 3 respectively. Table 1 shows the bias in calculating feature-specific R^2 for the first three features as well as the sum of all feature-specific R^2 for all three models with $n = 1000$, $p = 100$, and $\sigma_\epsilon = 1.5$. The estimation results of the three models in other settings are plotted in Appendix C. Overall, Q-SHAP provides a more stable and accurate calculation of feature-specific R^2 than the other two methods.

We divide all features into two groups, signal features (the first three) and nuisance features (the rest). For each group, we calculated the mean absolute error (MAE) by comparing feature-specific R^2 values to the theoretical ones in each of the 1,000 datasets and averaged MAE over the 1,000 datasets, shown in Figure 2. Note that, by limiting memory to 2GB, SAGE can only report R^2 for the data sets with sample size at 500 and 1,000.

For both signal and nuisance features, Q-SHAP and SAGE exhibit consistent behavior across all models. In contrast, SPVIM tends to bias the calculation, especially for small sample sizes, indicated by the rapid increase of MMAE when the sample size goes down. Among signal features, Q-SHAP has better accuracy than SAGE, followed by SPVIM in general. All methods tend to have better accuracy when sample size increases.

For the nuisance features, only SPVIM is biased away from 0. On the other hand, both Q-SHAP and SAGE have almost no bias for nuisance features across different sample sizes. For all three methods, R^2 of signal features tends to have a larger bias than nuisance features.

We compared the computational time of the three different methods by running all algorithms in parallel on a full node consisting of two AMD CPUs@2.2GHz with 128 cores and 256 GB memory. We unified the environment with the help of a Singularity container [Kurtzer et al., 2017] built under Python version 3.11.6. Due to the large size of the simulation, we limit all methods to a maximum wall time of 4 hours per dataset on a single core, with memory limited to 2 GB. The running times are shown in Figure 3. Both SAGE and SPVIM demanded a long time to compute even with only 100 features. Q-SHAP is hundreds of times faster than both SAGE and SPVIM in general and is the only method that can be completed when the dimension is 500 in constrained computation time and memory.

6 REAL DATA ANALYSIS

We illustrate the utility of Q-SHAP by applying it to three datasets: (1) Healthcare data that includes eight features for each of the 1338 subjects besides their healthcare insurance

Table 1: Estimation bias (SE) of X_1 -specific, X_2 -specific, X_3 -specific, and the sum of all feature-specific R^2 values across the three models with $n = 1,000$, $p = 100$, and $\sigma_\epsilon = 1.5$. Bolded values indicate the smallest bias in magnitude.

	Method	X_1 -specific R^2	X_2 -specific R^2	X_3 -specific R^2	Sum of all R^2
Model a	Q-SHAP	0.006 (0.011)	0.008 (0.014)	0.015 (0.015)	0.031 (0.014)
	SAGE	-0.015 (0.015)	-0.017 (0.016)	-0.021 (0.016)	-0.053 (0.020)
	SPVIM	0.049 (0.035)	0.069 (0.039)	0.116 (0.051)	0.242 (0.208)
Model b	Q-SHAP	0.008 (0.017)	0.002 (0.009)	0.009 (0.016)	0.024 (0.019)
	SAGE	-0.026 (0.016)	-0.014 (0.011)	-0.024 (0.014)	-0.062 (0.019)
	SPVIM	0.111 (0.046)	0.049 (0.033)	0.098 (0.048)	0.256 (0.219)
Model c	Q-SHAP	0.005 (0.017)	0.002 (0.009)	0.005 (0.015)	0.021 (0.016)
	SAGE	-0.025 (0.016)	-0.014 (0.012)	-0.024 (0.014)	-0.062 (0.020)
	SPVIM	0.107 (0.047)	0.048 (0.032)	0.098 (0.046)	0.260 (0.206)

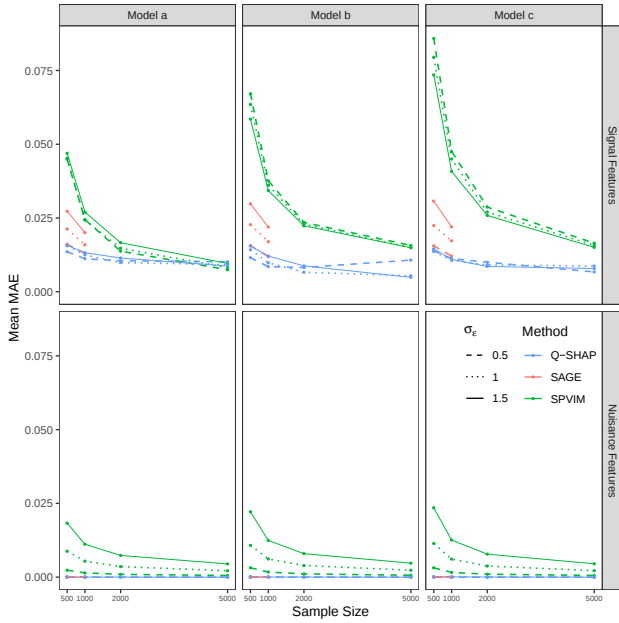


Figure 2: The mean absolute error (MAE) averaged over 1,000 datasets with $p = 100$

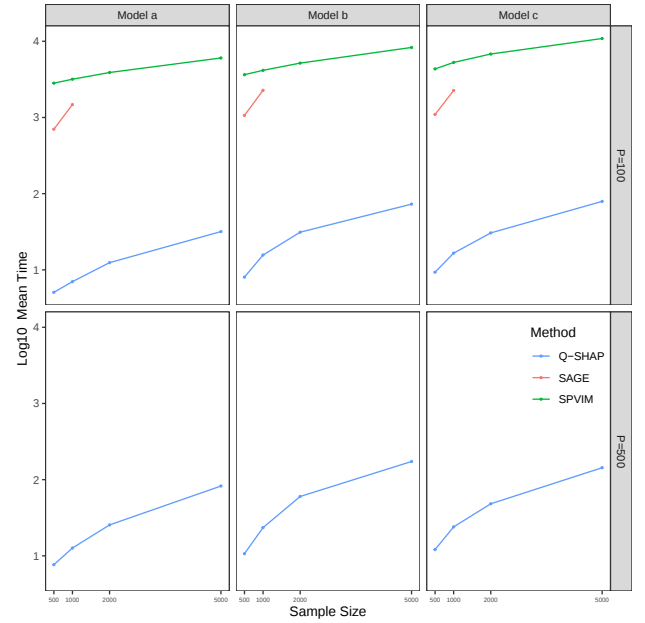


Figure 3: The running time (seconds) in simulation study

costs¹; (2) Prostate data that includes expression levels of 17,261 genes for 551 samples, available from UCSC Xena [Goldman et al., 2020], and cancer-indicating Gleason score, available from TCGAbiolinks [Colaprico et al., 2016]; and (3) S&P 500 data that includes prices of NVIDIA and other 469 stocks for 1,258 business days from February 8, 2013 to February 7, 2018². Here we predicted the daily return rate of NVIDIA stock from those of other 469 stocks in S&P 500 data, and the Gleason score, adjusted for age and race, using the gene expression features in Prostate data.

¹<https://www.kaggle.com/datasets/arunjangir245/healthcare-insurance-expenses/data>

²<https://www.kaggle.com/datasets/camnugent/sandp500/data>

For each data, we constructed the tree ensemble using XGBoost with tuning parameters optimized via 5-fold cross-validation and random search in a parameter space specified with the number of trees in $\{50, 100, 500, 1000, 1500, 2000, 2500, 3000\}$, maximum depth in $\{1, 2, \dots, 6\}$, and learning rate in $\{0.01, 0.05, 0.1\}$. For Prostate and S&P 500 data, as both SAGE and SPVIM cannot manage the large numbers of features, we first applied Q-SHAP to the tree ensembles built on all features and then selected the top 100 features to reconstruct the tree ensembles, whose feature-specific R^2 were calculated using all three methods. As shown in Table 2, although feasible for limited number of features, both SAGE and, in particular, SPVIM took much longer time to calculate the feature-specific R^2 values.

Table 2: The running time in real data analysis

	Q-SHAP	SAGE	SPVIM
Healthcare	32 s	6 m	22 m
Prostate	128 s	43 m	67 h
S&P 500	41 s	26 h	138 h

For the healthcare data, the tree ensemble in predicting healthcare insurance expenses reports the total R^2 at 0.86. As shown in Table 3, both Q-SHAP and SAGE ranked the eight features similarly, although SAGE tends to report slightly smaller values. On the other hand, SPVIM differs from the other two methods with its ranking list, demonstrated by the highlighted features in Table 3. In fact, SPVIM reported some much larger feature-specific R^2 values with some much smaller ones which were even negative.

With 100 features in the rebuilt tree ensembles for the prostate and S&P 500 data, we observe inconsistencies between Q-SHAP and SAGE, while the feature ranking from Q-SHAP matched that of models built on full feature set, see Table 4 and Table 5. The rebuilt tree ensembles reported total R^2 of 0.995 for the prostate data and 0.733 for the S&P 500 data. Notably, for the prostate data, the sum of all 100 feature-specific R^2 values from Q-SHAP matches the total R^2 value, whereas SAGE yields a slightly lower sum of 0.938. In contrast, although SPVIM tends to overstate some feature-specific R^2 values, the total sum of its feature-specific R^2 is only -0.22, substantially lower than the total value at 0.995. Overall, the real data analysis aligns with the simulation study, confirming that SAGE tends to underestimate the feature-specific R^2 , SPVIM shows instability, and Q-SHAP outperforms both in computational efficiency and the accuracy of feature-specific R^2 .

Table 3: Feature-specific R^2 in the healthcare data. Highlighted are features with R^2 larger than the preceding one.

Feature	Q-SHAP	SAGE	SPVIM
smoker_yes	0.481	0.455	0.565
age	0.323	0.293	0.327
children	0.026	0.021	0.012
bmi	0.021	0.019	0.051
sex_male	0.003	0.000	-0.033
region_southwest	0.002	0.001	0.044
region_southeast	0.001	0.000	0.048
region_northwest	0.000	0.000	-0.020

7 CONCLUSION

The coefficient of determination, aka R^2 , measures the proportion of the total variation explained by available features. Its additive decomposition, following Shapley [1953], pro-

Table 4: Feature specific R^2 in the prostate data. Highlighted are features with R^2 larger than the preceding one.

Gene	Orginal	Model with Selected Feaures		
	Q-SHAP	Q-SHAP	SAGE	SPVIM
SLC7A4	0.105	0.112	0.063	0.031
FAM72A	0.051	0.063	0.045	-0.025
CENPA	0.039	0.039	0.024	-0.175
KIAA0319L	0.034	0.036	0.029	0.033
CBX2	0.028	0.034	0.025	0.063
COL5A2	0.018	0.022	0.022	-0.033
SPATA4	0.018	0.019	0.021	-0.004
DOCK6	0.016	0.017	0.010	-0.150
KIF18B	0.015	0.023	0.012	0.185
TSEN15	0.014	0.015	0.021	-0.004
OTHERS	0.572	0.614	0.667	-0.140

Table 5: Feature-specific R^2 in the S&P 500 data. Highlighted are features with R^2 larger than the preceding one.

Stock	Orginal	Model with Selected Feaures		
	Q-SHAP	Q-SHAP	SAGE	SPVIM
AMAT	0.100	0.111	0.075	-0.141
AMD	0.077	0.085	0.071	0.102
MCHP	0.072	0.069	0.034	0.158
ADI	0.068	0.072	0.048	-0.026
TXN	0.041	0.043	0.029	-0.142
MU	0.035	0.041	0.039	-0.232
ADM	0.022	0.026	0.013	-0.004
AGN	0.021	0.023	0.018	0.120
NEM	0.016	0.015	0.010	0.202
PBCT	0.015	0.013	0.011	-0.137
Others	0.157	0.234	0.153	-0.565

vides an ideal evaluation of each feature’s attribute to explain the total variation. Shapley values are defined by differences involving all feature subsets, a seemingly NP-hard problem in general. Recently, several methods [Lundberg and Lee, 2017b, Yang, 2021, Bifet et al., 2022] have been developed to leverage the structure of tree-based models and provide computationally efficient algorithms to decompose the predicted values. However, decomposing R^2 demands the decomposition of a quadratic loss reduction by multiple trees. We have shown in Section 4 that we can attribute the total loss reduction by the tree ensemble to each single tree, and the tree-specific loss reductions are subject to further decomposition to each feature. However, decomposing a quadratic loss of a single tree needs work with the squared terms of predicted values, invalidating previously developed methods for predicted values. Leveraging structural property of trees and theoretical results of polynomials, we developed the Q-SHAP algorithm to consolidate calculations cross models and calculate Shapley values of squared pre-

dicted valued in polynomial time. The algorithm works not only for R^2 but also for general quadratic losses. Ultimately, it may provide a framework for more general loss functions via approximation.

Acknowledgements

This research was partially supported by NIH grants R01GM131491, R01AG080917, and R01AG080917-02S1, NCI grants R03 CA235363 and P30CA062203, and UCI Anti-Cancer Challenge funds from the UC Irvine Comprehensive Cancer Center. The results shown in Section 6 are in part based upon data generated by the TCGA Research Network: <https://www.cancer.gov/tcga>. The content is solely the responsibility of the authors and does not necessarily represent the official views of the National Institutes of Health or the Chao Family Comprehensive Cancer Center.

References

- Clément Bénard, Gérard Biau, Sébastien Da Veiga, and Erwan Scornet. SHAFF: Fast and consistent SHAPley eFFect estimates via random Forests. In *International Conference on Artificial Intelligence and Statistics*, pages 5563–5582. PMLR, 2022.
- Albert Bifet, Jesse Read, Chao Xu, et al. Linear TreeShap. *Advances in Neural Information Processing Systems*, 35: 25818–25828, 2022.
- Leo Breiman. Random Forests. *Machine Learning*, 45: 5–32, 2001.
- Niklas Bussmann, Paolo Giudici, Dimitri Marinelli, and Jochen Papenbrock. Explainable machine learning in credit risk management. *Computational Economics*, 57 (1):203–216, 2021.
- Siu Lun Chau, Robert Hu, Javier Gonzalez, and Dino Sejdinovic. RKHS-SHAP: Shapley values for kernel methods. *Advances in Neural Information Processing Systems*, 35: 13050–13063, 2022.
- Tianqi Chen and Carlos Guestrin. XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 785–794, 2016.
- Antonio Colaprico, Tiago C Silva, Catharina Olsen, Luciano Garofano, Claudia Cava, Davide Garolini, Thais S Sabedot, Tathiane M Malta, Stefano M Pagnotta, Isabella Castiglioni, et al. TCGAAbiolinks: an R/Bioconductor package for integrative analysis of TCGA data. *Nucleic Acids Research*, 44(8):e71, 2016.
- Ian Covert, Scott M Lundberg, and Su-In Lee. Understanding global feature contributions with additive importance measures. *Advances in Neural Information Processing Systems*, 33:17212–17223, 2020.
- Jerome H Friedman. Greedy function approximation: A gradient boosting machine. *Annals of Statistics*, pages 1189–1232, 2001.
- Keith O Geddes, Stephen R Czapor, and George Labahn. *Algorithms for Computer Algebra*. Springer Science & Business Media, 1992.
- Mary J Goldman, Brian Craft, Mim Hastie, Kristupas Repečka, Fran McDade, Akhil Kamath, Ayan Banerjee, Yunhai Luo, Dave Rogers, Angela N Brooks, et al. Visualizing and interpreting cancer genomics data via the Xena platform. *Nature Biotechnology*, 38(6):675–678, 2020.
- Ali Gorji, Andisheh Amrollahi, and Andreas Krause. Amortized shap values via sparse fourier function approximation. *arXiv preprint arXiv:2410.06300*, 2024.
- Henry W Gould. *Combinatorial Identities*. Morgantown Printing and Binding Co, 1972.
- Hemant Ishwaran. Variable importance in binary regression trees and forests. *Electronic Journal of Statistics*, 1:519–537, 2007.
- Adam Karczmarz, Tomasz Michalak, Anish Mukherjee, Piotr Sankowski, and Piotr Wygocki. Improved Feature Importance Computation for Tree Models based on the Banzhaf Value. In *Proceedings of the Thirty-Eight Conference on Uncertainty in Artificial Intelligence*, pages 969–979, 2022.
- Guolin Ke, Qi Meng, Thomas Finley, Taifeng Wang, Wei Chen, Weidong Ma, Qiwei Ye, and Tie-Yan Liu. Lightgbm: A highly efficient gradient boosting decision tree. *Advances in Neural Information Processing Systems*, 30, 2017.
- Gregory M Kurtzer, Vanessa Sochat, and Michael W Bauer. Singularity: Scientific containers for mobility of compute. *PLOS ONE*, 12(5):e0177459, 2017.
- Stan Lipovetsky and Michael Conklin. Analysis of regression in game theory approach. *Applied Stochastic Models in Business and Industry*, 17(4):319–330, 2001.
- Scott M Lundberg and Su-In Lee. A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30, 2017a.
- Scott M Lundberg and Su-In Lee. Consistent feature attribution for tree ensembles. *Workshop on Human Interpretability in Machine Learning*, 2017b.

Scott M Lundberg, Gabriel Erion, Hugh Chen, Alex De-Grave, Jordan M Prutkin, Bala Nair, Ronit Katz, Jonathan Himmelfarb, Nisha Bansal, and Su-In Lee. From local explanations to global understanding with explainable ai for trees. *Nature Machine Intelligence*, 2(1):56–67, 2020.

Christoph Molnar. *Interpretable Machine Learning*. lulu.com, 2020.

Art B Owen and Clémentine Prieur. On Shapley value for measuring importance of dependent inputs. *SIAM/ASA Journal on Uncertainty Quantification*, 5(1):986–1002, 2017.

Liudmila Prokhorenkova, Gleb Gusev, Aleksandr Vorobev, Anna Veronika Dorogush, and Andrey Gulin. Catboost: unbiased boosting with categorical features. *Advances in Neural Information Processing Systems*, 31, 2018.

Lloyd S Shapley. A value for N-person games. *Contributions to the Theory of Games*, 2(28):307–317, 1953.

Eunhye Song, Barry L Nelson, and Jeremy Staum. Shapley effects for global sensitivity analysis: theory and computation. *SIAM/ASA Journal on Uncertainty Quantification*, 4(1):1060–1083, 2016.

Gregor Stiglic, Primož Kocbek, Nino Fijacko, Marinka Zitnik, Katrien Verbert, and Leona Cilar. Interpretability of machine learning-based prediction models in healthcare. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 10(5):e1379, 2020.

Brian Williamson and Jean Feng. Efficient nonparametric statistical inference on population feature importance using Shapley values. In *International Conference on Machine Learning*, pages 10282–10291. PMLR, 2020.

Jilei Yang. Fast Treeshap: Accelerating SHAP value computation for trees. *Advances in Neural Information Processing Systems*, 34, 2021.

Fast Calculation of Feature Contributions in Boosting Trees (Supplementary Material)

Zhongli Jiang¹

Min Zhang¹

Dabao Zhang¹

¹Department of Epidemiology & Biostatistics, University of California, Irvine, CA, USA

A PROOFS

We first establish the following lemma.

Lemma 1

$$\sum_{k=0}^{p-n} \binom{|F|+k}{k} \binom{p-1-|F|-k}{p-n-k} = \binom{p}{n}.$$

Proof. Using Gould's identity [Gould, 1972], we have

$$\begin{aligned} & \sum_{k=0}^{p-n} \binom{|F|+k}{k} \binom{p-1-|F|-k}{p-n-k} \\ &= \sum_{k=0}^{p-n} \binom{k}{k} \binom{p-k-1}{p-n-k} \\ &= \sum_{k=0}^{p-n} \binom{p-k-1}{n-1} \\ &= \binom{p}{n}, \end{aligned}$$

where we used the Hockey-Stick Identity in the last step.

Proof of Proposition 1. Through expansion and Lemma 1, we have

$$\begin{aligned} & \sum_{k=0}^{p-n} \frac{\binom{p-n}{k} \binom{n-1}{|F|}}{p \binom{p-1}{|F|+k}} \\ &= \frac{1}{p} \sum_{k=0}^{p-n} \frac{(|F|+k)!(p-n)!(p-1-|F|-k)!(n-1)!}{(p-1)!(p-n-k)!k!|F|!(n-1-|F|)!} \\ &= \frac{1}{p} \frac{(p-n)!(n-1)!}{(p-1)!} \sum_{k=0}^{p-n} \binom{|F|+k}{k} \binom{p-1-|F|-k}{p-n-k} \\ &= \frac{1}{p} \frac{(p-n)!(n-1)!}{(p-1)!} \binom{p-1}{n-1} \frac{p}{n} \\ &= \frac{1}{n}. \end{aligned}$$

Proof of Proposition 2. We first rewrite the two polynomials

$$\begin{cases} C(z) = V(z)c, \\ P(z) = V(z)a, \end{cases}$$

where $V(z)$ is the Vandermonde matrix for vector z , and c and a are the coefficients of polynomials $C(z)$ and $P(z)$, respectively. Then the inner product

$$\begin{aligned} C(z) \cdot P(z) &= P(z) \cdot C(z) \\ &= a^T c \\ &= (V(z)^{-1} P(z))^T c \\ &= P(z)^T (V(z)^T)^{-1} c. \end{aligned}$$

Letting

$$z = \omega,$$

and noting that the Vandermonde matrix evaluated at ω is symmetric, we have

$$(V(\omega)^T)^{-1} = V(\omega)^{-1} = \frac{1}{n} V(\omega^{-1}),$$

whose multiplication with c is just the Inverse Fast Fourier transformation (IFFT) over c [Geddes et al., 1992]. Hence the proposition holds.

B THEORETICAL R^2 VALUES IN SIMULATED MODELS

The theoretical total and feature-specific R^2 in the three models are shown in Table 6.

Table 6: Theoretical R^2 in Simulated Models

Model	σ_ϵ	R^2			
		Total	X_1	X_2	X_3
a	0.50	0.9864	0.2094	0.2863	0.4907
	1.00	0.9477	0.2012	0.2750	0.4715
	1.50	0.8894	0.1888	0.2581	0.4425
b	0.50	0.9860	0.4390	0.1341	0.4129
	1.00	0.9459	0.4212	0.1286	0.3961
	1.50	0.8860	0.3945	0.1205	0.3710
c	0.50	0.9868	0.4288	0.1450	0.4130
	1.00	0.9491	0.4124	0.1395	0.3972
	1.50	0.8925	0.3878	0.1312	0.3735

C ADDITIONAL SIMULATION RESULTS

C.1 BOXPLOTS OF THE FIRST THREE FEATURE-SPECIFIC AND TOTAL R^2 VALUES

We have compared the performance of three different methods, i.e., our proposed Q-SHAP, SAGE by Covert et al. [2020], and SPVIM by Williamson and Feng [2020], in calculating the feature-specific R^2 as well as the sum of all feature-specific R^2 for the three models specified in Section 5, with different settings, i.e., $n \in \{500, 1000, 2000, 5000\}$, $p \in \{100, 500\}$, and $\sigma_\epsilon \in \{0.5, 1, 1.5\}$. The results are shown in Fig. 4-11. Note that the results of SAGE are unavailable in Fig. 6-11 because it cannot report those R^2 with our limited computational resources, and the results of SPVIM are unavailable in Fig. 8-11 because it demands too much time to complete the computation when $p = 500$.

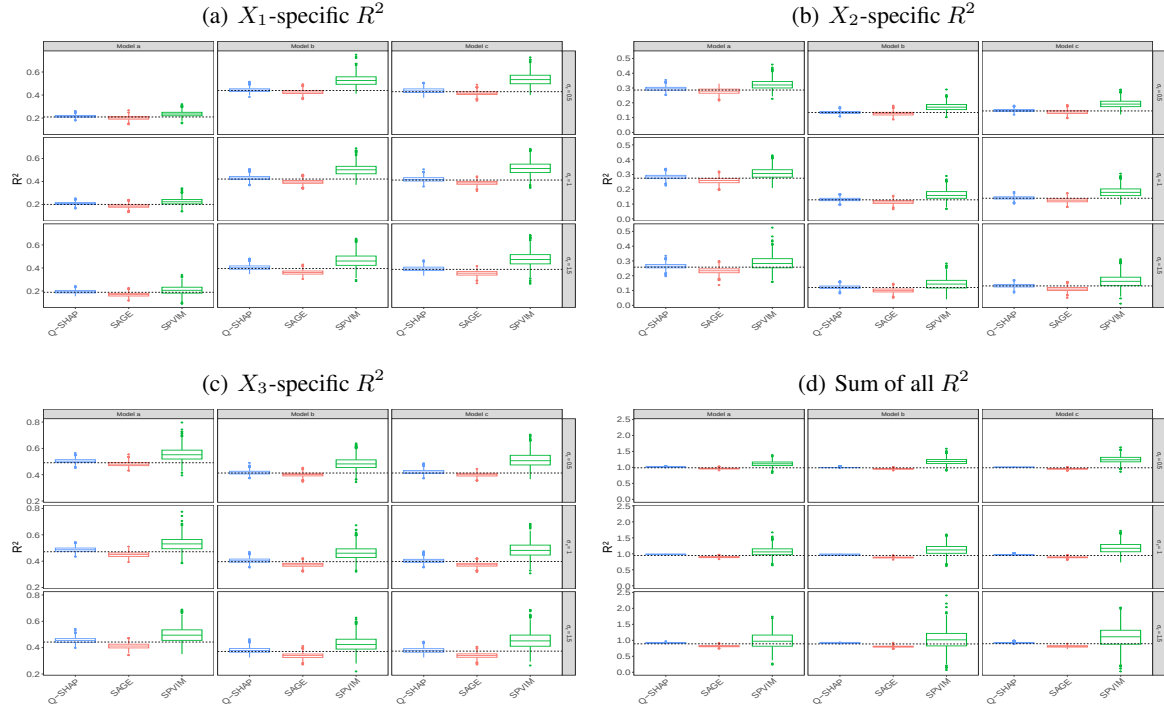


Figure 4: Boxplots of (a) X_1 -specific, (b) X_2 -specific, (c) X_3 -specific, and (d) the sum of all feature-specific R^2 in the three models with $n = 500, p = 100$. The dashed lines show the theoretical R^2 .

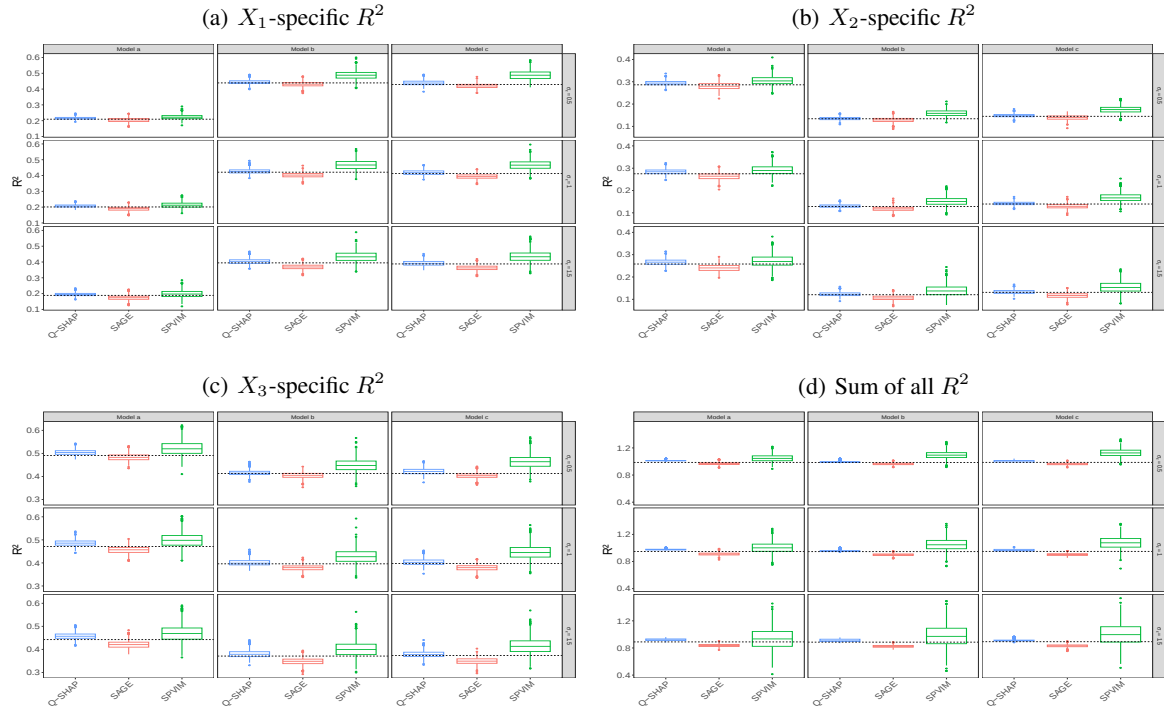


Figure 5: Boxplots of (a) X_1 -specific, (b) X_2 -specific, (c) X_3 -specific, and (d) the sum of all feature-specific R^2 in the three models with $n = 1000, p = 100$. The dashed lines show the theoretical R^2 .

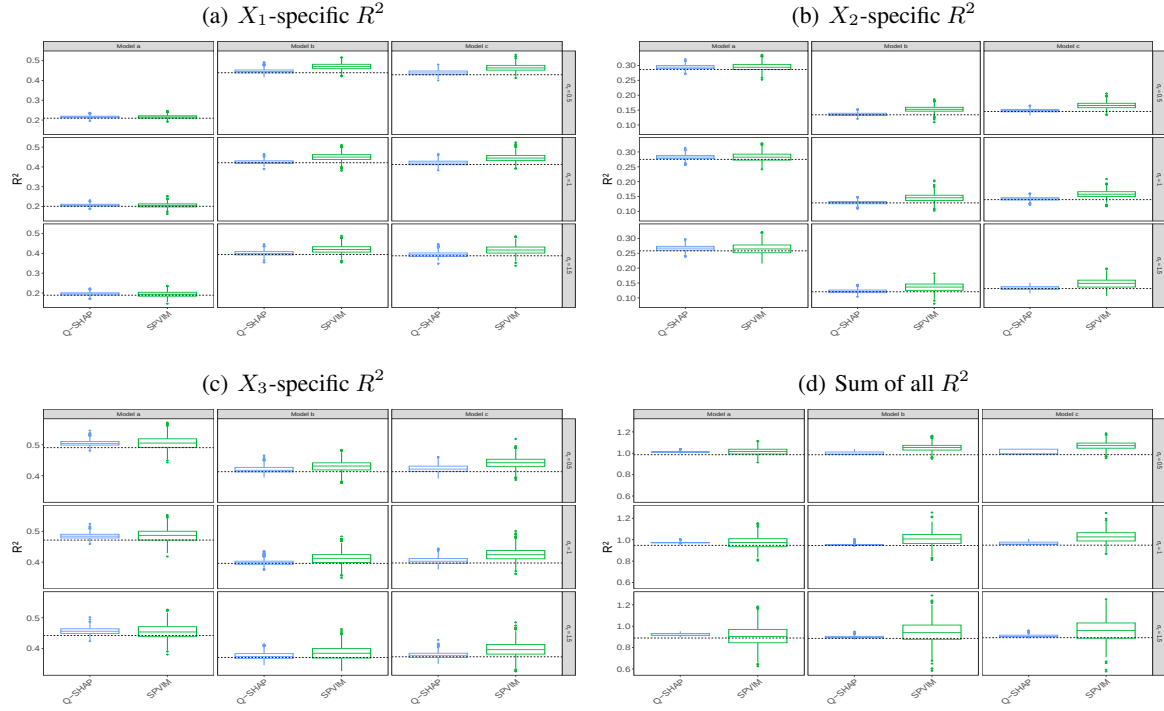


Figure 6: Boxplots of (a) X_1 -specific, (b) X_2 -specific, (c) X_3 -specific, and (d) the sum of all feature-specific R^2 in the three models with $n = 2000$, $p = 100$. The dashed lines show the theoretical R^2 .

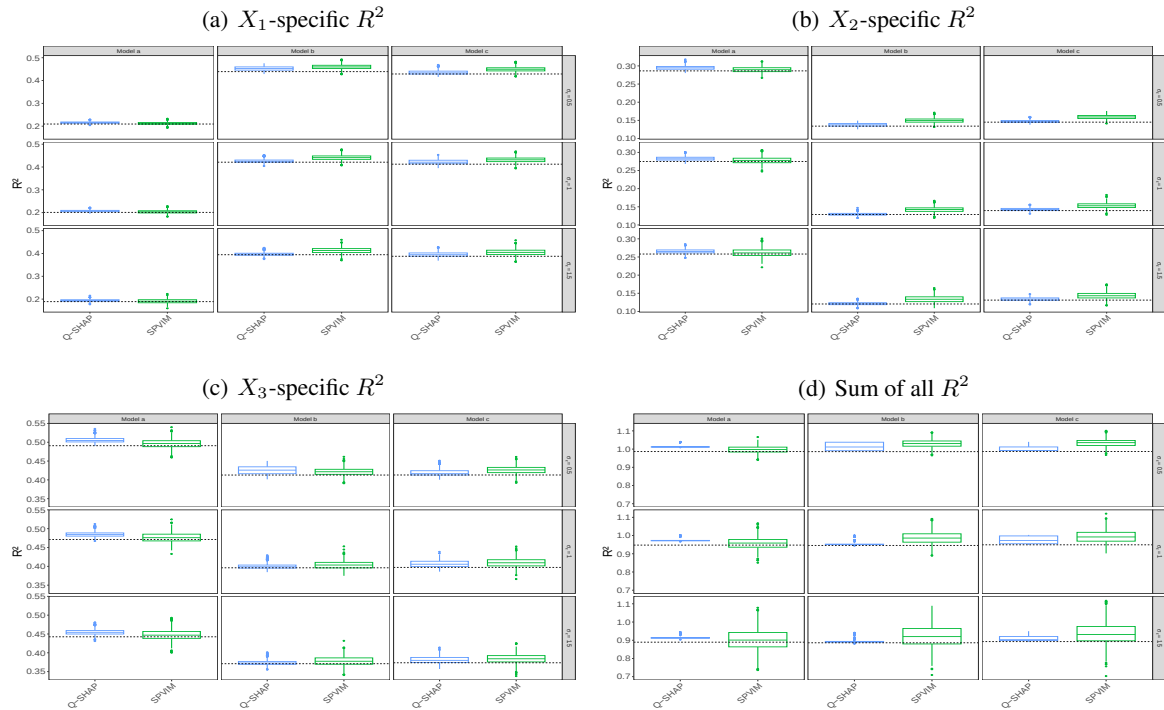


Figure 7: Boxplots of (a) X_1 -specific, (b) X_2 -specific, (c) X_3 -specific, and (d) the sum of all feature-specific R^2 in the three models with $n = 5000$, $p = 100$. The dashed lines show the theoretical R^2 .

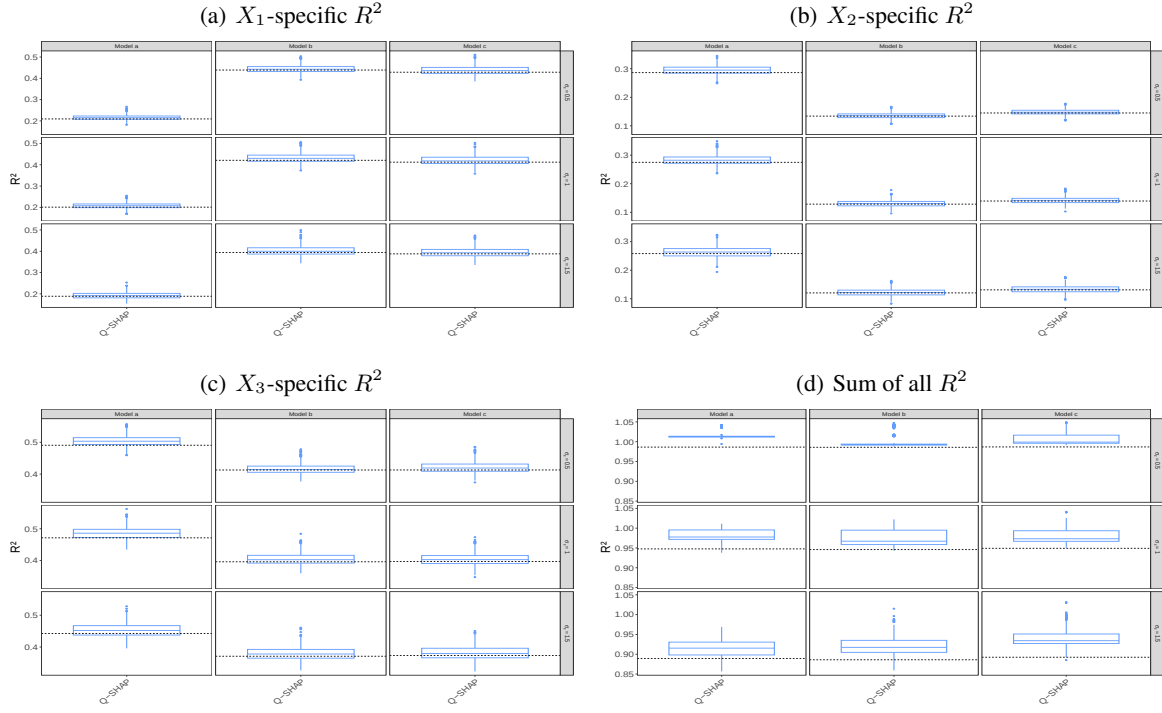


Figure 8: Boxplots of (a) X_1 -specific, (b) X_2 -specific, (c) X_3 -specific, and (d) the sum of all feature-specific R^2 in the three models with $n = 500$, $p = 500$. The dashed lines show the theoretical R^2 .

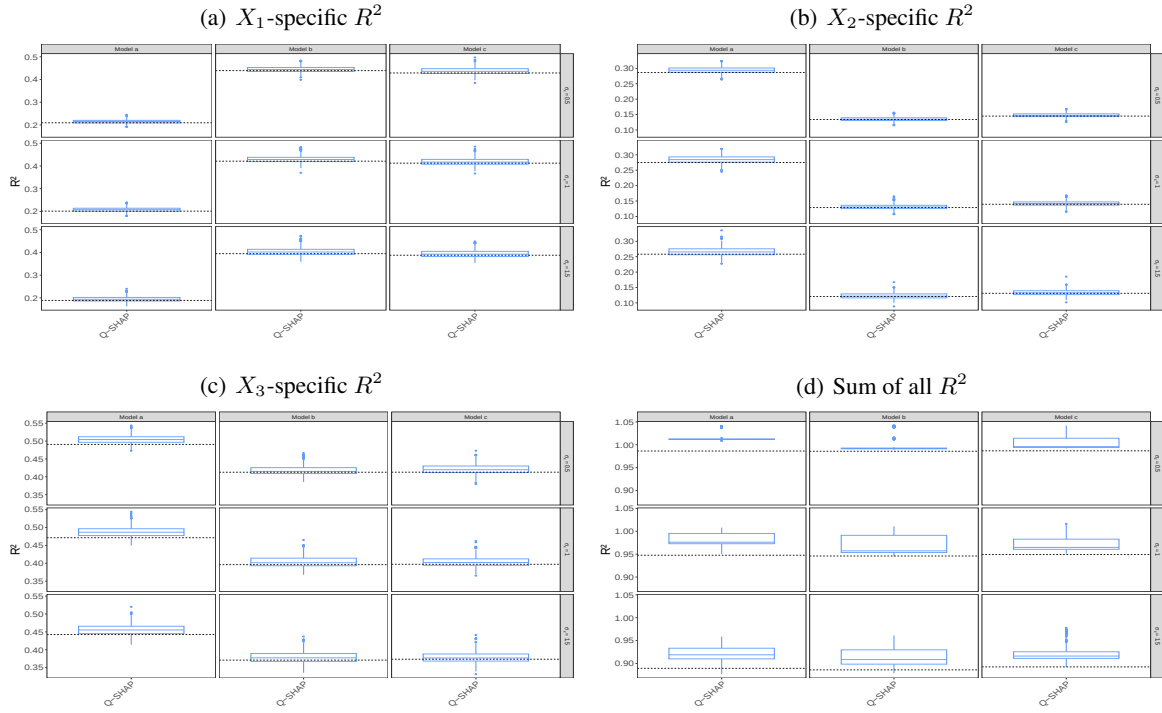


Figure 9: Boxplots of (a) X_1 -specific, (b) X_2 -specific, (c) X_3 -specific, and (d) the sum of all feature-specific R^2 in the three models with $n = 1000$, $p = 500$. The dashed lines show the theoretical R^2 .

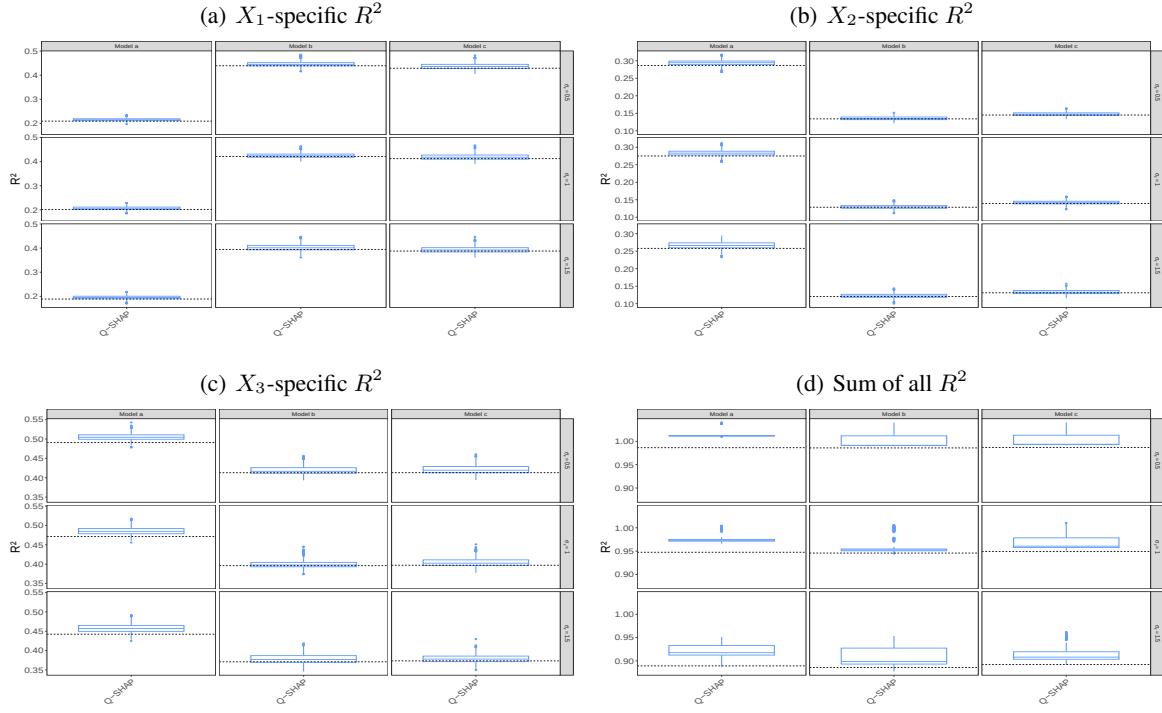


Figure 10: Boxplots of (a) X_1 -specific, (b) X_2 -specific, (c) X_3 -specific, and (d) the sum of all feature-specific R^2 in the three models with $n = 2000$, $p = 500$. The dashed lines show the theoretical R^2 .

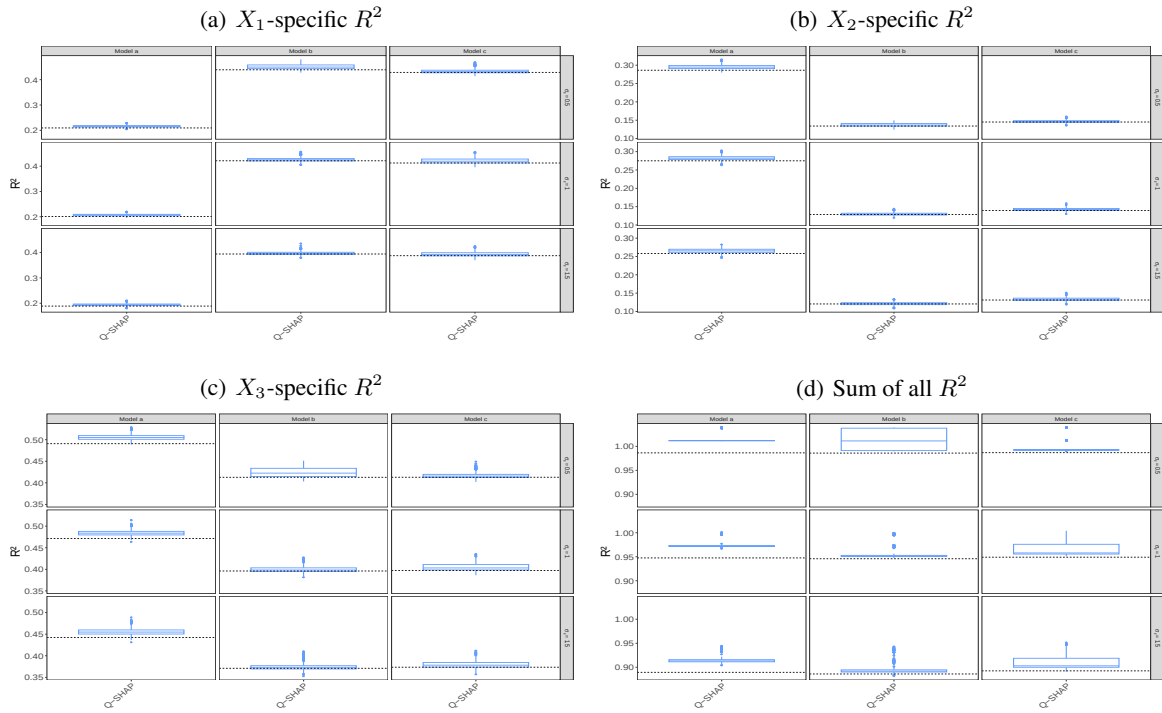


Figure 11: Boxplots of (a) X_1 -specific, (b) X_2 -specific, (c) X_3 -specific, and (d) the sum of all feature-specific R^2 in the three models with $n = 5000$, $p = 500$. The dashed lines show the theoretical R^2 .

C.2 PLOTS OF THE MEAN ABSOLUTE ERROR (MAE)

Similar to Fig. 2, we show in Fig. 12 the mean absolute error (MAE) of feature-specific R^2 for both signal and nuisance features averaged over 1,000 datasets when $p = 500$. Note that the results of SAGE and SPVIM are unavailable because none of them can complete the computation for $p = 500$ with limited computational resources.

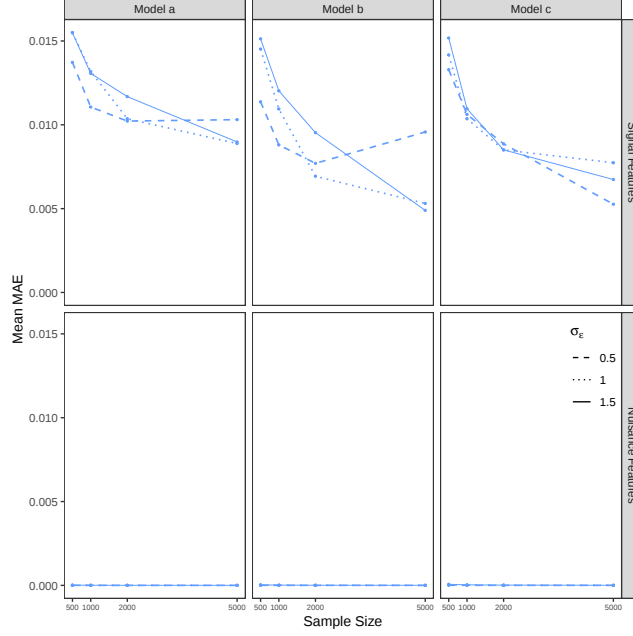


Figure 12: The mean of absolute error (MAE) of the feature-specific R^2 by Q-SHAP averaged across 1,000 datasets with $p = 500$

D COMPLEXITY OF THE ALGORITHM

Here we assume that the dataset includes n samples as well as p features, and a total of T trees are constructed with the maximum tree depth D and maximum tree leaves L . We denote S the number of permutations taken in SAGE with $S = 10^{20}$ by default. Then, when the trees are constructed by XGBoost, the complexity of SAGE is $O(TDSpn)$ [Covert et al., 2020], and the complexity of SPVIM is $O(TDpn^2 \log n)$ [Williamson and Feng, 2020]. Instead, the complexity of Q-SHAP is $O(TL^2D^2n)$ which doesn't rely on the number of features p .

Let's first consider the complexity of Q-SHAP in Algorithm 1 for a single tree and one sample. As shown in Algorithm 1, the two outer loops that iterate through the tree leaves, result in a complexity of $O(L^2)$. Within the inner loop, the computation involves the number of features induced by each pair of leaves, leading to $O(D)$ operations. The evaluation of $t[j]$ involves the computation of $C(z) \cdot P(z)$, which takes $O(D)$ operations since the number of union features between two leaves is bounded by $2D$. Combining these, the overall complexity for one tree and one sample is $O(L^2D^2)$. Thus, for the whole dataset, the complexity of Q-SHAP scales to $O(nL^2D^2)$ for a single tree. With the advancements introduced in Section 4, Q-SHAP has a total complexity of $O(TnL^2D^2)$ for the ensemble of T boosting trees.

The property that the complexity of Q-SHAP doesn't rely on the number of features is a prominent advantage of Q-SHAP and is critical in analyzing high-dimensional data. Such an advantage is achieved via Proposition 1, which eliminates dependence on p by leveraging the internal structure of the tree. Furthermore, unlike SAGE and SPVIM, which require extensive sampling, Q-SHAP directly utilizes the tree's weight function introduced in Section 3.2, eliminating the need for any sampling.